



Revista Internacional de Ciencias Sociales y Humanidades, SOCIOTAM

ISSN: 1405-3543

hmcappello@yahoo.com

Universidad Autónoma de Tamaulipas
México

HINOJOSA RIVERO, Guillermo

El tratamiento estadístico de las redes semánticas naturales

Revista Internacional de Ciencias Sociales y Humanidades, SOCIOTAM, vol. XVIII, núm. 1, -, 2008,
pp. 133-154

Universidad Autónoma de Tamaulipas
Ciudad Victoria, México

Disponible en: <http://www.redalyc.org/articulo.oa?id=65411190007>

- Cómo citar el artículo
- Número completo
- Más información del artículo
- Página de la revista en redalyc.org

redalyc.org

Sistema de Información Científica

Red de Revistas Científicas de América Latina, el Caribe, España y Portugal

Proyecto académico sin fines de lucro, desarrollado bajo la iniciativa de acceso abierto

EL TRATAMIENTO ESTADÍSTICO DE LAS REDES SEMÁNTICAS NATURALES

Guillermo HINOJOSA RIVERO
Universidad Iberoamericana Puebla, México

RESUMEN

La técnica llamada "Redes Semánticas Naturales" es una herramienta útil para el estudio de los significados que tienen ciertas palabras o expresiones en un grupo social determinado. En este artículo proponemos una serie de técnicas cuantitativas para el manejo de los datos que arroja un estudio de redes semánticas.

Primero, se hace una crítica del manejo cuantitativo tradicional en los estudios de Redes Semánticas Naturales. Luego, se proponen nuevas maneras para determinar los conjuntos SAM y para calcular diversos índices que permitirán comparar grupos diferentes en investigaciones realizadas por diferentes autores en aspectos como la homogeneidad del grupo y el peso semántico de diferentes conceptos.

Finalmente, se dan algunas sugerencias para determinar los tamaños de las muestras necesarios para hacer inferencias a poblaciones y para determinar niveles de significancia de las diferencias entre dos grupos que respondieron a los mismos estímulos.

Palabras clave: redes semánticas, estadística de redes semánticas, comparación de grupos, índices de homogeneidad de grupos.

**STATISTICAL ANALYSIS
OF NATURAL SEMANTIC NETWORKS**

ABSTRACT

The so called 'Natural semantic networks' technique is a useful tool to study the meanings of certain words or idioms within a certain social group. In this paper we propose several quantitative techniques to analyze data from semantic networks studies.

First we review critically the quantitative techniques mostly used in these studies. Then, several new ways are proposed to find the SAM sets, and to calculate index numbers of group homogeneity and semantic weights to compare data coming from different studies made by different authors.

Lastly, some suggestions are given to find the appropriate sample sizes to infer population values and to compare groups responding to the same stimuli words.

Key words: Semantic networks, statistics for semantic networks, group comparing, group homogeneity index.

La difusión acelerada de la información en la “sociedad del conocimiento” implica el peligro de que todo lo que se difunda sea considerado conocimiento seguro, sólido. Los usuarios de la información tienen cada vez menos tiempo para verificar la calidad de la información que reciben. Entonces, es una responsabilidad de quienes producen información y conocimientos, asegurar la buena calidad de lo que difunden.

La información proveniente de estudios sociales a los que se les aplicó algún tratamiento estadístico corre el doble peligro de ser más creíble, puesto que implica números (y el apoyo en estadística em-

pleada inadecuadamente). En particular, siempre nos ha parecido que la herramienta de investigación social llamada “redes semánticas naturales” podría aumentar su poder de crear conocimiento al utilizar análisis estadísticos apropiados a los datos que obtiene. En lo que resta de este texto propondremos algunas técnicas estadísticas para el análisis de este tipo de redes.

LA TÉCNICA DE LAS REDES SEMÁNTICAS NATURALES

Dicha técnica de investigación social inicialmente propuesta por Figueroa, González y Solís en 1981 (citado en Valdez, 2000) es una herramienta útil para el estudio de los significados que tienen ciertas palabras o expresiones en un grupo social determinado. En teoría, dicha técnica permitiría comparar dos o más grupos de acuerdo con el significado que le asignan los grupos a ciertos conceptos claves de interés para el investigador. También se abre la posibilidad de estudiar un grupo humano de acuerdo con los significados que le asigna a varios conceptos. Para ejemplos del uso de esta técnica, vea Varela, *et al.* (2000); González y Valdez (2005). En este artículo proponemos una serie de técnicas cuantitativas para el manejo de los datos que arroja un estudio de redes semánticas.

