

Avances en la Ingeniería del Lenguaje y del Conocimiento

Research in Computing Science

Series Editorial Board

Editors-in-Chief:

Grigori Sidorov (Mexico)
Gerhard Ritter (USA)
Jean Serra (France)
Ulises Cortés (Spain)

Associate Editors:

Jesús Angulo (France)
Jihad El-Sana (Israel)
Alexander Gelbukh (Mexico)
Ioannis Kakadiaris (USA)
Petros Maragos (Greece)
Julian Padget (UK)
Mateo Valero (Spain)

Editorial Coordination:

María Fernanda Ríos Zacarias

Research in Computing Science es una publicación trimestral, de circulación internacional, editada por el Centro de Investigación en Computación del IPN, para dar a conocer los avances de investigación científica y desarrollo tecnológico de la comunidad científica internacional. **Volumen 125**, noviembre 2016. Tiraje: 500 ejemplares. *Certificado de Reserva de Derechos al Uso Exclusivo del Título* No.: 04-2005-121611550100-102, expedido por el Instituto Nacional de Derecho de Autor. *Certificado de Licitud de Título* No. 12897, *Certificado de licitud de Contenido* No. 10470, expedidos por la Comisión Calificadora de Publicaciones y Revistas Ilustradas. El contenido de los artículos es responsabilidad exclusiva de sus respectivos autores. Queda prohibida la reproducción total o parcial, por cualquier medio, sin el permiso expreso del editor, excepto para uso personal o de estudio haciendo cita explícita en la primera página de cada documento. Impreso en la Ciudad de México, en los Talleres Gráficos del IPN – Dirección de Publicaciones, Tres Guerras 27, Centro Histórico, México, D.F. Distribuida por el Centro de Investigación en Computación, Av. Juan de Dios Bátiz S/N, Esq. Av. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, C.P. 07738, México, D.F. Tel. 57 29 60 00, ext. 56571.

Editor responsable: *Grigori Sidorov, RFC SIGR651028L69*

Research in Computing Science is published by the Center for Computing Research of IPN. **Volume 125**, November 2016. Printing 500. The authors are responsible for the contents of their articles. All rights reserved. No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, recording or otherwise, without prior permission of Centre for Computing Research. Printed in Mexico City, in the IPN Graphic Workshop – Publication Office.

Avances en la Ingeniería del Lenguaje y del Conocimiento

**David Pinto
Darnes Vilariño
Mireya Tovar (eds.)**



Instituto Politécnico Nacional
"La Técnica al Servicio de la Patria"



Instituto Politécnico Nacional, Centro de Investigación en Computación
México 2016

ISSN: 1870-4069

Copyright © Instituto Politécnico Nacional 2016

Instituto Politécnico Nacional (IPN)
Centro de Investigación en Computación (CIC)
Av. Juan de Dios Bátiz s/n esq. M. Othón de Mendizábal
Unidad Profesional “Adolfo López Mateos”, Zacatenco
07738, México D.F., México

<http://www.rcs.cic.ipn.mx>

<http://www.ipn.mx>

<http://www.cic.ipn.mx>

The editors and the publisher of this journal have made their best effort in preparing this special issue, but make no warranty of any kind, expressed or implied, with regard to the information contained in this volume.

All rights reserved. No part of this publication may be reproduced, stored on a retrieval system or transmitted, in any form or by any means, including electronic, mechanical, photocopying, recording, or otherwise, without prior permission of the Instituto Politécnico Nacional, except for personal or classroom use provided that copies bear the full citation notice provided on the first page of each paper.

Indexed in LATINDEX, DBLP and Periodica

Printing: 500

Printed in Mexico

Editorial

Esta edición especial de la Revista Research in Computing Science contiene una serie de contribuciones originales que han sido seleccionadas a partir de un proceso de evaluación ciega doble (double blind), lo cual significa que los nombres de los autores de los artículos y los nombres de los revisores son ambos desconocidos. Este procedimiento es ejecutado en aras de proveer una evaluación anónima, que derive en artículos de mayor calidad para este volumen; particularmente, en esta ocasión la tasa de rechazo fue del 34%, cuidando que en todos los casos, al menos dos especialistas del comité revisor hicieran una evaluación de la pertinencia, originalidad y calidad de cada artículo sometido.

