

## Sistemas recomendadores de objetos de aprendizaje: cuestiones involucradas para su mejora

### Learning object recommender systems: issues for their improvement

Claudia Deco<sup>1</sup>, Cristina Bender<sup>2</sup>

<sup>1,2</sup> Universidad Nacional de Rosario, Argentina

<sup>1</sup>Correo electrónico: [cdeco@uca.edu.ar](mailto:cdeco@uca.edu.ar), [deco@fceia.unr.edu.ar](mailto:deco@fceia.unr.edu.ar)

<sup>2</sup>Correo electrónico: [cbender@uca.edu.ar](mailto:cbender@uca.edu.ar), [bender@fceia.unr.edu.ar](mailto:bender@fceia.unr.edu.ar)

Recibido: 6 de junio de 2017

Aceptado: 2 de noviembre de 2017

---

#### Resumen:

Los Sistemas Recomendadores de Objetos de Aprendizaje son importantes para ayudar a la formación del profesional debido a que permiten encontrar los recursos educativos que mejor se ajustan al estilo de aprendizaje y perfil del usuario. En este artículo se presentan diversas cuestiones relacionadas con la mejora de la usabilidad de repositorios de recursos educativos para ayudar en la gestión de dichos repositorios. Un primer aspecto es dar soporte a las tareas de recopilación de documentos que realiza el administrador del repositorio con el objetivo de detectar documentos plausibles de ser cargados en estos repositorios junto con sus metadatos de interés. Además, se propone facilitar la carga de recursos educativos al repositorio utilizando extracción automática de metadatos. De esta forma, se pretende aumentar la población de los documentos en estos repositorios acompañando el desarrollo de Repositorios Institucionales de Acceso Abierto que es una prioridad en el marco de las políticas del Ministerio de Ciencia, Tecnología e Innovación y el Consejo Interuniversitario de Argentina. Otro aspecto es mejorar la búsqueda de información trabajando en la recomendación automática de objetos de aprendizaje considerando el perfil del usuario y también la valoración colaborativa de grupos de usuarios similares.

Palabras clave: Objetos de aprendizaje, sistemas recomendadores, metadatos.  
Abstract:

Learning Object Recommender Systems are important to help training professionals because they can find the educational resources that best fit the learning style and the profile of the user. In this article, several issues related to improving the usability of repositories of educational resources are presented to assist in the management of these repositories. First, to support the document collection tasks carried out by the repository administrator in order to detect plausible documents to be loaded in these repositories together with their metadata of interest. In addition, proposals to facilitate the loading of educational resources to the repository using the automatic extraction of metadata are presented. In this way, there should be an increase of the population of documents in these repositories, accompanying the development of Open Access Institutional Repositories, which is a priority within the framework of the policies of the Ministry of Science, Technology and Innovation and the Inter-University Council of Argentina. Another aspect is to improve the information search by working on the automatic recommendation of learning objects considering the user profile and also the collaborative assessment of similar groups of users.

Keywords: Learning objects, recommender systems, metadata.

Licencia Creative Commons



## Introducción

En los últimos años el desarrollo de los sistemas recomendadores ha crecido notablemente, siendo muy útiles en distintos campos y en particular en la educación. En el dominio de la Educación la utilización de recursos educativos acordes al estilo de aprendizaje y perfil del usuario ayudan a la formación del futuro profesional. En este contexto, el usuario puede ser tanto un docente que busca material para preparar una clase como un estudiante que requiere material para aprender un tema.

En las universidades públicas de Argentina el desarrollo de Repositorios Institucionales de Acceso Abierto es una prioridad en el marco de las políticas del Ministerio de Ciencia, Tecnología e Innovación. El objetivo de estos repositorios es facilitar el almacenamiento, la búsqueda y la reutilización de recursos educacionales. Un Objeto de Aprendizaje es todo recurso digital que apoya a la educación y que puede ser reutilizado, actualizado, combinado, separado, referenciado y sistematizado [8]. Un Repositorio de Objetos de Aprendizaje es una gran colección de objetos de aprendizaje estructurada como una base de datos con metadatos asociados y que generalmente están en la Web. Ejemplos de estos repositorios son ARIADNE ([www.ariadne-eu.org](http://www.ariadne-eu.org)), LAFLOR (Federación Latinoamericana de Repositorios, [ariadne.cti.espol.edu.ec/FederatedClient](http://ariadne.cti.espol.edu.ec/FederatedClient)) y OER Commons (Open Educational Resources, [www.oercommons.org](http://www.oercommons.org)).