De manera muy resumida, la técnica, explicada con detalle por Valdez (*op. cit.*) y por Vera Noriega, *et al.* (2005), consiste en lo siguiente: se seleccionan una o más palabras estímulo de las cuales se quiere saber el significado que le dan los sujetos miembros de algún grupo en particular. Se les pide que definan la palabra estímulo mediante un mínimo de cinco palabras sueltas, que pueden ser verbos, adverbios, adjetivos, sustantivos, nombres o pronombres, sin utilizar artículos ni proposiciones. Cuando los sujetos han hecho su lista de palabras definidoras se les pide que, de manera individual, las jerarquicen de acuerdo con la cercanía o importancia que tiene cada una de las palabras con la palabra estímulo. Le asignarán el número uno a la palabra más cercana al estímulo, el dos a la siguiente, y así sucesivamente, hasta agotar todas las palabras definidoras.

El vaciado de datos y el cálculo de resultados se hace de la siguiente manera: se hace una relación de todas las palabras dichas

por los sujetos. Cada palabra tiene a su derecha diez casillas para registrar la jerarquía que cada persona le dio. La tabla siguiente aclara este procedimiento.

Tabla 1. Ejemplo del vaciado de datos con las dos primeras palabras usadas para definir la palabra estímulo, manzana.

Estímulo	Jerarquía									
<i>Manzana</i>	1	2	3	4	5	6	7	8	9	10
<i>Fruta</i>										
<i>Roja</i>										

La *Tabla 1* muestra un caso ficticio en el que la palabra *manzana* sirve de estímulo y los miembros del grupo responden con las palabras *fruta* y *roja*, entre otras. La tabla indica que la palabra *fruta* fue colocada en primer lugar por tres personas y en segundo lugar por dos personas; una persona la puso en cuarto lugar; una más en quinto, etc. La palabra *roja* fue puesta en primer lugar por dos personas, en segundo lugar por tres personas, etc.

Después de construir una tabla con todas las palabras definidoras que fueron usadas, se calculan los dos principales valores de la red semántica:

- A) El valor J, que es el total de palabras definidoras generadas por los sujetos. De acuerdo con Valdez (*op. cit.*), este valor es un indicador de la riqueza semántica de la red: a mayor cantidad de palabras, mayor riqueza semántica.
- B) El valor M, que indicaría el peso semántico de cada palabra definidora.

Para calcular el valor M se multiplica la frecuencia de aparición en cada lugar de la jerarquía por el valor semántico que se le da a esa jerarquía y se suman los resultados obtenidos para las diez posiciones jerárquicas. Valdez (*op. cit.*) y otros autores asignan un valor semántico de 10 a la jerarquía 1, valor de 9 a la jerarquía 2, y así sucesivamente, hasta darle un valor semántico de 1 a la jerarquía 10.

Entonces, en el ejemplo, el valor M de la palabra *fruta* se calcula multiplicando 3 apariciones en primer lugar por 10, que es el valor semántico de la jerarquía 1, lo que da 30, más 2 por 9 del segundo lugar (igual a 18), más 1 por 7 del cuarto lugar (=7), más 1 por 6 del cuarto lugar (=6), más 2 por 5 (=10), más 1 por 3 (=3), más 1 por 1 (=1), lo que da un total de 75. El mismo cálculo para la palabra *roja* da 68.

Tabla 2. Cálculo de los valores M para el ejemplo de la *Tabla 1*.
En el renglón superior se colocaron los pesos semánticos asignados a cada posición jerárquica. En la columna de la extrema derecha se pusieron los valores M calculados.

Valor semántico	10	9	8	7	6	5	4	3	2	1	Valor M
Estímulo	Jerarquía										
<i>Manzana</i>	1	2	3	4	5	6	7	8	9	10	
<i>Fruta</i>											75
<i>Roja</i>											68

Es importante entender cómo se calculan los valores M y la racionalidad tras la asignación de los valores semánticos a cada una de las posiciones jerárquicas. Valdez (*op. cit.*) razona que las palabras puestas en primer lugar son las que más peso deben tener al calcular. Al considerar un máximo de 10 posiciones jerárquicas, simplemente invierte los pesos. Si las jerarquías van de 1 a 10, los pesos respectivos asignados van de 10 a 1. Con esto da más peso a los primeros lugares y menos a los últimos.

Después de calcular el valor M de todas las palabras, se toman las diez palabras definidoras con mayor valor M y se considera que este conjunto de palabras refleja el significado que el grupo le da a la palabra estímulo. A este conjunto se le llama conjunto SAM (*Semantic Association Memory*). Los conjuntos SAM son la base para comparaciones y correlaciones posteriores intra e inter grupos.

PROBLEMAS ESTADÍSTICOS DE LA TÉCNICA

Desde el punto de vista estadístico, los valores J y M tienen algunas propiedades que los hacen poco sólidos como medidas psicométricas. El valor J, que se interpreta comúnmente como “la riqueza de la red”, es un valor poco robusto porque depende de dos condiciones muy variables: A) El tamaño del grupo de sujetos. Si el grupo crece, el valor J crece también hasta, quizá, alcanzar un valor asintótico. Pero no se puede saber de qué tamaño debe ser el grupo para que J alcance el valor asintótico que podría utilizarse como “riqueza de la red”. B) Basta con la presencia de un sujeto que dé respuestas atípicas para que J aumente. Estos inconvenientes de J se reflejan en el poco uso que se le da a este valor en los trabajos que utilizan esta técnica (pero vea Valdez, 1996, para un ejemplo de inferencia estadística a partir de valores J). Isabel Reyes (1993) propone que el valor J se llame simplemente “tamaño de la red”, ya que no está claro qué refleja dicho valor.

