

Gestión de resúmenes para dispositivos móviles

Managing summaries for mobile devices

Fernando Llopis, José M. Gómez, Elena Lloret, Patricio Martínez, Yoan Gutiérrez

Departamento de Lenguajes y Sistemas Informáticos

Universidad de Alicante

Apdo. de Correos 99

E-03080, Alicante, Spain

{llopis,jmgomez,elloret,patricio,ygutierrez}@dlsi.ua.es

Resumen: Los dispositivos móviles han cambiado notablemente la forma en la que los usuarios acceden a la información disponible en Internet. Estos dispositivos permiten un acceso instantáneo desde cualquier lugar, pero tienen una serie de limitaciones importantes sobre los ordenadores personales. Su limitada pantalla, así como en ocasiones la limitada capacidad de recepción de la información, dado el coste, hacen que la selección de información a acceder sea todavía más importante. La generación automática de resúmenes multi-documentos es una alternativa de interés que es el objeto de este artículo. Así en este artículo se presentan y evalúan un modelo de generación automática de resúmenes, junto con un sistema de recuperación de información basado en pasajes.

Palabras clave: Gestión de resúmenes, Recuperación de Información, Recuperación de Pasajes, Turismo, COMPENDIUM, IR-n

Abstract: Mobile devices have significantly changed the way users access the information available on Internet. These devices allow instant access anytime and anywhere, but they have a number of important limitations with respect to personal computers. The limited screen space and, sometimes, the limited capacity to receive the information, make the selection of information even more important. Automatic summary generation from multi-document summarization is an interesting alternative which is the subject of this paper. Therefore, in this article is presented and evaluated a model of automatic summarization with an information retrieval system based on passages.

Keywords: Summarization, Information Retrieval, Passage Retrieval, Tourism, COMPENDIUM, IR-n

1 Introducción

La evolución de los hábitos y necesidades de los usuarios de Internet ha sido una constante. Una vez que los usuarios podían acceder a millones de documentos, el principal problema a resolver fue el de cómo facilitar la búsqueda de los documentos relevantes para el usuario entre todos ellos. La aparición de los buscadores como Yahoo, Google, Bing y muchos más resolvió, de una forma razonablemente satisfactoria, este problema. Una persona sentada cómodamente frente a la pantalla del ordenador podía realizar una búsqueda de documentos. El esquema de prácticamente todos los buscadores es el de proporcionar una lista de referencias a 10/20 páginas de internet acompañados con una selección de fragmentos de textos incluidos en

ellas y que suelen contener palabras de la consulta del usuario.

La siguiente tarea del usuario era la de ir accediendo a aquellas páginas que podía considerar relevantes, en base a esos fragmentos de texto que le mostraba el buscador. El proceso se convertía en una repetición de estas tareas hasta que el usuario daba por satisfechas sus necesidades de información, o si decidía refinar su búsqueda inicial dado el poco éxito de la misma.

Los principales objetos de la investigación en la recuperación de información se basaban en ser lo más eficientes y eficaces a la hora de seleccionar los documentos relevantes ante las consultas del usuario. Pero algo ha cambiado notablemente en la forma en los dispositivos

con los que los usuarios acceden a Internet¹. Los dispositivos móviles son uno de los medios más utilizados para el acceso a la información. La principal ventaja de los dispositivos móviles es que prácticamente están disponibles en cualquier momento. Los principales inconvenientes a la hora de facilitar el procesamiento de la información son dos: el primero el tamaño de la pantalla, cuatro o cinco veces más pequeña que los monitores que suelen acompañar a los ordenadores de sobremesa o portátiles; el segundo son las limitaciones con las que la información puede llegar al dispositivo móvil. La progresiva aparición de las redes 4G ha solucionado en parte el problema de la escasa velocidad en accesos a la información desde un dispositivo móvil, pero siguen sin solucionar el coste todavía alto que cobran las compañías por la transmisión de información a los móviles. Hay otro aspecto que introdujo el llamado Internet 2.0, y es la importancia del usuario no solo como receptor pasivo de la información sino como generador activo de la misma a través de opiniones. Así, suelen ser más valoradas las aportaciones que los usuarios realizan sobre un tema, que lo que podría ser considerada la versión oficial. Esto ha supuesto la aparición de innumerables páginas que complementan la información con opiniones de todo tipo acerca de algún elemento.

