

LA PROBLEMÁTICA DE LA RI

De forma general – según Baeza-Yates [BAEZA-YATES] – el problema de la RI puede ser estudiado desde dos puntos de vista: el computacional y el humano. El primer caso tiene que ver con la construcción de estructuras de datos y algoritmos eficientes que mejoren la calidad de las respuestas. El segundo caso corresponde al estudio del comportamiento y de las necesidades de los usuarios.

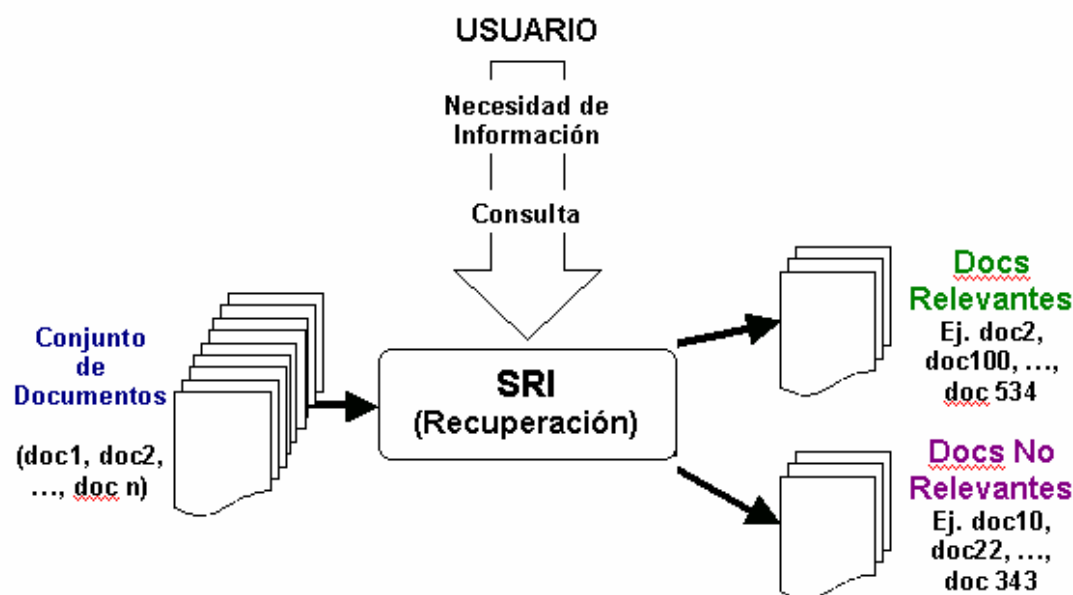


Figura 1 – La problemática de la RI

Si se analiza la problemática de la RI desde un alto nivel de abstracción (Figura 1) podemos establecer que:

- Existe una colección de documentos que contienen información de interés (sobre uno o varios temas)
- Existen usuarios con necesidades de información, quienes las plantean al SRI en forma de una consulta (en inglés, *query*. En adelante, ambas palabras se utilizarán indistintamente)
- Como respuesta, el sistema retorna – de forma ideal – referencias a documentos “relevantes”, es decir aquellos que satisfacen la necesidad expresada, generalmente en forma de una lista rankeada.