El cálculo del valor M tiene tres fuentes de problemas estadísticos: A) La arbitrariedad de asignar un valor semántico de 10 a la jerarquía 1, ¿por qué 10 y no 8 o 100? B) La incorrección de manejar valores ordinales, la jerarquización de las palabras, como si fueran valores intervalares que se pueden multiplicar y sumar. C) Depende del número de sujetos.

Veamos primero el problema del 10 arbitrario. Cuando uno asigna el valor semántico de 10 a las menciones en primer lugar y, como es lo usual, disminuye el peso semántico de uno en uno para las menciones en segundo, tercero, hasta décimo lugar, significa que se está dispuesto a aceptar ciertas relaciones matemáticas curiosas. Por ejemplo, aceptar que dos menciones en sexto lugar, con peso semántico de 5, equivalen a una mención en primer lugar (2 por 5 equivale a 1 por 10).

El problema real es que el orden de las palabras —y por tanto el conjunto SAM— depende del valor arbitrario que se dé a los pesos semánticos. Un ejemplo sencillo mostrará esta dependencia. Supóngase que al definir *manzana*, la palabra *jugosa* tuvo una mención en

primer lugar y la palabra *dura* tuvo dos menciones en quinto lugar. Al hacer el cálculo convencional, empezar con el 10 para el primer lugar y disminuir de uno en uno, *jugosa* tendría un valor M de 10 (1 por 10) y *dura* tendría un valor M de 12 (2 por 6). Al ordenar las palabras, *dura* quedaría arriba de *jugosa*. Por otro lado, un investigador puede decidir que, dado que los sujetos rara vez ponen más de cinco palabras definidoras, se puede dar un peso semántico de cinco a la mención en primer lugar y disminuir de uno en uno hasta dar un peso de uno a las menciones en quinto lugar. En tal caso, la palabra *jugosa* tendría un valor M de 5 (1 por 5) y la palabra *dura* tendría valor de 2 (2 por 1); *jugosa* quedaría ahora por arriba de *dura*. Vea la tesis de C. Cuétara (2004) para un ejemplo de pesos semánticos de 5 a 1.

En el ejemplo anterior vale la pena resaltar que el orden de las palabras no dependió de los datos, sino de una decisión del investigador. Aunque se adopte la convención de que todos los investigadores usen siempre la misma escala, quedará la duda acerca de qué tan bien refleja la realidad.

La segunda dificultad del valor M deriva de tratar valores ordinales como si fueran intervalares o de razón. Si las posiciones jerárquicas, que son datos ordinales, se multiplican por un escalar, el peso semántico, el resultado sigue siendo ordinal, en el mejor de los casos. Ahora bien, si los datos ordinales se suman, el resultado es un número sin sentido. Es como sumar dos medianos más un chico, pensando que el resultado es un grande y otro poquito. Ordenar las palabras definidoras según su valor M —que es una suma de ordinales— puede ser sólo un reflejo muy deformado de la red de significados del grupo. Por la misma razón, usar técnicas estadísticas como la correlación de Pearson, la prueba t o el análisis de varianza con valores M es un ejercicio con valor estadístico dudoso.

Finalmente, los valores M dependen del tamaño del grupo. Si el grupo crece indefinidamente, los valores M crecerán también indefinidamente. Aún salvando las dificultades mencionadas arriba, no es posible comparar los valores M obtenidos en una investigación con los obtenidos en otra, a menos que los grupos sean del mismo tamaño.

Estas dificultades parecen estar implícitamente reconocidas en las investigaciones hechas con la técnica de redes semánticas naturales, ya que los investigadores procuran igualar el tamaño de los grupos (Valdez, 1996) o hacen poco uso de los valores J y M después de haber definido el conjunto SAM (Iuit y Castillo, 1995). Más bien las conclusiones se basan en análisis cualitativos de los conjuntos SAM.

Reyes (1993) critica la manera usual de definir el conjunto SAM (seleccionar las diez palabras con mayores pesos semánticos) con base en tres consideraciones: A) Se desconoce el tamaño real de la red y, por tanto, el de la muestra que lo representa; B) El número arbitrario de diez palabras puede excluir palabras definidoras importantes (podemos añadir que también puede incluir palabras definidoras de poca importancia); C) Falta un sustento teórico para la delimitación del conjunto. Más adelante propondremos una manera de definir el conjunto SAM de forma que supera, creemos, esas dificultades.