Así una búsqueda del tipo "Iglesias Alicante" en la conocida página de viajes Trip Advisor devolvería un resultado como en la Figura 1.

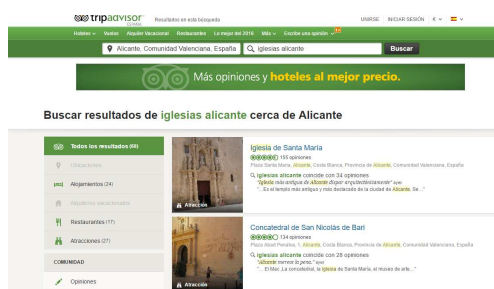


Figura 1: Búsqueda de Iglesias Alicante en Trip Advisor

Ahora le tocaría al usuario navegar a través de los dos enlaces que se proponen. Al seleccionar en una de las primeras, se pueden acceder a las opiniones que los usuarios han ido escribiendo sobre la misma como se puede

ver en la Figura 2.

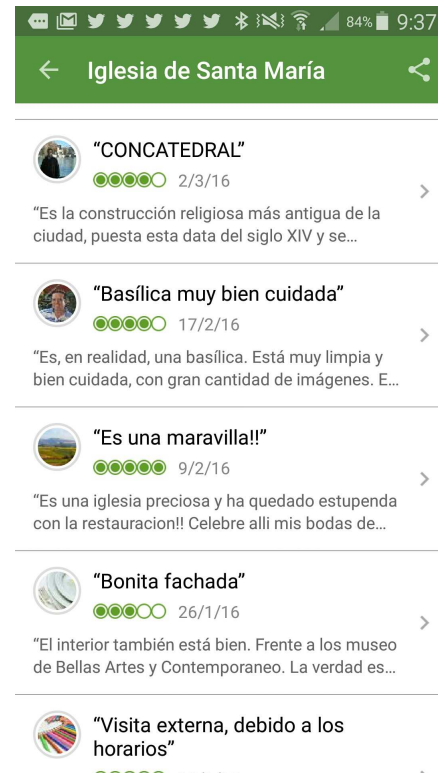


Figura 2: Opiniones de usuarios en aplicación móvil

Para facilitar la búsqueda a los usuarios se trataría de conseguir en primer lugar obtener información relevante ante una petición del usuario y en vez de mostrarle ordenados una serie de documentos o enlaces en función de su relevancia, se propone realizar un procesamiento adicional a los documentos iniciales con el objeto de identificar la parte más relevante de los mismos y mostrar en la pantalla del dispositivo móvil del usuario un pequeño fragmento de texto que resuma o identifique partes clave de los documentos relevantes. Con un esquema parecido al que se muestra en la Figura 3.



Figura 3: Resumen de opiniones

La idea es resumir las opiniones de los usuarios en base a los términos más relevan-

¹<http://www.idc.com/getdoc.jsp?containerId=prUS40664915>

tes en el total de las opiniones. Hay términos que se repiten como *bonita*, *Barroco*, *siglo XIV* y *horarios*. Los tres primeros son positivos y el último algo negativo por los problemas de acceso que indican algunos usuarios. Con esto se facilita al usuario decidir si un lugar merece la pena visitarlo leyendo tres o cuatro frases. Siempre tiene la opción, si así lo desea, de leer todas las opiniones, pero la opción de leer un resumen es cómodo y rápido, sobre todo si debe seleccionar por ejemplo que monumentos debe visitar con restricciones de tiempo.