Las contribuciones presentes en este volumen son el resultado de una selección de los mejores artículos que fueron previamente presentados en el simposio en Ingeniería del Lenguaje y del Conocimiento (LKE'2016), en particular en la cuarta edición de esta serie de eventos. Esta conferencia ha sido organizada en el seno de la Facultad de Ciencias de la Computación de la Benemérita Universidad Autónoma de Puebla (BUAP) por cuatro años consecutivos, y nace como una iniciativa del laboratorio de Ingeniería del Lenguaje y del Conocimiento con la finalidad de ofrecer un espacio académico y de investigación, en el cual sea posible reportar trabajos relacionados con el área. Este evento promueve la cooperación entre diferentes grupos de investigación, pues permite el intercambio de resultados científicos, prácticos y la generación de nuevo conocimiento.

Esperamos que este volumen sea de utilidad para el lector y los autores de los artículos seleccionados encuentren en esta edición especial un espacio de intercambio científico productivo que enriquezca la colaboración entre estudiantes y académicos en el ámbito de la ingeniería del lenguaje y del conocimiento.

Deseamos agradecer a la Red Temática en Tecnologías del Lenguaje y a la Sociedad Mexicana de Inteligencia Artificial por los apoyos brindados.

El proceso de revisión y selección de artículos se llevó a cabo usando el sistema libremente disponible llamado EasyChair, <http://www.easychair.org>.

David Pinto
Darnes Vilariño
Mireya Tovar
Guest Editors

Facultad de Ciencias de la Computación,
LKE-FCC-BUAP, México

Noviembre 2016

Table of Contents

	Page
SEEG-BUAP: Un corpus de señales electro-encefalográficas con emociones inducidas	9
<i>Illiana Morales, Darnes Vilariño, Cristina Aguilar, Josefa Somodevilla</i>	
Diseño de un robot móvil controlado por un agente reactivo en tiempo real	17
<i>Griselda Saldaña-González, Jorge Cerezo-Sánchez, Mario M. Bustillo-Díaz, Apolonio Ata Pérez, Beatriz Bernabé Loranca, Gerardo Martínez Guzmán, Rogelio González Velázquez</i>	
Clasificación de polaridad con un diccionario mediante el algoritmo de Bayes	27
<i>Yuridiana Alemán, Darnes Vilariño, Josefa Somodevilla</i>	
Desarrollo de un sistema aumentativo y adaptativo para apoyar el proceso de lectoescritura para alumnos con discapacidad visual de educación superior.....	37
<i>Carmen Cerón, Etelvina Archundia, Beatriz Beltrán, Patricia Cervantes, José Galindo</i>	
Prototipo de una aplicación móvil accesible para apoyar a la localización de adultos mayores mediante un código compartido entre usuarios.....	51
<i>Etelvina Archundia, Carmen Cerón, Beatriz Beltrán, Patricia Cervantes, Fernando Rodríguez</i>	
Herramienta web para almacenar y visualizar objetos distribuidos	63
<i>José M. Hernández-Salinas, Carlos R. Jaimez-González</i>	
Sistema web personalizable con diseño adaptable para la administración de reservaciones en hoteles.....	75
<i>Betzabet García-Mendoza, Carlos R. Jaimez-González</i>	
Serialización de objetos PHP a XML	87
<i>Lidia N. Hernández-Piña, Carlos R. Jaimez-González</i>	
Sistema web para registro y administración de cursos de educación continua	97
<i>Carlos R. Jaimez-González, Wulfrano A. Luna-Ramírez</i>	