Los Metadatos son un conjunto de atributos necesarios para describir las principales características de un recurso. LOM (Learning Object Metadata) es el estándar de IEEE que especifica la sintaxis y la semántica de un conjunto mínimo de metadatos necesario para identificar, administrar, localizar y evaluar un objeto de aprendizaje ([www.standards.ieee.org/findstds/standard/1484.12.1-2002.html](http://www.standards.ieee.org/findstds/standard/1484.12.1-2002.html)). Este estándar organiza los metadatos en una forma jerárquica, agrupándolos en nueve categorías:

General: contiene metadatos descriptivos del objeto como un todo. Por ejemplo: título, idioma, etcétera.

Ciclo de vida: contiene características relacionadas con la historia y el estado actual del objeto. Por ejemplo: versión, fecha, entre otros.

Meta-metadatos: contienen información sobre los metadatos del objeto. Por ejemplo, el idioma de los metadatos.

Técnica: contiene requisitos y características técnicas del objeto. Por ejemplo: formato, tamaño, ubicación, etcétera.

Educacional: contiene características pedagógicas y educacionales del objeto. Por ejemplo: nivel de interactividad, dificultad, tipo de recurso, etcétera.

Derechos: son metadatos con las condiciones de uso para la explotación del recurso. Por ejemplo si tiene o no costo.

Relación: corresponde a las características que definen la vinculación del objeto descrito con otros objetos de aprendizaje.

Anotación: son metadatos con comentarios sobre el uso educativo del objeto.

Clasificación: contiene la descripción sobre el sistema de clasificación utilizado para describir el objeto mediante palabras claves.

En los Repositorios Institucionales de Acceso Abierto los docentes-investigadores almacenan diferentes tipos de producción de recursos educativos que utilizan en sus prácticas académicas [6]. Algunos tipos de recursos son publicaciones científicas, materiales de cátedra, apuntes, presentaciones, guías de trabajo, producción en extensión, obras de arte, etc. RepHip ([rephip.unr.edu.ar/](http://rephip.unr.edu.ar/)) es el repositorio de la Universidad Nacional de Rosario. Este repositorio está implementado en la plataforma DSpace ([www.dspace.org](http://www.dspace.org)) que es la plataforma adoptada en el Sistema Nacional de Repositorios Digitales de Argentina.

En este trabajo se presentan cuestiones relacionadas con la mejora de la usabilidad de los repositorios de recursos educativos. En este sentido, se plantean cuestiones vinculadas a mejorar la búsqueda de objetos de aprendizaje, facilitar el almacenamiento de estos objetos, y asistir a la recopilación de documentos para poblar el repositorio. El resto del trabajo se organiza de la forma siguiente: A continuación se presentan cuestiones vinculadas a la mejora de la búsqueda de Objetos de Aprendizaje; la sección siguiente contiene cuestiones relacionadas con la carga de los recursos en repositorios; luego se plantean cuestiones vinculadas con la recopilación automática de documentos para su carga en repositorios.

#### Mejora de la búsqueda de objetos de aprendizaje

En el dominio de la educación existe gran cantidad y diversidad de material, disponible en la web, que puede contribuir al proceso enseñanza-aprendizaje. Pero un mismo material no es el adecuado para todos los usuarios, porque los usuarios poseen distintas características y preferencias personales, que deberían ser consideradas en el momento de la búsqueda. En este contexto, el usuario puede ser tanto un docente que busca material para preparar una clase como un estudiante que requiere material para aprender un tema.