El objetivo es ser capaz de obtener resúmenes desde diversos documentos. Para ello vamos a basarnos en los modelos propuestos en la generación automática de resúmenes cuya misión fundamental es la de reducir la dimensión de un texto manteniendo la información relevante del mismo, incluida dentro del ámbito del procesamiento del lenguaje natural (PLN).

Los modelos de resúmenes que presentaremos en este artículo se centran en los de tipo extractivo, los cuales son los que se generan a partir de la selección y extracción de frases del texto original. La frase es la unidad mínima de información con el significado suficiente para ser presentada al usuario y que éste pueda entender razonablemente su significado.

En nuestro trabajo se han realizado experimentos para analizar la forma óptima de utilizar sistemas de resúmenes pero que utilicen fuentes de texto diversas. Para llevar a cabo esto, consideramos todos los documentos de estas distintas fuentes como un único texto que, posteriormente, fue el que se resumió. Esta primera experimentación la hemos preferido realizar sobre una colección de textos y preguntas que va a ser evaluada dentro de las conferencias NTCIR.

La experimentación realizada en este artículo se basa en el uso combinado de un sistema de recuperación de Información basado en pasajes, el IR-n (Llopis and Vicedo, 2001) y el sistema de generación automática de resúmenes COMPENDIUM (Lloret and Palomar, 2012). Para las pruebas iniciales hemos utilizado la colección propuesta en la tarea Mobile Click del NTCIR-12.

El artículo se estructura del siguiente modo: en primer lugar se hace un análisis de la situación actual de los modelos de resúmenes de opiniones en la sección 2 para luego

presentar el sistema generador de resúmenes automáticos COMPENDIUM y el sistema de recuperación de Pasajes IR-n en las secciones 3 y 4. Las secciones 5 y 6 describen la experimentación y resultados obtenidos por los sistemas. Finalmente, se exponen las conclusiones obtenidas en la sección 7.

2 *Análisis del contexto actual*

En la mayoría de los enfoques de sistemas de resúmenes de comentarios se integra un motor de minería de opiniones que tiene como objetivo fundamental identificar y clasificar las opiniones presentes en los comentarios diferenciando positivas, negativas o neutras.

Por ejemplo, el enfoque propuesto en (Kokkoras et al., 2008) genera múltiples resúmenes de opiniones, considerando la información de metadatos que suelen acompañar. El proceso de resumen incluye como primera etapa la creación de un diccionario que contiene el vocabulario específico del dominio. En una segunda etapa, las frases se puntúan con respecto a su relevancia utilizando métodos estadísticos basados en recuentos de frecuencias y la presencia de las palabras en el diccionario obtenido en la fase anterior.

En la tercera etapa, la puntuación asignada inicialmente a una sentencia se ajusta teniendo en cuenta la contribución de los metadatos. Los metadatos incluyen información que no está directamente presente en el texto de la revisión, pero pueden proporcionar información sobre la utilidad de la revisión, o el nivel técnico de la crítica, siendo por tanto útil para determinar la importancia de una oración.

En la etapa final, la calidad del resumen generado se mejora mediante la eliminación de frases redundantes. Esto se realiza mediante la comparación de las palabras en la frase resumen con las palabras de las frases restantes, y el recálculo de la puntuación de las frases disminuyendo la de aquellas que contienen palabras que ya aparecen en el resumen.

En el enfoque mencionado anteriormente, las opiniones son consideradas como un texto único. El problema es que una valoración puede contener diferentes opiniones con respecto a varios aspectos de ese tema. Por ejemplo, acerca de un hotel, un usuario puede escribir un comentario expresando que “el hotel estaba muy bien situado, pero el servicio fue un poco decepcionante”; o para un teléfono

móvil, los usuarios pueden expresar sus opiniones sobre su precio, su batería, su diseño, su pantalla, etc.