En las investigaciones que utilizan los valores M, parece ser que el único uso que se les da es ordenar las palabras definidoras para obtener el conjunto SAM. A pesar de las dificultades estadísticas implicadas en el cálculo, los conjuntos SAM obtenidos a partir de esos valores sí pueden tener una buena coincidencia con los conjuntos que se obtienen a partir de otros métodos; como el que propondré a continuación.

SOLUCIÓN DE LOS PROBLEMAS DESCRITOS

La *Tabla 3* presenta la parte más representativa del vaciado de datos para las palabras definidoras mencionadas ante el estímulo “*Ser mujer*” por un grupo de 24 mujeres de nuevo ingreso a la universidad (datos de C. Cuétara, 2004). La primera columna lista las primeras 21 palabras del total de 59 palabras que fueron mencionadas ($J=59$). No se muestran 38 palabras que sólo fueron mencionadas una vez. El máximo de palabras dadas por las personas fue 5, por lo que sólo se consideran 5 posiciones jerárquicas indicadas en el primer renglón. En el segundo renglón se muestran los pesos semánticos que se asignan a cada posición si se hace de la manera usual,

dando un peso de 10 a la primera posición. El tercer renglón muestra los pesos semánticos que dio Cuétara, asignando 5 a la primera posición. De las columnas 1 a 5 se ven las frecuencias con que cada palabra fue mencionada en cada una de las 5 posiciones jerárquicas. Por ejemplo, la palabra *femenina* fue mencionada por 6 personas en la primera posición, por 1 en la cuarta y por 2 en la quinta. La columna M (5) muestra los valores M que se obtienen con los pesos semánticos de 5 para la primera posición. La columna M (10) muestra esos mismo valores con los pesos semánticos que inician en 10. La columna Frecuencia muestra el número de veces que cada palabra fue mencionada, sin tomar en cuenta la posición. Finalmente, la columna Md muestra las medianas de las posiciones jerárquicas que le asignaron a cada palabra quienes la mencionaron. La mediana es la medida de tendencia central apropiada para datos ordinales, y se determina con la posición del dato que divide por la mitad al conjunto de datos.

Tabla 3. Vaciado parcial de datos de las respuestas dadas ante el estímulo *Ser Mujer* por un grupo de 24 mujeres, estudiantes de nuevo ingreso en la universidad. Ver el texto para una explicación del contenido de las columnas y del origen de los datos.

Posición	1°	2°	3°	4°	5°				
Peso 10	10	9	8	7	6				
Peso 5	5	4	3	2	1				
Definidora						M (5)	M (10)	Frecuencia	Md
Femenina	6			1	2	34	79	9	1
Amiga	1	2	2	2	2	25	70	9	3
Amor	4		1			23	48	5	1
Inteligencia	1	4				21	46	5	2
Madre	2	1		2		18	43	5	2
Comprensión	1		2	2		15	40	5	3
Fuerte	1	1		2	1	14	39	5	4
Responsable		2	1	1		13	33	4	2.5
Madurez		1	1		2	9	29	4	4
Respeto	1	2				13	28	3	2

cont.

Tabla 3. Vaciado parcial de datos de las respuestas dadas ante el estímulo *Ser Mujer* por un grupo de 24 mujeres, estudiantes de nuevo ingreso en la universidad. Ver el texto para una explicación del contenido de las columnas y del origen de los datos (cont.).

Cuidado	2	1				11	26	3	2
Buena			3			9	24	3	3
Casarte			3			9	24	3	3
Integridad	2					10	20	2	1
Posición	1°	2°	3°	4°	5°				
Peso 10	10	9	8	7	6				
Peso 5	5	4	3	2	1				
Definidora						M (5)	M (10)	Frecuencia	Md
Sensible	1	1				9	19	2	1.5
Cariño	1			1		7	17	2	2.5
Confianza		1	1			7	17	2	2.5
Capaz			1	1		5	15	2	3.5
Sincera			1	1		5	15	2	3.5
Trabajadora		1			1	5	15	2	3.5
Obligaciones					2	2	12	2	5

La *Tabla 3* fue ordenada de acuerdo con dos criterios: frecuencia decreciente y, en caso de empate, mediana creciente. Este ordenamiento coincide totalmente con el que se obtendría con los valores *M* (10) decrecientes, pero muestra algunas discrepancias con el ordenamiento según los valores *M* (5) (vea los valores *M* de las palabras *madurez* a *integridad* en la parte media de la tabla). Se puede demostrar matemáticamente que mientras mayor es el peso semántico asignado a la primera posición, disminuyendo de uno en uno para las posiciones subsecuentes, mayor coincidencia habrá entre los ordenamientos obtenidos con base en las frecuencias y los obtenidos con base en los valores *M*. Esta coincidencia se alcanzará más rápido mientras menos sujetos haya. De modo que, en la coincidencia obtenida en este caso, de 24 sujetos puede no encontrarse con grupos de 100 sujetos o más.