Hacia una herramienta web para visualizar la ejecución de programas escritos en el lenguaje Java	109
<i>Miguel Castillo-Cortes, Carlos R. Jaimez-González</i>	
Análisis del logro académico de estudiantes en el nivel medio superior a través de minería de datos centrada en el usuario	121
<i>Gabriel Maldonado, Guillermo Molero-Castillo, José Rojano-Cáceres, Alejandro Velázquez-Mena</i>	
Análisis basados en n-gramas para la clasificación del nivel de Burnout en académicos universitarios	135
<i>Mauricio Castro, David Pinto, Darnes Vilariño, Mireya Tovar</i>	
Una propuesta de clasificación automática de polaridad para notas periodísticas.....	145
<i>José A. Baez, Orlando Ramos, María J. Somodevilla, Ivo H. Pineda, Darnes Vilariño, Concepción Pérez de Celis</i>	
Búsqueda paralela exhaustiva aplicada a cadenas de ADN y ARNi	157
<i>Jesús García-Ramírez, Ivo H. Pineda T., María J. Somodevilla, Mario Rossainz, Concepción Pérez de Celis</i>	
Un modelo ontológico espacial para el tratamiento de desórdenes de conducta infantil	167
<i>María J. Somodevilla, Andrea M. P. Tamborrell, Ivo H. Pineda, Concepción Pérez de Celis</i>	

SEEG-BUAP: Un corpus de señales electro-encefalográficas con emociones inducidas

Illiana Morales, Darnes Vilariño, Cristina Aguilar, Josefa Somodevilla

Benémerita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, Puebla, México

{illiana.mrls.t,dvilarinoayala,mariajsomodevilla}@gmail.com
crisaguilar@hotmail.com
<http://www.lke.buap.mx/>

Resumen. En este artículo se presenta un corpus de señales electro-encefalo-gráficas etiquetado manualmente con emociones inducidas a través de videos y encabezados de noticias. El corpus, al que se le ha denominado SEEG-BUAP, fue construido en base a una población de 40 personas a quienes se le mostró una serie de videos que detonaban algún tipo de emoción del estilo “alegría”, “tristeza”, etc. Se espera que la colección de datos construida y puesta a disposición pública sirva como base para el desarrollo de modelos de clasificación automática basadas en emociones básicas detectadas a través de señales electro-encefalográficas.

Palabras clave: Corpus etiquetado, señales electro-encefalográficas, clasificación automática.

SEEG-BUAP: A Corpus of Electroencephalographic Signals with Induced Emotions

Abstract. In this paper we present a manually annotated corpus of electroencephalographic signals with emotions induced through videos and news headlines. The corpus, named SEEG-BUAP, has been constructed using 40 persons to whom a video has been shown, expecting to trigger some kind of emotion such as “happiness”, “sadness”, etc. We expect that the dataset constructed, will be freely available for further use in the construction of classification models for automatic detection of basic human emotions employing electroencephalographic signals.

Keywords. Manually annotated corpus, electroencephalographic signals, machine learning.

1. Introducción

El electroencefalograma (EEG) permite sondear la actividad neuronal, ya sea en entornos clínicos o de investigación. Médicamente, es una prueba estándar

para el diagnóstico de epilepsia, así como una serie de otras condiciones relacionadas con traumatismo y patología [10,13].

En el ámbito de la investigación, un EEG se utiliza para estudiar las respuestas neuronales ante ciertos estímulos externos como en [2], donde al desarrollar actividades como bailar y tocar guitarra, obtuvieron los datos de manera manual para un grupo de alumnos de Salsa en pareja. En otro trabajo [11], se interpretan las señales electroencefalográficas en la pronunciación del habla imaginada, en particular un vocabulario de 5 palabras, donde cada una de las palabras fue repetida 33 veces en 27 diferentes sujetos de prueba, sin embargo, en dicho trabajo no se menciona un corpus de prueba y/o entrenamiento disponible públicamente.