Ante esto, las técnicas de minería de opiniones comenzaron a desarrollar enfoques que teniendo en cuenta las diferentes características del producto, servicio, lugar, etc. Esto se conoce como multi-aspecto o la minería opinión multi-faceta, apareciendo también la tarea de resúmenes de opinión multi-aspecto, cuando el objetivo es resumir la revisión, en relación con las diferentes características. En (Abulaish et al., 2009) se genera el resumen teniendo en cuenta los resultados de un sistema de minería de opinión, y que representa de una manera visual, a través de un gráfico. El enfoque de minería de opiniones propuesto extrae diferentes características de la revisión, junto con los modificadores asociados a la opinión de que disponga para cada uno de los aspectos que considera (por ejemplo, el precio - excelente, alto, caro, barato). Tras haber evaluado la lista de las opiniones y modificadores asociados para cada característica extraída, se establece su polaridad usando (Baccianella and Sebastiani, 2010), asociando a cada elemento a tres puntuaciones numéricas, que lo definen como positivo o negativo. En (Ly et al., 2011), las diferentes características expresadas en los comentarios se llaman facetas. Estas facetas son identificados a través de una etapa de identificación de las facetas de productos, que comprende etiquetado gramatical, análisis de dependencias y el recuento de la frecuencia. Una vez que se han determinado las facetas de productos, el enfoque de integración se limita únicamente a las frases con opiniones (positivas o negativas), obtenidos a través de un enfoque de minería de opiniones. Estas frases de opinión resultantes se agrupan en función de su similitud contenido, distinguiendo entre frases positivas y negativas.

Finalmente se selecciona la frase más representativa de cada grupo para formar el resumen.

Pero en ocasiones las opiniones no marcan en sí aspectos positivos o negativos. El hecho de que un iglesia del siglo XIV sea bonita, es positivo, pero a su vez incorpora información acerca de que es del siglo XIV lo que puede ser positivo para un interesado por las iglesias de esa época y negativo para una persona que desea visitar edificios modernos.

El modelo que planteamos en este artículo

se centra en el resumen de textos provenientes de varias fuentes, como puedan ser los comentarios pero que traten de responder a una petición concreta del usuario.

3 Sistema de resúmenes COMPENDIUM

COMPENDIUM es un sistema automático que es capaz de realizar resúmenes de uno o varios documentos. El sistema utiliza diferentes técnicas para la identificación, selección y extracción de la información más relevante. Para ello se utilizan cinco etapas en la que la salida de cada una de ellas es la entrada de la siguiente hasta producir el resultado final.

- **Análisis lingüístico.** Esta etapa consiste en el preproceso del texto, realizando un análisis lingüístico básico, utilizando herramientas y recursos externos. Este pre-procesamiento incluye la segmentación, tokenización, etiquetado y el análisis sintáctico, así como la identificación y borrado de stopwords.
- **Detección de redundancias.** El objetivo de esta etapa es identificar la información redundante en los documentos de origen, con el objeto de no incluirla en el resumen. Para este propósito, las técnicas de Textual Entailment (TE) han demostrado ser las adecuada para esta etapa, ya que pueden determinar si el significado de un fragmento de texto puede deducirse de otro (Glickman, 2006).
- **Identificación de temas principales.** El objetivo de esta etapa es determinar los temas principales del documento/s que se pretende resumir.

En COMPENDIUM, los temas de un documento están representados por la frecuencia de los términos que contiene. A raíz de esta declaración, se supone que los términos más frecuentes de un documento son indicativos de los temas incluidos en él. Por lo tanto, las oraciones que contienen términos frecuentes se puntúan más alto tal como se comentará en la etapa de *detección de relevancia*.

Sin embargo, es importante señalar que las “stopwords”, que fueron identificadas previamente en fases anteriores no se tienen en cuenta para el cálculo de la frecuencia de un documento, por tanto, no

forma parte de los temas de un documento.

- **Detección de relevancia.** En esta etapa, se calcula y asigna un peso a cada frase, en función de su relevancia.