Los Sistemas Recomendadores ayudan a los usuarios a encontrar productos de acuerdo a sus características y preferencias. Su objetivo es explorar y filtrar las mejores opciones a partir de un perfil de usuario. Entre las aplicaciones potenciales de estos sistemas, el dominio de la educación es un buen candidato. Los Sistemas Recomendadores pueden utilizar diferentes técnicas para recomendar. Así, se

tienen distintos enfoques tales como basados en contenido, de filtrado colaborativo e híbridos, entre otros. Los sistemas basados en contenido se basan en la descripción del contenido del producto y el perfil del usuario. Los sistemas de filtrado colaborativo se basan en la similitud de los productos o de los usuarios. Los sistemas híbridos combinan distintas técnicas para mejorar resultados, siendo la combinación más utilizada el filtrado colaborativo en conjunto con la recomendación basada en contenido.

Los sistemas recomendadores basados en contenido se apoyan en la descripción del contenido del recurso y el perfil del usuario. Esta descripción se realiza mediante los metadatos que son un conjunto de atributos que describen las características de un recurso. Estos metadatos siguen, como se mencionó anteriormente, el estándar LOM (Learning Object Metadata) de IEEE. El perfil del usuario contiene sus preferencias y características, y puede construirse de forma implícita o explícita. Algunos datos que pueden obtenerse implícitamente o automáticamente pueden ser el país donde está ubicado el usuario, si la consulta la está realizando desde una universidad, etc. La obtención de datos en forma explícita se puede realizar mediante, un cuestionario (Figura 1).



# Sistema Recomendador de Objetos de Aprendizaje

---

## Realice su búsqueda

**Restricciones:**

TÍTULO:  (se puede o no hacer entre comillas, según sea "palabra exacta")  
 ASG:    
 PALETO: ☐  
 CASTELLANO:    
 BTPACCTA: Año    
 AS:    
 ESTO:    
 CURSO:    
 NIVEL: Medio

**Resultados:**

DESCRICIÓN:   
 CUERPO:

**Mostrando:**

REPETICIÓN:

**Recomendados**

Figura 1: Ejemplo de cuestionario para la toma explícita de datos del usuario.

Un sistema recomendador basado en contenido se presentó en [2], donde se presenta una arquitectura para este tipo de sistemas (Figura 2).

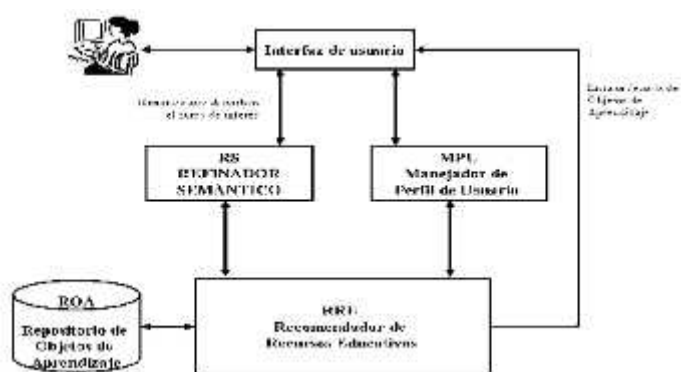


Figura 2: Arquitectura de un sistema recomendador.

El objetivo es recuperar los mejores recursos acordes al perfil del usuario. Para esto, se utiliza un sistema de reglas para inferir cuánto un recurso satisface las preferencias de un usuario. De esta forma, el sistema devuelve una lista ordenada de recursos (Figura 3), colocando en los primeros lugares los objetos más acordes al perfil del usuario que realiza la búsqueda.



Figura 3: Lista de recursos ordenada según el perfil del usuario.

En dicho trabajo se utilizó el repositorio Ariadne porque tiene metadatos LOM y un alto porcentaje cargado de metadatos educacionales. Para la experimentación intervinieron alumnos y docentes de la Licenciatura en Ciencias de la Computación de la Universidad Nacional de Rosario y se utilizaron recursos de acceso libre y disponible en idioma inglés, francés o español, y el tema de búsqueda utilizado fue "Java". Con respecto a los resultados, las pruebas con este caso de uso fueron satisfactorias. Se observó que la precisión del recomendador depende de la existencia y de la calidad de los metadatos de los objetos.