PROPUESTA

La propuesta central de este artículo es que en las investigaciones que utilizan redes semánticas naturales se ordenen las palabras con base en el criterio de frecuencia decreciente en primer lugar y los empates se resuelvan con el criterio de la mediana creciente. Este ordenamiento tiene varias ventajas sobre el que se obtiene con los valores M y ninguna de sus desventajas:

- Se elimina la arbitrariedad de los pesos semánticos que pueden producir resultados variables, que dependen del número de sujetos y del peso asignado a la primera posición.
- No se cometen incorrecciones estadísticas al multiplicar y sumar datos ordinales.
- Las medianas no se afectan al variar el número de posiciones jerárquicas que resultan de las respuestas de los sujetos.
- Se puede trabajar con cinco o diez, o con el máximo número de palabras dadas por los sujetos, y la mediana no tendrá variaciones.
- Las frecuencias se pueden convertir a porcentajes y con esto permitir las comparaciones entre grupos de diferentes tamaños medidos en diferentes momentos y lugares.
- Las frecuencias y porcentajes son datos en escala de razón, por lo que se pueden analizar con las técnicas estadísticas más poderosas.
- Las frecuencias de dos palabras diferentes son directamente comparables entre sí, pudiéndose calcular razones y proporciones entre ambas frecuencias.
- Finalmente, el ordenamiento basado en las frecuencias de las palabras definidoras permite encontrar criterios para delimitar los conjuntos SAM.

Tabla 4. Frecuencias y porcentajes de las 21 palabras que tuvieron frecuencia mayor o igual a 2.
Es el mismo conjunto de datos que se muestra en la Tabla 3.

Definidora	Frecuencia	Porcentaje
Femenina	9	37.5
Amiga	9	37.5
Amor	5	20.8
Inteligencia	5	20.8
Madre	5	20.8
Comprensión	5	20.8
Fuerte	5	20.8
Responsable	4	16.7
Madurez	4	16.7
Respeto	3	12.5
Cuidado	3	12.5
Buena	3	12.5
Casarte	3	12.5
Integridad	2	8.3
Sensible	2	8.3
Cariño	2	8.3
Confianza	2	8.3
Capaz	2	8.3
Sincera	2	8.3
Trabajadora	2	8.3
Obligaciones	2	8.3

La *Tabla 4* presenta las mismas palabras que la *Tabla 3*, con sus frecuencias respectivas y el porcentaje de sujetos que mencionaron cada palabra. Este porcentaje se calcula dividiendo la frecuencia entre el número de sujetos, que en este caso fueron 24, y multiplicando

el resultado por 100. Se puede ver que las dos palabras con mayor frecuencia fueron mencionadas por el 37.5% de los sujetos.

En la tabla se omitieron 38 palabras definidoras cuya frecuencia fue 1, lo cual equivale a un porcentaje de 4.1. Esta tabla proporciona una base para seleccionar el conjunto SAM. El investigador puede decidir incluir en el conjunto SAM todas aquellas palabras con porcentaje mayor a 20%. En tal caso, su conjunto estaría constituido por 7 palabras. Si decide incluir todas las palabras con porcentaje mayor a 15%, tendría 9 palabras. Es irremediable que al aumentar el tamaño del conjunto disminuya su representatividad. Si se quiere aumentar la representatividad del conjunto seleccionado, el número de palabras disminuirá. (Esto es congruente con el principio de lógica clásica, según el cual, “si se aumenta la extensión, disminuye la intensidad, y viceversa”, lo que se puede decir de todos y cada uno de los miembros de un grupo, disminuye cuando aumenta el tamaño del grupo).

REPRESENTATIVIDAD DE LAS MUESTRAS

En el párrafo anterior se introdujo el concepto de representatividad como si fuera equivalente al porcentaje de personas que mencionaron una palabra. Si bien representatividad y porcentaje no son equivalentes, la representatividad de un conjunto SAM está en función de los porcentajes de las palabras incluidas en el conjunto. Un índice de representatividad como el que sugiere Reyes (1993) con el nombre de *Índice de Consenso Grupal*, puede definirse a partir de los porcentajes. Una posibilidad de dicho índice es el promedio de los porcentajes de las 10 palabras con mayor frecuencia. Otra posibilidad es simplemente el porcentaje de la palabra con mayor frecuencia.

Para asegurar que el grupo examinado es representativo de una población y poder generalizar los resultados, es necesario apegarse a los lineamientos de la teoría del muestreo. La forma de seleccionar al grupo es el requisito más importante para hacer inferencias válidas. El tamaño mínimo de la muestra puede determinarse con ayuda de la *Tabla 5*, compuesta a partir de las ecuaciones de muestreo (Triola, 2004). La columna de la derecha indica el rango de precisión

dentro del cual se encuentra la proporción “real” en que aparecería una palabra definidora dada en la población. Solamente se calcularon cinco rangos: de $\pm 1\%$ a $\pm 5\%$.