Este peso toma la frecuencia de los términos calculada en la etapa anterior, y se combina con otra característica basada en *El principio de la cantidad de código* (Givón, 1990).

El fundamento de esta teoría es que una determinada información en un texto, independientemente de que sea nueva o dada, debe recibir una codificación tal que resalte más o menos según el grado de relevancia que tiene esta información en el texto. Si la información es más relevante recibirá una “cantidad mayor de codificación”, esta información se codificará con mayor peso léxico; y, al revés, si es menos relevante se utilizará una “cantidad menor” para su codificación de menor peso léxico.

La puntuación total de la frase se calcula considerando la frecuencia de las palabras que forman parte de los sintagmas nominales que contiene una frase y la longitud de los mismos, de manera que frases que tengan sintagmas nominales con palabras más frecuentes obtendrán mayor relevancia respecto al resto.

- **Generación del resumen.** El objetivo de esta etapa es la de generar un resumen con una longitud específica. Esta longitud se expresa en forma de una tasa de compresión (es decir, el porcentaje de información que el resumen contiene con respecto al documento de origen). Dada una tasa de compresión, las frases más relevantes (es decir las que tienen una puntuación más alta) serán las seleccionadas para formar el resumen final hasta dicha longitud deseada. Con el fin de minimizar los posibles problemas con la coherencia, las frases seleccionadas se presentan en el mismo orden en que están en el documento de origen.

4 Sistema IR-N

Los sistemas de Recuperación de Pasajes (PR) son sistemas de Recuperación de Información que determinan la relevancia de un documento con respecto de una pregunta en

base a la similitud de diferentes fragmentos (pasajes) de dicho documento con respecto a la misma pregunta. Estos modelos no solo permiten mejorar la localización de documentos relevantes sino que además permiten localizar con mayor exactitud la parte realmente relevante del documento (Kaszkiel and Zobel, 1997).

Los sistemas de PR se clasifican en función de cómo determinan los pasajes de cada documento. El sistema IR-n (Llopis and Vicedo, 2001) es un sistema de PR que define los pasajes en base a un número determinado de frases. Esto permite dotar a los pasajes de cierto contenido sintáctico.

Así el modelo de trabajo del sistema IR-n es el siguiente

1. Un documento se divide en frases. Cada pasaje puede estar formado por hasta un número fijo de frases. Para la experimentación realizada en este trabajo se ha trabajado en pasajes de una única frase.
2. La relevancia de un pasaje o frase p con respecto a una pregunta o query q se obtiene de la siguiente forma:

$$Relevancia = \sum_{t \in p \wedge q} W_{p,t} * W_{q,t} \quad (1)$$

Siendo:

$$W_{p,t} = \log_e(f_{p,t} + 1),$$

$f_{p,t}$ es el número de apariciones del término t en el pasaje. p ,

$$W_{q,t} = \log_e(f_{q,t} + 1) * idf,$$

$f_{q,t}$ es el número de apariciones del término t en pregunta q ,

$$idf = \log_e(n/f_t + 1),$$

n es el número total de documentos de la colección

f_t es el número de documentos donde aparece el término t

La fórmula es muy parecida a la definida en (Salton, 1989) y conocida como “la del Coseno”. La principal diferencia es que en el sistema IR-n se omite los cálculos de normalización por el tamaño del documento, dado que estamos trabajando con frases, y se entiende que cada una de ellas es una unidad comparable a cualquier otra.

El planteamiento de utilizar el sistema IR-n como un sistema de generación de resúmenes se basa en el concepto de favorecer las frases que contienen los términos que se repiten más frecuentemente en todos los documentos. Para ello se plantea la propuesta de una vez seleccionados todos los documentos generar una pregunta con la concatenación de todos ellos y luego evaluar la relevancia para cada documento de forma individual con respecto a esa pregunta generada.