Un problema que se presentó en esta experimentación fue la falta de información en los metadatos educacionales dado que muchos estaban sin datos y también se

observó falta de calidad de muchos de los metadatos cargados. Por ejemplo, se encontraron objetos donde el metadato Idioma del objeto estaba establecido usando el título, pero el cuerpo del objeto estaba en otro idioma. Otro ejemplo fue un documento clasificado como Texto en Tipo de Recurso, pero dicho documento era un código de programa y debería haber sido clasificado como Ejercicio o Ejemplo.

Los metadatos son muy importantes para la recuperación personalizada, pero preparar estos recursos con metadatos adecuados es laborioso. Por lo tanto es importante desarrollar sistemas para la extracción automática de metadatos.

### Carga de objetos en los repositorios

El proceso de carga de objetos de aprendizaje y sus metadatos en un repositorio es una tarea ardua para el usuario y por esto hay baja carga de objetos y muchas veces no hay información en los metadatos o la misma es de baja calidad [7].

En la plataforma DSpace ([www.dspace.org/](http://www.dspace.org/)) para cargar un recurso se debe seleccionar primero una Comunidad y luego una Colección. Por ejemplo, para cargar una tesina de la Licenciatura en Ciencias de la Computación, se debe elegir primero la Comunidad del Departamento de Ciencias de la Computación y dentro de ésta la Colección de Tesinas. Luego, se describe dicho objeto mediante metadatos obligatorios y opcionales. Finalizado este paso, se está en condiciones de subir los archivos que componen el recurso. Finalmente se acepta la licencia del repositorio, se revisan los metadatos ingresados y se confirma el depósito.

Este proceso resulta tedioso para los usuarios por lo que para facilitar la carga se propuso mejorar la interfaz y modificar el flujo de carga estándar de DSpace [1]. La arquitectura de esta propuesta se presenta en la Figura 4.

En esta arquitectura, se pide al usuario en el módulo Submit Step que suba al repositorio el archivo de interés antes de cargar los metadatos. Para facilitar la carga de los metadatos se agregó el módulo Metadata Step con una herramienta para la extracción automática de ciertos metadatos. En particular interesaba extraer los metadatos Título, Autores, Palabras clave, Resumen e Idioma. Para esto se analizaron distintas herramientas para la extracción de información en documentos de texto. Cada herramienta extrae distintos tipos de metadatos, tiene sus propios objetivos y arquitectura y usa distintas técnicas. Se seleccionaron para el análisis las siguientes herramientas: Alchemy ([www.alchemyapi.com/](http://www.alchemyapi.com/)), KEA ([www.nzdl.org/kea/](http://www.nzdl.org/kea/)), MrDlib ([www.mr-dlib.org/](http://www.mr-dlib.org/)), ParsCit ([wing.comp.nus.edu.sg/parsCit/](http://wing.comp.nus.edu.sg/parsCit/)) y PDFBox ([pdfbox.apache.org/](http://pdfbox.apache.org/)). Alchemy analiza documentos de texto o HTML, extrae Autor, Palabras Claves e Idioma. KEA extrae Palabras Claves. MrDlib extrae Título y Autores. ParsCit genera un XML e identifica Título, Autor, Resumen y Palabras Claves. PDFBox busca Autor, Título,

### Filiación y Palabras Clave.

Luego de este análisis se optó por utilizar la combinación de ParsCit+Alchemy, con la cual se logró incrementar la calidad de las palabras claves retornadas pasando de un 56% de resultados correctos obtenidos con Alchemy a un 70%. En el Metadata Step, el primer paso (1) consiste en transformar el archivo en un documento de texto; en el paso (2) se lo analiza con ParsCit para extraer Título y Resumen; los pasos (3) y (4) utilizan la interface web de Alchemy para la extracción de los metadatos Autor, Palabras Clave e Idioma, que son enviadas en el paso (5) a una etapa de consolidación de todos los metadatos obtenidos. Finalmente en el módulo Review Step, el usuario valida los metadatos sugeridos por el sistema, y los confirma o los corrige si fuera necesario.