Las tres columnas indican la confianza que puede tenerse en que la proporción 'real' se encuentre precisamente en un rango dado. Es claro que, a mayor precisión en la predicción, menor es el rango de variación y el tamaño de la muestra debe ser mayor. También, para lograr mayor confianza en la predicción, se requiere un mayor tamaño de la muestra. Por ejemplo, si se selecciona una muestra de 601 sujetos, la precisión de la predicción será de $\pm 4\%$ con una confianza de 95%. Se asume que el tamaño de la población es muy grande, como la población de una ciudad o de un país. Para poblaciones limitadas, como la población de una universidad, deben calcularse los tamaños de las muestras en cada caso.

Tabla 5. Tamaño mínimo necesario de una muestra representativa para obtener un rango de variación dado con grados de confianza de 90%, 95% y 99%. La tabla se calculó asumiendo la varianza máxima para proporciones ($p=q=0.5$; varianza=0.25). Otros valores de p y q resultarán en menores tamaños de muestra.

Rango de variación	Confianza		
	90%	95%	99%
$\pm 1\%$	6764	9604	16588
$\pm 2\%$	1691	2401	4147
$\pm 3\%$	752	1068	1844
$\pm 4\%$	423	601	1037
$\pm 5\%$	271	385	664

COMPARACIÓN ESTADÍSTICA DE MUESTRAS

Los porcentajes de las palabras definidoras permiten hacer comparaciones estadísticas entre dos o más grupos diferentes que definen la misma palabra, o entre dos o más palabras definidas por el mismo grupo. Las comparaciones estadísticas pueden complementar los análisis cualitativos usuales para comparar conjuntos SAM.

Para un ejemplo de análisis cualitativo puede verse el trabajo de Valdez, Oudhof y Posadas (1998).

Dos grupos diferentes pueden compararse, aun si definen palabras diferentes. El objeto de la comparación sería averiguar cuál de los dos grupos presenta un mayor consenso grupal. Una manera de hacerlo es definir los conjuntos SAM con el mismo criterio para ambos grupos. Por ejemplo, incluir en el conjunto todas las palabras con porcentajes de consenso mayores a 15%. El grupo que tenga el conjunto mayor será el que muestre más consenso grupal. Otra manera puede ser calcular el porcentaje promedio de las 10 palabras superiores (podemos llamar a este promedio como C10; mnemónico de Consenso Grupal de 10 palabras): el grupo que tenga el mayor C10 mostrará el mayor consenso.

Evaluar si la diferencia entre dos índices C10 es significativa puede hacerse mediante la aplicación de la prueba de Mann-Whitney o la de Kolmogorov-Smirnov para dos muestras independientes, comparando los 10 primeros porcentajes de cada grupo. La aplicación de la prueba *t* no parece ser una alternativa deseable, porque como se comparan los puntajes extremos, los 10 mayores, no se puede asegurar la normalidad de los datos, necesaria para esta prueba.

Es posible hacer exámenes estadísticos para evaluar si dos grupos diferentes le dan o no el mismo significado a un concepto dado. La *Tabla 6* muestra las palabras definidoras, con sus frecuencias, medianas y porcentajes para el concepto *Ser mujer* dadas por un grupo de mujeres próximas a graduarse en la universidad. En este caso, el tamaño del grupo fue 65, el total de palabras dadas fue 112 ($J=112$). En la *Tabla 6* se muestran sólo las palabras con frecuencia mayor a 2. Datos de Cuétara (2004). A manera de ejemplo, podemos comparar los datos de la *Tabla 4* y la *Tabla 6*, que contienen las respuestas dadas por dos grupos de mujeres universitarias, principiantes y próximas a graduarse, ante el mismo concepto: *Ser Mujer*.

Tabla 6. Frecuencias, medianas y porcentajes de las 49 palabras definidoras de *Ser Mujer* que tuvieron frecuencia mayor o igual a 2. Grupo de mujeres universitarias por graduarse, N=65.

Definidor	Frecuencia	Md	Porcentaje
Responsable	21	3	32.3
Amor	20	2	30.8
Inteligencia	15	1	23.1
Fortaleza	11	2	16.9
Comprensión	11	4	16.9
Femenina	9	3	13.8
Sensible	9	3	13.8
Madre	8	1.5	12.3
Trabajadora	8	3.5	12.3
Seguridad	7	2	10.8
Fiel	7	3	10.8
Tierna	7	4	10.8
Capaz	6	2	9.2
Cariño	6	2.5	9.2
Débil	6	5	9.2
Amiga	5	2	7.7
Profesional	5	2	7.7
Respeto	5	2	7.7
Compromiso	5	3	7.7
Entrega	5	3	7.7
Ahorrativa	5	5	7.7
Carrera	5	5	7.7
Independiente	4	2	6.2
Realización	4	2	6.2
Familia	4	3.5	6.2
Adaptable	4	4	6.2

cont.