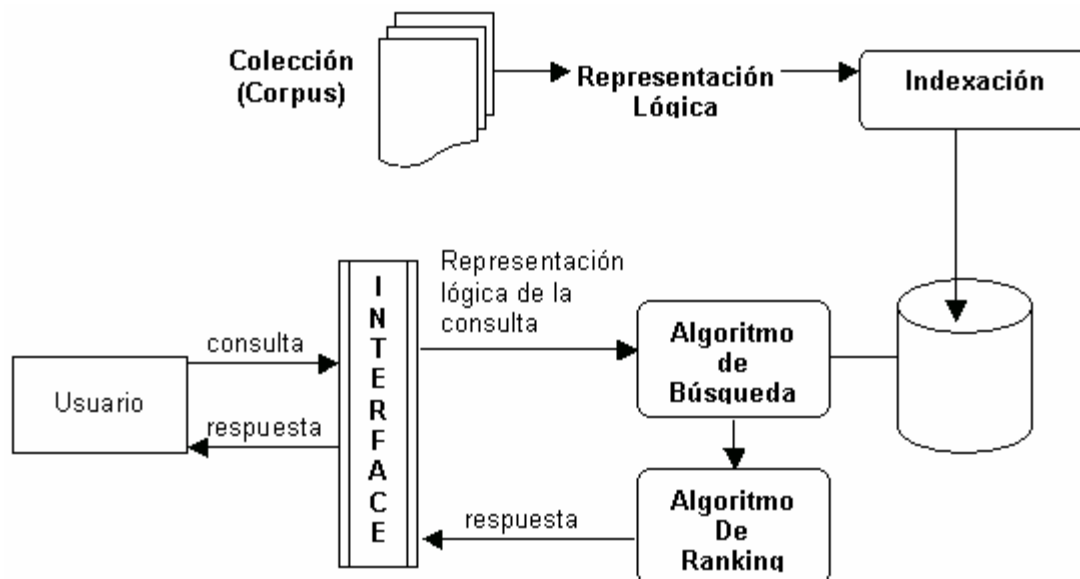


Figura 2 – Arquitectura básica de un SRI

Como podemos observar, se inicia desde un conjunto de documentos de texto, los cuales están compuestos por sucesiones de palabras que forman estructuras gramaticales (por ejemplo, oraciones y párrafos). Tales documentos están escritos en lenguaje natural y expresan ideas de su autor sobre un determinado tema. El conjunto de todos los documentos con los que se trata y sobre los que se deben realizar operaciones de RI se denomina **corpus**, **colección** o **base de datos textual o documental**. Para poder realizar operaciones sobre un corpus, es necesario obtener primero una **representación lógica** de todos sus documentos, la cual puede consistir en un conjunto de términos, frases u otras unidades (sintácticas o semánticas) que permitan – de alguna manera – caracterizarlos. Por ejemplo, la representación de los documentos mediante un conjunto de sus términos se la conoce como “bolsa de palabras” (*bag of words*).

A partir de la representación lógica existe un proceso (**indexación**) que llevará a cabo la construcción de estructuras de datos (normalmente denominadas **índices**) que la almacene y soporte búsquedas eficientes. Es importante destacar que una vez construidos los índices, los documentos del corpus pueden ser eliminados del sistema ya que éste retornará las referencias a los mismos porque cuenta con la información necesaria para hacerlo. En tal caso, el usuario será el encargado de localizar el documento para consultarlo. A los sistemas que funcionan bajo este modelo se los denomina “sistemas referenciales”, en contraste con los que sí almacenan y



Como mencionamos, la recuperación de información intenta resolver el problema de encontrar documentos relevantes que satisfagan la necesidad de información del usuario. Sin embargo, se ha planteado la dificultad para llevar a cabo esta tarea debido a la imposibilidad de expresar exactamente tal necesidad. Además, la noción de relevancia es un juicio subjetivo [RIJSBERGEN] y depende de diferentes factores relacionados más cercanamente con el usuario. La relevancia de un documento respecto a un *query* se refiere a cuánto el primero responde al segundo. De igual manera, luego el usuario evalúa qué tanto, es decir, en qué medida, se satisface su necesidad de información [KORFHAGE].

Es por ello, que se plantea la relevancia como similitud, para poder comparar documentos con consultas y – bajo ciertos criterios – definir una medida de distancia entre ambos. Por lo tanto, se puede plantear la idea que “un documento es relevante a una consulta si son similares”, donde la medida de similitud puede estar basada en diferentes criterios (coincidencias de términos, significado de éstos, frecuencia de aparición de términos y distribución del vocabulario, entre otros).

Martínez Méndez y otros [MARTINEZ] resaltan la dificultad para determinar la relevancia o no de un documento respecto de una consulta. Plantean – por ejemplo – que dos personas pueden juzgar un mismo documento de diferente manera y que es difícil establecer los criterios para la evaluación de la relevancia. Finalmente, mencionan la idea de relevancia parcial, es decir, cuando solo una parte del documento se considera relevante.

Por otro lado, como el *query* no describe exactamente la necesidad de información del usuario, algunos autores [KORFHAGE] definen el concepto de “pertinencia”, donde se incluyen las restricciones impuestas por el SRI. Este concepto está relacionado con la utilidad del documento para el usuario [MARTINEZ], de acuerdo a la necesidad de información original que guió su búsqueda, independientemente si es en parte o todo el documento.

Sin embargo – y a pesar de las dificultades para determinarla – el concepto genérico de relevancia es aceptado ampliamente por la comunidad de RI para evaluar la respuesta de un SRI respecto de una consulta de un usuario, la cual – como ya mencionamos – surge a partir de una necesidad de información.