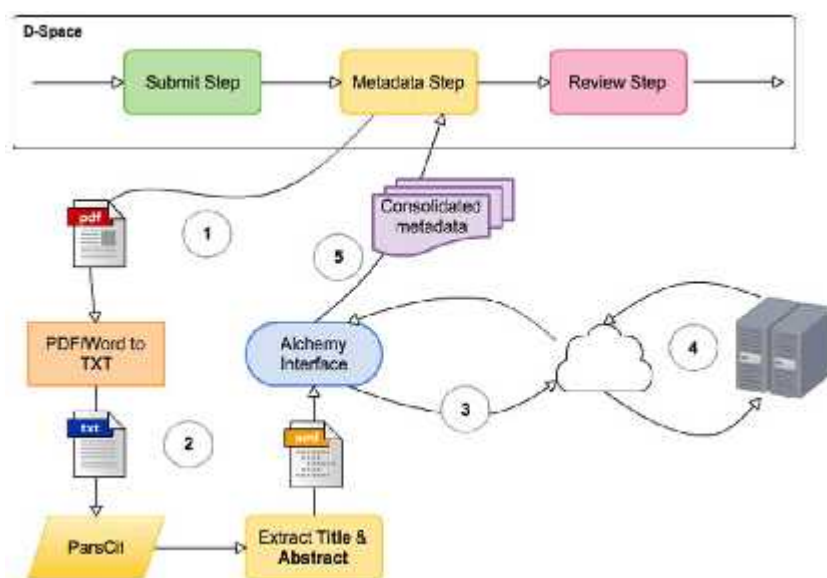


Figura 4: Arquitectura propuesta y su interacción con DSpace.

Para la experimentación de la extracción se utilizó un corpus de 760 documentos del repositorio RepHip teniendo en cuenta la diversidad temática y de tipo de archivo. Los resultados arrojaron que se simplifica el proceso de carga de objetos digitales educativos, disminuyendo así el trabajo del usuario y mejorando la cantidad y la calidad de los metadatos cargados.

### Recopilación de documentos para la carga en repositorios

Para poblar estos repositorios es necesario contar con políticas y estrategias de difusión y desarrollar herramientas informáticas para detectar objetos digitales educativos publicados en sitios web institucionales y que puedan ser cargados al repositorio. Actualmente esta tarea de recopilación es tediosa y es realizada manualmente por personal de ciencias de la información encargado del repositorio. El administrador, luego de detectar posibles documentos para cargar en el repositorio, se comunica con el autor del documento, invitándolo a subir dicho recurso, si aún no lo ha hecho.

Para ayudar en esta tarea de poblar el repositorio se realizó una recopilación automática de documentos junto con cierta información relevante tal como Título, Autores y sus datos de contacto, de docentes e investigadores de la Universidad Nacional de Rosario [3]. El objetivo de obtener datos de contacto, es que el administrador del repositorio se contacte con el autor a fin de sugerirle depositar su recurso en el repositorio. El problema encontrado en esta automatización es que muchas veces los datos de contacto requeridos (autor, email y filiación) no están en el documento, pero sí pueden encontrarse en otras páginas del mismo sitio web. La arquitectura propuesta explota esta característica para mejorar la automatización del proceso.

Algunos datos a extraer se buscan en el documento de texto y otros en páginas web vecinas del sitio. La entrada es una lista de URLs de sitios donde se desea realizar la búsqueda. La salida es una lista de documentos encontrados junto con la información extraída. Para manejar estos datos se optó por una base de datos orientada a grafos (BDoG) porque la estructura de los sitios web a crawlear puede ser representada por un grafo [5]. Las BDoG representan la información en un grafo, formados por nodos (entidades) y aristas (relaciones entre ellos) y permite usar la teoría de grafos para recorrerlo.

La arquitectura propuesta se presenta en la Figura 5.

En la arquitectura se cuenta con el módulo Coordinator que se encarga de coordinar al resto de los módulos. El Crawler es el módulo responsable de realizar las tareas específicas de web crawling.