Tabla 6. Frecuencias, medianas y porcentajes de las 49 palabras definidoras de Ser Mujer que tuvieron frecuencia mayor o igual a 2. Grupo de mujeres universitarias por graduarse, N=65 (cont.).

Definidor	Frecuencia	Md	Porcentaje
Belleza	4	4	6.2
Integridad	3	1	4.6
Privilegio	3	1	4.6
Esposa	3	2	4.6
Ética	3	2	4.6
Honestidad	3	2	4.6
Humana	3	3	4.6
Igualdad	3	3	4.6
Paciente	3	3	4.6
Sincera	3	3	4.6
Dedicada	3	4	4.6
Apoyo	3	5	4.6
Lucha	2	1	3.1
Alegre	2	2	3.1
Culta	2	2	3.1
Educada	2	2	3.1
Tarea	2	2	3.1
Libre	2	2.5	3.1
Creativa	2	3	3.1
Autosuficiente	2	3.5	3.1
Hogar	2	3.5	3.1
Compañía	2	4	3.1
Experiencia	2	4.5	3.1

La comparación cuantitativa entre los dos grupos puede iniciarse con los indicadores de consenso intra-grupo. Si definimos como conjunto SAM el formado por las palabras con $p > 15\%$, encontramos que el grupo de principiantes tiene un conjunto SAM ($p > 15\%$) de 9

palabras y el grupo de graduadas tiene un conjunto de 5 palabras, lo que indica un mayor consenso en el grupo de principiantes. La misma conclusión se obtiene a partir de otros puntos de corte para el conjunto SAM —20% y 10%—, pero no para 30%.

El índice C10 —frecuencia promedio de las diez palabras más frecuentes— es 24.2% para el grupo de principiantes y 18.3% para el grupo de graduadas. Nuevamente indica mayor consenso en el primer grupo. Pero tanto la prueba de Mann Whitney como la de Kolmogorov-Smirnov indican que las diferencias cuantitativas entre los 10 puntajes mayores de los dos grupos no son significativas.

Podemos hacer otros análisis más finos para comparar los conjuntos SAM de ambos grupos. Para calcular una correlación es necesario igualar los conjuntos a correlacionar. Para esto se toman las palabras de ambos conjuntos SAM y se anotan los porcentajes respectivos para las palabras en que ambos conjuntos coincidan. Para las palabras en las que no coinciden los conjuntos, se buscan los porcentajes en las tablas generales de resultados. Si una palabra que se encuentra en un conjunto SAM no aparece en la tabla general del otro grupo, se le anota un porcentaje de 0%.

La *Tabla 7* muestra, en las tres primeras columnas, las palabras de ambos conjuntos SAM ($p > 15\%$), junto con los puntajes respectivos en cada grupo. En el penúltimo renglón se muestra el coeficiente de correlación de Pearson que resultó ser negativo y no significativo, lo cual indica que los dos grupos de puntajes no están asociados: los conjuntos son diferentes. El último renglón muestra el valor del coeficiente tau de Kendall, que quizá sea el índice más apropiado para comparar conjuntos SAM. También es negativo y no significativo. Para la interpretación correcta de este índice, consúltase el libro de Siegel (1988) o el de Hays (1973).

Puesto que los dos conjuntos son diferentes, y no se puede predecir uno a partir del otro, podemos examinar más de cerca en dónde están las diferencias. Para cada una de las palabras seleccionadas se hizo una prueba de significación de diferencia de proporciones (Bruning y Kintz, 1977; Triola, 2004). Las dos últimas columnas de la *Tabla 7* muestran los resultados: las palabras *femenina*,

amiga y *madurez* tuvieron un porcentaje significativamente mayor en el grupo de principiantes. Entonces, las mayores diferencias entre ambos grupos se encuentran ligadas a los pesos que les dan a esas tres palabras.

Las mismas comparaciones entre conjuntos y entre palabras específicas de cada conjunto pueden hacerse para comparar las respuestas de un grupo ante dos conceptos-estímulo diferentes.

Tabla 7. Correlación y comparación de los porcentajes de cada palabra que aparece en los conjuntos SAM ($p>15\%$) de ambos grupos.

Palabra	Principiantes	Graduadas	Z	Significación
Femenina	37.5	13.8	2.471	$p<0.005$
Amiga	37.5	7.7	3.426	$p<0.005$
Amor	20.8	30.8	-0.931	n.s
Inteligencia	20.8	23.1	-0.230	n.s
Madre	20.8	12.3	1.008	n.s
Comprensión	20.8	16.9	0.425	n.s
Respeto	16.7	7.7	1.248	n.s
Madurez	16.7	1.5	2.768	$p<0.005$
Responsabilidad	16.7	32.3	-1.453	n.s
Fortaleza	4.2	16.9	-1.557	n.s
Correlación de Pearson	-0.185			n.s
Coefficiente tau de Kendall	-0.077			n.s

CONCLUSIÓN

La propuesta de sustituir los valores M por frecuencias y medianas permite lograr ordenamientos confiables y válidos de las palabras definidoras según sus pesos semánticos. Al convertir las frecuencias a porcentajes se pueden establecer criterios para determinar los conjuntos SAM representativos, así como para establecer índices de consenso intra-grupo; se abre la posibilidad de determi-

nar los tamaños de los grupos por examinar de acuerdo con los porcentajes de error tolerables.