Para poder realizar tareas importantes relacionadas con la clasificación automática es necesario contar con conjuntos de datos de entrenamiento adecuados. En particular, cuando se diseña un corpus deben tomarse en cuenta distintos aspectos tales como la finalidad, los límites y el tipo de corpus [12]. En la actualidad es posible encontrar algunos corpora disponibles y relacionados con EEG, por ejemplo, conjunto de datos llamado “EEG Motor Movement/Imagery” cuenta con mas de 1500 registros de EEG de 1 y dos minutos, obtenidos de 109 voluntarios a los cuales se les pidió realizar ciertas tareas como: abrir y cerrar los puños con ojos abiertos, o abrir y cerrar los puños con los ojos cerrados [4]. Estos datos fueron capturados por 64 canales de registro EEG utilizando el sistema BCI200 [8]. Por otro lado, la base de datos CHB-MIT Scalp EEG Database fue obtenida en el Hospital de Niños de Boston, y se compone de registros EEG de personas pediátricas con convulsiones. Esta base de datos contiene información de 22 personas (5 hombres de edades entre 3-22, y 17 mujeres de edades entre 1.5-19) [9]. La base de datos EPILEPSIAE (una colección europea de información sobre epilepsia) contiene los registros EEG de 275 pacientes de los centros de epilepsia del Hospital Universitario de Friburgo, Alemania, del Hospital de la Universidad de Coimbra, Portugal, y del Hospital de la Pitier-Salpêtrière París, Francia [7]. Finalmente, el corpus HUT-EEG liberó recientemente 14 años de registros EEG clínicos recogidos en el Hospital de la Universidad de Temple. La colección consta de 16,986 sesiones de 10,874 sujetos únicos; los registros han sido organizados y se combinan con informes clínicos que describen los pacientes y las exploraciones [6].

Relativamente poca de esta información está disponible públicamente a la comunidad de investigación en una forma que sea útil para el aprendizaje automático [6], y la que se encuentra disponible está dirigida a estudios médicos. Hasta donde tenemos conocimiento, en la literatura no se cuenta con un corpus disponible públicamente de señales EEG con etiquetas de emociones. Es por eso, que en este artículo se describe la creación de un corpus de señales EEG con emociones inducidas mediante contenido multimedia.

El contenido restante de este artículo está estructurado como sigue. En la Sección 2. se presentan conceptos relacionados con la electro-encefalografía, en particular, los métodos de obtención de las señales EEG. En la sección 3. se presenta el proceso de adquisición de datos, describiendo no solamente como

se registran los datos sino también las emociones seleccionadas. Los resultados obtenidos se muestran y discuten en la Sección 4.. Finalmente, en la sección 5. se dan las conclusiones y trabajo a futuro.

2. Electroencefalografía

El electro-encefalograma es el registro de la actividad eléctrica que se genera en el cerebro por las neuronas en el interior del mismo; siendo ésta obtenida mediante electrodos superficiales en la base del cráneo [5]. A medida que el registro de la actividad eléctrica del cerebro se realiza mas cerca del cerebro, se puede obtener una señal mas limpia.

En este trabajo, la adquisicion de las señales EEG se realiza mediante la diadema de Emotiv EPOC+ [3], la cual ha sido principalmente diseñada para aplicaciones de investigación contextualizadas prácticas. Emotiv EPOC+ ofrece acceso datos de señales EEG con una calidad relativamente alta, utilizando software propietario con un SDK avanzado. Aunque también es posible acceder a esta información mediante software libre o lenguajes de programación como Python. Este dispositivo utiliza sensores superficiales que se ubican en el cuero cabelludo para leer la información electro-encefalográfica; así, el dispositivo cuenta con 14 canales de lectura, mas dos de referencia ubicados detrás de las orejas (ver Figura 1). Los sensores con los que cuenta la diadema de Emotiv Systems se colocan en la parte occipital, parietal y frontal de la cabeza; su nomenclatura indica la región de la misma donde están ubicados: Frontal (F), Central (C), Parietal (P), Occipital (O), Temporal (T) y Fronto-Parietal (FP). La diadema se complementa con un receptor USB, el cual permite que las señales adquiridas mediante los sensores sean enviadas a una computadora. Las señales electroencefalográficas se encuentran en el rango de frecuencias de 0 a 50 Hz y se pueden clasificar en cuatro bandas de frecuencia: Delta, Theta, Alpha y Beta.