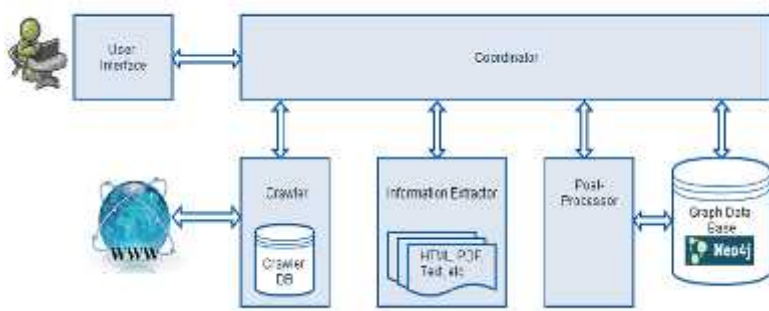


Figura 5: Arquitectura propuesta para la recopilación de documentos

El Extractor de Información extrae la información de los recursos recuperados por el crawler. Este módulo se compone de varios submódulos especializados: PDFBox se encarga de extraer el texto de un archivo PDF; ParsCit analiza la estructura y reconoce título, autores, emails y filiaciones; AlchemyAPI se encarga de extraer palabras claves y reconoce entidades organización y persona como potenciales valores de filiación y autor.

La base de datos orientada a grafos (BDoG) genera una estructura de grafo donde

los nodos contienen las URLs recuperadas por el crawler, y las hojas contienen las URLs que no tienen otras URLs salientes o bien tienen los objetos de interés. Para cada nodo se almacena la URL padre y las URLs hijas. A modo de ejemplo se muestra en la Figura 6 la base de datos resultante del crawling de un sitio de un docente-investigador de la UNR.

El módulo Post-Procesador encuentra y vincula información adicional que no está en la hoja y se puede encontrar en un nodo no muy lejano. Este módulo recorre las hojas y realiza recorridos ascendentes en la jerarquía de la estructura del sitio web; por ejemplo para buscar email de contacto o posibles datos de autores y filiación.

Se implementó un prototipo en Java utilizando Crawler4j ([edu.uci.ics.crawler4j](http://edu.uci.ics.crawler4j)) y Neo4j (<https://neo4j.com/>) y se desarrollaron aplicaciones para las herramientas de extracción Apache PDFBox, Alchemy API y ParsCit.

Para evaluar el desempeño del prototipo se lo ejecutó en algunos sitios académicos correspondientes a la página de un docente-investigador y a la página de una materia de la Licenciatura en Ciencias de la Computación.

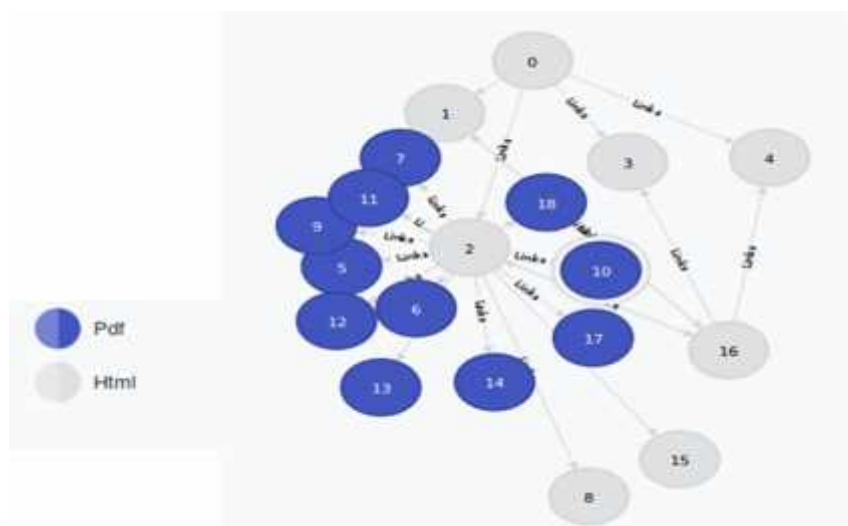


Figura 6: BDoG resultante del crawling de un sitio de un docente-investigador.