Por ejemplo la pregunta 1 era “*Iglesias en Alicante*” se lanzaría una búsqueda con un buscador web tradicional o en el caso de ejemplo se seleccionarían las opiniones de cada usuario de la web TripAdvisor. A cada uno de estos elementos les llamaremos iUnits.

- iUnit 1 Es la construcción religiosa más antigua de la ciudad, puesta esta data del siglo XIV y se encuentra asentada sobre el mismo lugar que antaño ocupaba una mezquita... para los alicantinos es la CONCATEDRAL. La visita es gratuita... pero mejor consultar horario.
- iUnit 2 Externamente la iglesia está bien, pero es una pena de no ponerla más al alcance de los turistas. No indican en el exterior los horarios, lo cual dice muy poco a favor de los cuidadores de dicho edificio, Debe de estar abierta de 11 a 12:30 y de 19 a 20:30, en los horarios de misa.
- iUnit 3 Alicante reúne varias iglesias realmente preciosas y la de Santa María es la que más me gusta, su fachada, sus imágenes procesionales, sus pinturas, altar mayor, capillas, techo, me pareció bellísima y que debe visitarse.

Posteriormente el sistema concatena todas las iUnits obtenidas generando una nueva consulta, cuya relevancia se evalúa con cada una de las frases de cada iUnit por separado obteniendo el valor de relevancia Rel_{IRn} tal como se puede ver en la figura 4:

El resumen se formaría por las iUnits con valor de relevancia mayor. Aunque dado el modelo que propone el sistema IR-n se podría aplicar de forma sencilla a recuperar las frases más relevantes en general de todas las iUnits.

5 Evaluación

Actualmente estamos trabajando en la recopilación de un corpus de prueba que pueda

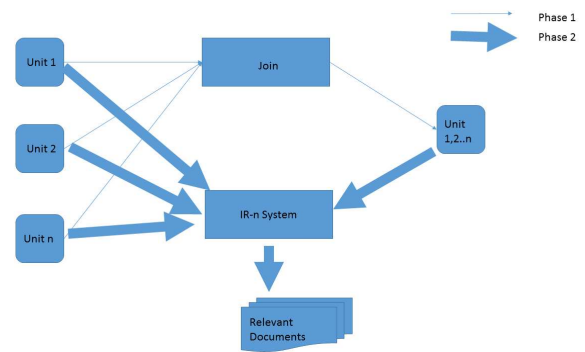


Figura 4: Sistema de Referencia: IR-n

servir para las evaluar la bondad del modelo dentro del sector turístico, pero aprovechando la tarea Mobile Click de las conferencias NTCIR hemos podido realizar una serie de experimentaciones sobre generación y ordenación de frases sobre documentos de información general.

Una de las tareas propuestas presentaba el siguiente planteamiento: se disponía de una consulta Q que se utilizaba para obtener documentos relevantes utilizando buscadores web tradicionales obteniendo una serie de iUnits (documentos de una frase). El objetivo era ordenar las iUnits en función de la relevancia con respecto a la pregunta original.

Por ejemplo la pregunta 1 era “*Pregunta 1(Q1) michael jackson death*” y las iUnits (iU) que ponía la organización a disposición de los participantes eran entre otras:

- Q1 iU1 family concerned about murray role.
- Q1 iU2 giving singer nightly doses of propofol.
- Q1 iU3 murray first met jackson in las vegas.

Los sistemas se evaluaban en base a la conocida medida Q-measure (Sakai, 2007), la que valora el orden en el que las frases relevantes son devueltas por los participantes teniendo en cuenta que la tarea exigía devolver todas las iUnits originales.

5.1 Resultados

Para las pruebas se disponía de la colección de preguntas y unidades utilizadas el año anterior. Así hemos decidido utilizar el modelo IR-n como caso base. Posteriormente hemos probado con el sistema COMPENDIUM

con varios porcentajes de comprensión (10 %, 20 % y 40 %), y utilizando o no los módulos de Textual Entailment que permiten eliminar información redundante.