En la siguiente sección se discute la metodología utilizada para la adquisición de los datos para la construcción del corpus SEEG-BUAP.

3. Adquisición de los datos

La construcción de un corpus de información etiquetada manualmente requiere indudablemente de una metodología que permita garantizar la calidad de los datos. Así, en esta sección se describe la manera en que el corpus fue creado.

3.1. Emociones seleccionadas

Para la confección del corpus SEEG-BUAP se tomaron en cuenta las seis emociones básicas propuestas por Ekman [1], puesto que estas emociones se encuentran presentes en todos los individuos y pueden ser detonadas a través de estímulos. A continuación se enumeran cada una de ellas.

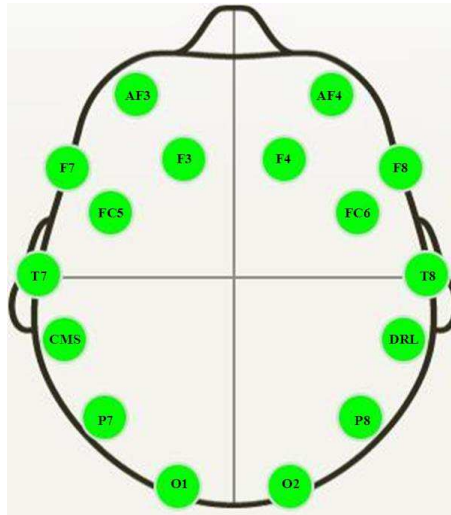


Fig. 1. Sensores de la diadema Emotiv Epoc+.

1. Alegría,
2. Tristeza,
3. Ira,
4. Miedo,
5. Asco,
6. Sorpresa.

Disparar dichas emociones en los seres humanos no es una tarea sencilla, sin embargo, nuestra metodología plantea inducir la emoción a través de videos.

3.2. Etiquetas de los videos

En una primera captura de datos se seleccionaron 3 videos que detonaran la emoción de “Alegría”. Se escogieron videos de comediantes nacionales e internacionales con una duración entre 3 y 5 minutos.

Los videos fueron manualmente etiquetados como se muestra en la Tabla 1. En aquellos intervalos de segundos donde presumiblemente se detonaba cierta emoción, se utilizó la etiqueta correspondiente. En el caso de no detonarse ninguna de las seis emociones básicas, la etiqueta para ese intervalo de segundos es “Ninguna”. Con ayuda de un sistema computacional desarrollado durante este proyecto fue posible realizar captura de imágenes del rostro del sujeto de prueba para comprobar que la emoción haya estado presente, quedando la captura de datos como se muestra en la Figura 2.

Una vez etiquetado cada video con el lapso de tiempo en donde una posible emoción debía ser inducida a una persona. El siguiente paso es adquirir las señales electro-encefalográficas y construir el corpus. La siguiente sección describe este proceso.

Tabla 1. Ejemplo de etiquetas en los videos.

Miedo		
Segundo inicial	Segundo final	Etiqueta
1	52	"Ninguna"
53	73	"Miedo"
74	79	"Ninguna"
80	106	"Miedo"
107	127	"Ninguna"



Fig. 2. Reproducción de video y adquisición de datos.

3.3. Sistema de adquisición de señales

Para la adquisición automática de los datos se desarrolló un sistema computacional que permite, utilizando dos procesos concurrentes, reproducir los videos y capturar los datos etiquetados automáticamente. Este sistema se desarrolló en Python, utilizando la API que Emotiv Systems proporciona a los desarrolladores. Al utilizar la API, el sistema permite recuperar las señales en bruto de los 16 sensores de la diadema, con una frecuencia de muestreo de 128 muestras por segundo.

Los datos obtenidos son almacenados en un formato *csv* (como se puede ver en la Tabla 2). Cada fila contiene el segundo en que fue capturado, el número

de muestra que representa (de un máximo de 128), el valor asociado al sensor y la etiqueta de la emoción.

Tabla 2. Ejemplo de datos recuperados (datos redondeados a cero decimales).