También es posible comparar los resultados de grupos diferentes: en sus aspectos puramente cuantitativos con las pruebas de Mann-Whitney y Kolmogorov-Smirnov. En los aspectos de orden los grupos se pueden comparar con el índice de correlación de Pearson y la tau de Kendall. La comparación de los porcentajes de cada palabra puede hacerse con la prueba de diferencia de dos proporciones.

La estadística aporta herramientas que ayudan al razonamiento del investigador; de ninguna manera lo sustituyen. No debemos confiar en las conclusiones estadísticas, totalmente mecánicas, que nos producen la ilusión de estar investigando y razonando. Las técnicas aquí propuestas no son sustituto de los análisis cualitativos que hasta ahora se han practicado en la investigación con redes semánticas naturales. Más bien deben complementarlos, al orientar y restringir las posibilidades de interpretación y al ayudar a tomar decisiones.

NOTA

¹. El autor agradece a la Mtra. Covadonga Cuétara por facilitarle los datos de su tesis doctoral, por haber estimulado la discusión de los métodos cuantitativos empleados, y por haber alentado la escritura de este artículo.

BIBLIOGRAFÍA

- BRUNING, J.L. y KINTZ, B.L. (1977). *Computational Handbook of Statistics*, Glenview, Ill Scott, Foresman and Company.
- CUÉTARA, C. (2004). *Representaciones de género y socioprofesionales: El caso de una universidad privada*, tesis inédita de doctorado en educación, Puebla, Pue., UIA Puebla.

- GONZÁLEZ, E.S. y VALDEZ M., J.L. (2005). "Significado psicológico de la depresión en médicos y psicólogos", *Psicología y Salud*, julio-diciembre, año/Vol. 15, N° 002. Consultada en <http://redalyc.uaemex/redalyc/pdf/291/29155210.pdf> (21 de mayo 2007).
- HAYS, W.L. (1973). *Statistics for the Social Sciences*, 2ª ed., Nueva York, Holt, Rinehart, and Winston.
- IUIT, B., J.I., CASTILLO, L.T. (1995). "Conceptos de familia, padre y madre en adolescentes yucatecos: Una comparación por nivel socioeconómico", *Psicología y Salud*, 5, pp. 79-91.
- REYES L., I. (1993). "Las redes semánticas naturales, su conceptualización y su utilización en la construcción de instrumentos", *Revista de Psicología Social y Personalidad*, Vol. 9 (1), pp. 81-97.
- SIEGEL, S. (1988). *Estadística no paramétrica aplicada a las ciencias de la conducta*, 11ª reimpresión, México, Trillas.
- TRIOLA, M.F. (2004). *Estadística*, 9ª ed., México, Pearson.
- VALDEZ, M., J.L. (1996). "La evaluación del autoconcepto a través de la técnica de redes semánticas", *Revista Mexicana de Psicología*, 13(2), pp. 175-185.
- (2000). *Las redes semánticas naturales*, Toluca, Edo. de México, Universidad Autónoma del Estado de México.
- VALDEZ M., J.L.; OUDHOF, B.H. y POSADAS M., M.A. (1998). "Significado psicológico de violencia, gobierno, democracia y EZLN, en diferentes niveles de escolaridad", *Revista Mexicana de Psicología*, 15(1), pp. 11-17.
- VARELA R., M.; PETRA, I.; GONZÁLEZ, E. y PONCE DE LEÓN, M.E. (2000). "Análisis semántico del concepto de enseñanza de profesores de medicina", *Revista de la Educación Superior*, Vol. XXIX (3), N° 116, octubre-diciembre. 2000. Consultada en http://www.anuies.mx/servicios/p_anuies/publicaciones/revsup/res116/menu1.htm (18 mayo 2007).
- VERA NORIEGA, J.A.; PIMENTEL, C.E. y BATISTA, F.J. (2005). "Redes semánticas: Aspectos teóricos, técnicos, metodológicos, analíticos", *Ra Ximhai*, septiembre-diciembre, Año/Vol. 1, N° 3, Universidad Autónoma Indígena de México. Consultada en <http://www.uaim.edu.mx/webraximhai> (18 mayo 2007).

HINOJOSA R., G.

Mtro. Guillermo HINOJOSA RIVERO
guillermo.hinojosa@iberopuebla.edu.mx
Universidad Iberoamericana Puebla
Departamento de Ciencias para el Desarrollo Humano
Director
Tel. 01 222 229-07-00 ext. 433