5.1.1 Sistema Referencia: IR-n

Nuestro sistema base fue el de utilizar únicamente el sistema IR-n sin ningún preprocesamiento de las iUnits ni eliminación de las stopwords (ya que el modelo de calculo de relevancia les da peso ínfimo al prácticamente hallarse en todas las iUnits). La pregunta entrada del sistema es la concatenación de todas las iUnits referidas a cada pregunta.

Así el sistema concatenaba todas las iUnits generando una consulta que se evaluaba contra cada iUnit por separado obteniendo el valor de relevancia. Para el ejemplo anterior la consulta generada sería la siguiente.

Gen1 family concerned about murray role. giving singer nightly doses of propofol. murray first met jackson in las vegas.

5.1.2 Resultados con COMPENDIUM

Los sistemas de generación automática de resúmenes tienen como entrada un documento con la concatenación de todos los iUnits. A ese documentos se le aplica el sistema COMPENDIUM con 3 diferentes grados de comprensión (10 %, 20 % y el 40 %).

Para cumplir los requisitos de la tarea, se le daba un valor de relevancia Rel_{Com} 1 a aquellas frases que el sistema de resúmenes seleccionaba y un 0 a aquellas frases que no se encontraban en el resumen.

También probamos el módulo de detección de redundancia basado en textual entailment. Los resultados de todos los experimentos se pueden ver en la Tabla 1.

	Q-measure		
	10 %	20 %	40 %
Baseline (IR-n)		0,8621	
COMPENDIUM	0,8196	0,8291	0,8403
COMPENDIUM+TE	0,8201	0,8218	0,8246

Tabla 1: Comparativa de resultados con sistemas aislados para el 10 %, 20 % y 40 % de ratio de resumen. En el baseline este ratio no se aplica.

El análisis de los resultados nos permitió obtener una serie de conclusiones:

- Textual Entailment empeoraba los resultados, al discriminar frases relevantes por estar incluidas en otras.

- En estos experimentos se obtuvieron mejores resultados con el sistema base (IR-n), ya que los sistemas de resúmenes no reordenaban las frases seleccionadas y simplemente las colocan en el orden de aparición en el texto original, siendo penalizadas en la medida Q-measure.

5.1.3 Combinando IR-n + COMPENDIUM

Dada la última conclusión obtenida, se optó por reordenar las frases obtenidas por el sistema de resúmenes teniendo en cuenta la relevancia que el sistema IR-n les había dado a cada una de ellas. Así, las primeras iban a ser las frase seleccionadas por el sistema de resúmenes, pero manteniendo el orden que el sistema IR-n había dado a cada una de ellas.

Para cada iUnit, se obtiene la siguiente relevancia:

$$Rel_{def} = Rel_{IRn} + (Rel_{Com} * 1000)$$

Los resultados se muestran en la Tabla 2. Como se puede observar, la combinación de ambos sistemas produce mejores resultados que utilizando sólo los sistemas de una manera independiente. El motivo es que IR-n reordena las frases que obtiene Compendium en lugar de utilizar el orden inicial. La medida Q-measure premia esta acción

	Q-measure
Baseline (IR-n)	0.8621
COMPENDIUM	0.8403
COMPENDIUM + IRn	0.8648

Tabla 2: Results con la combinación IR-n+COMPENDIUM

6 Evaluación

6.1 iUnit Ranking Subtask

Los resultados utilizando el modelo COMPENDIUM+IR-n sobre la colección de test se pueden ver en la Tabla 3. Se obtuvo el mejor resultado con IR-n+COMPENDIUM.

7 Conclusiones

Teniendo en cuenta las necesidades actuales de los usuarios actuales que utilizan sus dispositivos móviles para la obtención de información a través de Internet, este tipo de tareas tienen un enorme interés. El objetivo es ser capaz de reducir y condensar la información de tal forma que el usuario puede obtener de un simple vistazo la respuesta a sus dudas.