Si bien es el primer modelo desarrollado y aún se lo utiliza, no es el preferido por los ingenieros de software para sus desarrollos. Existen diversos puntos en contra que hacen que cada día se lo utilice menos y – además – se han desarrollado algunas extensiones, bajo el nombre modelo booleano extendido [WALLER] [SALTON_b], que tratan de mejorar algunos puntos débiles.

b) MODELO VECTORIAL

Este modelo fue planteado y desarrollado por Gerard Salton [SALTON_c] y – originalmente – se implementó en un SRI llamado SMART. Aunque el modelo posee más de treinta años, actualmente se sigue utilizando debido a su buena performance en la recuperación de documentos.

Conceptualmente, este modelo utiliza una matriz documento–término que contiene el vocabulario de la colección de referencia y los documentos existentes. En la intersección de un término t y un documento d se almacena un valor numérico de importancia del término t en el documento d ; tal valor representa su *poder de discriminación*. Así, cada documento puede ser visto como un vector que pertenece a un espacio n -dimensional, donde n es la cantidad de términos que componen el vocabulario de la colección. En teoría, los documentos que contengan términos similares estarán a muy poca distancia entre sí sobre tal espacio. De igual forma se trata a la consulta, es un documento más y se la mapea sobre el espacio de documentos. Luego, a partir de una consulta dada es posible devolver una lista de documentos ordenados por distancia (los más relevantes primero). Para calcular la semejanza entre el vector consulta y los vectores que representan los documentos se utilizan diferentes fórmulas de distancia, siendo la más común la del coseno.

Obsérvese el siguiente ejemplo donde se representa a un documento d y a una consulta c :

Documento: “*La República Argentina ha sido nominada para la realización del X Congreso Americano de Epidemiología en Zonas de Desastre. El encuentro se realizará...*”

Consulta: “*argentina congreso epidemiología*”

	argentina	...	congreso	epidemiología	...
--	-----------	-----	----------	---------------	-----

d_1	0.5		0.3	0.2	
...					
d_n					
Consulta	0.4		0.3	0.3	

Matriz término-documento con pesos normalizados entre 0 y 1

c) MODELO PROBABILÍSTICO

Fue propuesto por Robertson y Spark-Jones [ROBERTSON] con el objetivo de representar el proceso de recuperación de información desde el punto de vista de las probabilidades. A partir de una expresión de consulta se puede dividir una colección de N documentos (figura 4) en cuatro subconjuntos distintos: REL conjunto de documentos relevantes, REC conjunto de documentos recuperados, RR conjunto de documentos relevantes recuperados y NN el conjunto de documentos no relevantes no recuperados.

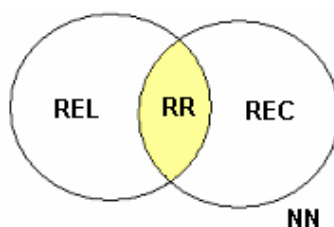


Figura 4. División de la colección

El resultado ideal de una consulta se da cuando el conjunto REL es igual REC. Como resulta difícil lograrlo en primera intención, el usuario genera una descripción probabilística del conjunto REL y a través de sucesivas interacciones con el SRI se trata de mejorar la performance de recuperación. Dado que una recuperación no es inmediata dado que involucra varias interacciones con el usuario y que estudios han demostrado que su performance es inferior al modelo vectorial, su uso es bastante limitado.

d) MODELOS PARA DOCUMENTOS ESTRUCTURADOS

Los modelos clásicos responden a consultas, buscando sobre una estructura de datos que representa el contenido de los documentos de una colección, únicamente como listas de términos significativos. Un modelo de recuperación de documentos estructurados utiliza la organización de los mismos a los efectos de mejorar la performance y brindar servicios alternativos al usuario (por ejemplo, uso de memoria visual, recuperación de



- [SALTON] Salton, G. (1971) (editor). **The SMART Retrieval System – Experiments in Automatic Document Processing**. Ed. Prentice Hall Inc. Englewood Cliffs, NJ.
- [SMEDT] Smedt, K. Liseth, A. Hassel, M. y Dalianis, H. (2005). **How short is good? An evaluation of automatic summarization**. En Holmboe, H. (ed.) **Nordisk Sprogteknologi 2004**. Årbog for Nordisk Språkteknologisk Forskningsprogram 2000-2004, pp 267-287.
- [WADE] Wade, C. y Allan, J., (2005) **Passage Retrieval and Evaluation**, CIIR Technical Report.
- [WALLER] Waller, W. G. y Kraft, D. H. (1979) **A mathematical model for a weighted Boolean retrieval system**". Information Processing and Management, 15(5):235-245. 1979.