Como caso de uso, se describe la experimentación para la página web de un docente-investigador. A partir de una URL inicial se recuperaron todos los archivos en PDF mediante el Crawler y se generó la base de datos del sitio analizado. Utilizando ParsCit como un primer extractor de información (EI1), se alcanzaron los siguientes valores de Precisión y Cobertura: la extracción de los Títulos es óptima obteniendo un 100% tanto en Precisión como en Cobertura; para los Autores se obtuvo un 90% en Precisión y un 70% en Cobertura; y para las Filiaciones se obtuvieron un 85% en Precisión y un 73.3% en Cobertura.

Además, se planteó un Post-Procesamiento para los Autores y las Filiaciones, vinculando la información extraída del documento con información que se extrae

del sitio web. Para esto, se aplican dos extractores más: EI2 y EI3. El extractor EI2 recorre las páginas web ascendiendo en la jerarquía para buscar entidades autor o filiación, utilizando AlchemyAPI. El objetivo es encontrar entidades dentro del documento que pudieron no ser reconocidas por los extractores que operan sólo con el texto del documento. El extractor EI3 busca en un documento entidades de tipo autor o filiación encontradas en otros documentos del sitio, porque según la estructura de los objetos a veces los extractores, tales como ParsCit, lo extraen en uno y no en otro. Al tratarse de documentos de un mismo sitio es probable que una misma entidad autor o filiación esté presente en más de un objeto. En ambos extractores para la búsqueda de entidades se utilizan técnicas de matcheo aproximado. Para Autor se utiliza el algoritmo de Smith-Waterman, y para Filiación se utiliza el algoritmo de Levenshtein-Demearu. Para completar este post-procesamiento se realizó la unión de los resultados filtrando valores duplicados.

Con la unión de los extractores ParsCit, EI2 y EI3, se obtuvo que la Cobertura tanto para Autor como para Filiación mejora respecto al mejor extractor individual sin perder Precisión, aumentando en 25.5% para Autores y en 22% para Filiaciones. Así, concluimos que la incorporación de datos que se encuentran en el sitio web permite mejorar la cobertura en la extracción.

También se realizó un post-procesamiento para la extracción de los emails. En este caso se calcula la unión de los resultados de MailList, que busca mails en la primer hoja del documento aplicando expresiones regulares al texto, de ParsCitMailList, que busca mails encontrados por ParsCit, y de RelMailList, que busca mails en nodos relacionados. Se aplica también el procedimiento UnionMailList que realiza la unión de las listas anteriores, eliminando duplicados. La precisión para los extractores de mail es óptima. La cobertura es del 60% pero se obtiene para cada documento al menos un mail de contacto.

Se puede concluir que un problema en la extracción es que a veces los datos requeridos no se encuentran en el documento, pero sí pueden encontrarse en otras páginas del mismo sitio. Por lo tanto, la extracción de los metadatos Autores, Filiación y Mail mejora con el procesamiento de páginas vecinas al documento en cuestión. Así, la arquitectura propuesta explota esta característica para mejorar la automatización de la extracción de información

## Conclusiones

Dada la heterogeneidad de estilos cognitivos y de aprendizaje mostrados por los alumnos, en los distintos momentos didácticos (lectura comprensiva, análisis, síntesis creativa, aplicación práctica, evaluación, etc.) los sistemas recomendadores brindan sugerencias a los profesores/estudiantes durante su aprendizaje para maximizar su desempeño, proveyendo la recuperación y presentación personalizada de los recursos de aprendizaje en base a las

necesidades, intereses, preferencias y gustos del usuario. Por esto, la recomendación de objetos de aprendizaje es importante para ayudar a la formación de los futuros profesionales debido a que permiten encontrar los recursos que mejor se ajustan al estilo de aprendizaje y perfil del usuario.