	Q-measure
IRn + COMPENDIUM (*)	0.9027
TITEC	0.9003
UHYG	0.8994
ORG	0.8975
RISAR	0.8972
RISAR	0.8962
IRn	0.8959
COMPENDIUM	0.8934

Tabla 3: Resultados para tarea iUnit Ranking Subtask

Los resultados obtenidos en las tareas del NTCIR 12 nos animan a seguir en la línea de explotación combinada de los sistemas COMPENDIUM e IR-n. Es necesario realizar una serie de trabajos adicionales para mejorar la experiencia de usuario. Por un lado evitar la redundancia de alguna de la información generada, que si que es valorada tal como se aplicaba la tarea pero no es interesante para el usuario.

Esto requiere la inclusión de técnicas adicionales de procesamiento del lenguaje natural pero con la problemática que deben estructurarse de forma que no consuman un tiempo importante que pudiera desesperar al usuario de este tipo de tareas.

Agradecimientos

Investigación realizada gracias a la financiación de los proyectos: DIIM2.0 (PROMETEOII/2014/001) de la Generalitat Valenciana; TIN2015-65100-R, DIGITY (TIN2015-65136-C2-2-R) del Ministerio de Economía y Competitividad y SAM (FP7-611312) de la Unión Europea

Bibliografía

Abulaish, M., Jahiruddin, Doja, M., and Ahmad, T. (2009). Feature and opinion mining for customer review summarization. In Chaudhury, S., Mitra, S., Murthy, C., Sastry, P., and Pal, S., editors, *Pattern Recognition and Machine Intelligence*, volume 5909 of *Lecture Notes in Computer Science*, pages 219–224. Springer Berlin Heidelberg.

Baccianella, A. E. S. and Sebastiani, F. (2010). Sentiwordnet 3.0: An enhanced lexical resource for sentiment analysis and opinion mining. In *Proceedings of the Seventh conference on International Language Resources and Evaluation (LREC'10)*,

Valletta, Malta. European Language Resources Association (ELRA).

Givón, T. (1990). *Syntax: A functional-typological introduction, II*. John Benjamins.

Glickman, O. (2006). *Applied Textual Entailment*. PhD thesis.

Kaszkiel, M. and Zobel, J. (1997). Passage Retrieval Revisited. In *Proceedings of the 20th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, Text Structures, pages 178–185, Philadelphia, PA, USA.

Kokkoras, F., Lampridou, E., Ntonas, K., and Vlahavas, I. (2008). Mopis: A multiple opinion summarizer. In Darzentas, J., Vouros, G., Vosinakis, S., and Arnellos, A., editors, *Artificial Intelligence: Theories, Models and Applications*, volume 5138 of *Lecture Notes in Computer Science*, pages 110–122. Springer Berlin Heidelberg.

Llopis, F. and Vicedo, J. L. (2001). IR-n system, a passage retrieval system at CLEF 2001. In *Workshop of Cross-Language Evaluation Forum (CLEF 2001)*, Lecture notes in Computer Science, pages 244–252, Darmstadt, Germany. Springer-Verlag.

Lloret, E. and Palomar, M. (2012). COMPENDIUM: a text summarisation tool for generating summaries of multiple purposes, domains, and genres. *Natural Language Engineering*, 19(02):147–186.

Ly, D. K., Sugiyama, K., Lin, Z., and Kan, M.-Y. (2011). Product review summarization from a deeper perspective. In *Proceedings of the 11th Annual International ACM/IEEE Joint Conference on Digital Libraries*, JCDL '11, pages 311–314, New York, NY, USA. ACM.

Sakai, T. (2007). On the reliability of information retrieval metrics based on graded relevance. *IP&M*, 43(2):531–548.

Salton, G. A. (1989). *Automatic Text Processing: The Transformation, Analysis, and Retrieval of Information by Computer*. Addison Wesley, New York.