Se puede mejorar la usabilidad de los Repositorios mediante herramientas para ayudar en la Búsqueda, la Carga y la Recopilación de documentos. En este artículo se presentaron diversas cuestiones relacionadas con estos aspectos. Respecto a la Búsqueda, se planteó la utilización de sistemas recomendadores que brinden una respuesta más cercana a las características y preferencias del usuario. Este tipo de búsqueda se apoya en metadatos de los recursos y de los usuarios, que son muy importantes para la recuperación personalizada. Preparar los recursos con metadatos adecuados es laborioso, por lo tanto es importante desarrollar sistemas para la extracción automática de metadatos. Esta extracción automática permite, además, facilitar la Carga de recursos educativos. De esta forma, se puede aumentar la población de los documentos en repositorios acompañando el desarrollo de Repositorios Institucionales de Acceso Abierto. También, se propuso dar soporte a las tareas de Recopilación de documentos que realiza el administrador del repositorio con el objetivo de detectar documentos plausibles de ser cargados junto con sus metadatos de interés. En este caso, es importante obtener datos de contacto de autores para invitarlos a depositar sus documentos en los repositorios.

Actualmente, se está trabajando en la recomendación híbrida de objetos basada en el diálogo y en la crítica donde las preferencias se capturan en forma dinámica interactuando con el usuario y así mejorar los resultados de la recomendación [4]. Estos recomendadores generan un proceso de elicitación que parte de preferencias iniciales y que luego aumentan y/o se refinan a partir de la crítica que hacen los usuarios a ejemplos que el sistema va presentando. Estos sistemas han probado ser eficientes para la recomendación de bienes de alto costo o riesgo. En particular, la elección de un material educativo es una decisión que demanda tiempo y tiene consecuencias sobre un grupo importante de alumnos por lo que este tipo de sistemas es altamente aplicable.

## Referencias Bibliográficas

1. Casali A, Bender C, Deco C y Fontanarrosa S. Extracción Automática de Metadatos como Soporte para el Autoarchivo de Objetos Digitales en Repositorios. Revista Colombiana de Computación 2014; 15(2):135-160.
2. Casali A, Deco C, Bender C, Gerling V. Recommender System for personalized retrieval of Learning Objects. En: Santos OC, Boticario JG, Editores. Book of Educational Recommender Systems and Technologies: Practices and Challenges. Spain: aDeNu Research Group. UNED. 2011. p. 182-210.
3. Casali A, Deco C, Beltramone S. An Assistant to Populate Repositories: Gathering Educational Digital Objects and Metadata Extraction. IEEE Journal of Latin-American Learning Technologies 2016 Mayo; 11(2):1-8.
4. Giustozzi F, Casali A, Deco C, Lemos dos Santos H, Cechinel C. Recommender System of Educational Resources: a Critiquing Based Proposal. LACLO 2016. Proceedings de la 11 Conferencia Latinoamericana de Objetos y Tecnologías de Aprendizaje; 2016 Oct 3-7; San Carlos, Costa Rica; 2016 p 1-8.
5. Robinson I, Webber J, Eifrem E. Graph Databases. USA: O'Reilly Media Inc.; 2013.
6. San Martín P, Bongiovani P, Casali A, Deco C. Study on Perspectives Regarding Deposit on Open Access Repositories in the Context of Public Universities in the Central-Eastern Region of Argentina. Scholarly an Research Communication 2015 February; 6(1):1-13.
7. Sonntag M. Metadata in E-Learning Applications: Automatic Extraction and Reuse. En Hofer C, Chroust G, Editores. IDIMT-2004. 12th Interdisciplinary Information Management Talks. Austria: Universitätsverlag Rudolf Trauner. 2004. p. 219-231.
8. Wiley D. Connecting Learning Objects to Instructional Design Theory: A definition, a metaphor, and a taxonomy. En: Wiley DA, Editor. Instructional Use of Learning Objects. Editorial Association for Instructional Technology. 2002.

## Autores

Dr.C. Claudia Deco

Docente/Investigador. Departamento de Investigación Institucional, Facultad de Química e Ingeniería del Rosario, Pontificia Universidad Católica Argentina. Departamento de Sistemas e Informática, Universidad Nacional de Rosario, Argentina.

MSc.Cristina Bender

Departamento de Investigación Institucional, Facultad de Química e Ingeniería del Rosario, Pontificia Universidad Católica Argentina. Departamento de Sistemas e Informática, Universidad Nacional de Rosario, Argentina.