Segundo	Muestra	AF3	F7	F3	FC5	T7	P7	...	AF4	Emoción
37	126	4162	4152	4160	4182	4189	4173	...	4166	Alegría
37	127	4161	4150	4162	4180	4187	4170	...	4164	Alegría
37	128	4168	4153	4173	4183	4186	4175	...	4171	Alegría
38	1	4174	4175	4207	4198	4187	3329	...	4178	Ninguna
38	...	...	...	...	...	...	...	...	...	Ninguna
38	128	4237	4293	4242	4264	4179	4182	...	4185	Ninguna

El proceso de adquisición de datos tuvo que ser afinado constantemente ya que, por ejemplo, cuando inicialmente se realizó una captura de datos en un salón de clases a 11 sujetos de prueba, se pudo concluir que es necesario disminuir distracciones. En la Figura 3 se puede observar a un sujeto de prueba con distracciones que evidentemente alteran el resultado del experimento.



Fig. 3. Sujeto de prueba distraído.

En la captura final se citó a cada sujeto de prueba a acudir a un laboratorio y se le situó dentro de una cabina blanca, colocándole además audífonos para no generar distracciones. Se aumentó la lista reproducción de videos a 6, considerando cada una de las emociones básicas, obteniendo un total de 36 videos que fueron previamente etiquetados por anotadores humanos. De estos 36 videos se seleccionaron 1 de cada emoción para formar la lista de reproducción que vería el sujeto de prueba.

4. Constitución del corpus SEEG-BUAP

En total se realizaron los experimento con un grupo de 26 personas (16 hombres y 10 mujeres), con un rango de edad de entre 20 a 23 años, estudiantes universitarios con perfil de Computación en las siguientes condiciones: descansados y sin presión.

4.1. Filtrado de datos

En total se obtuvieron más de 1,100,000 muestras a partir de las 26 pruebas realizadas. Se descartaron 8 pruebas debido a inconvenientes con el hardware utilizado (dispositivo sin batería suficiente, problemas de conexión, etc.), y a que en algunas ocasiones la emoción no se detonó como se esperaba, lo cual pudo ser comprobado con las imágenes que se sacaron como apoyo a la creación de este corpus. Las pruebas utilizables son las que forman parte del corpus SEEG-BUAP. El corpus final cuenta con más de 650,000 muestras etiquetadas manualmente con la emoción correspondiente¹.

5. Conclusiones y trabajo a futuro

En este trabajo se presenta un corpus de señales EEG con etiquetado manual de emociones inducidas a través de videos. Para la correcta creación del corpus SEEG-BUAP, fueron necesarias varias pruebas, en cada una de ellas se pudo mejorar la calidad de los datos. El corpus puesto a disposición pública cuenta con mas de 650 mil muestras etiquetadas de acuerdo a una de seis posibles emociones básicas.

Como trabajo a futuro se busca utilizar el corpus creado para proponer un modelo de clasificación de emociones con señales EEG. El procesamiento de las señales EEG debe ser llevado a cabo cuidadosamente debido a que hasta el momento se cuenta únicamente con los datos planos (*raw*), es decir, tal cual se ha leído desde el dispositivo. Dado que los datos obtenidos son valores reales, se considera que un clasificador basado en una mezcla de gaussianas sería el que obtendría los mejores resultados, sin embargo, no se descarta el uso de otro tipo de clasificadores, como por ejemplo, máquinas de soporte vectorial. También se espera utilizar el corpus de imágenes para la creación de una base de conocimientos para retroalimentar facialmente un robot humanoide.

Referencias

1. Ekman, P., Friesen, W.V., O'Sullivan, M., Chan, A., Diacoyanni-Tarlatzis, I., Heider, K., Krause, R., LeCompte, W.A., Pitcairn, T., Ricci-Bitti, P.E., Scherer, K., Tomita, M., Tzavaras, A.: Universals and cultural differences in the judgments of facial expressions of emotion. *Journal of Personality and Social Psychology* 53(4), 712–717 (1987)

¹ El corpus puede ser descargado gratuitamente de <http://lke.cs.buap.mx/seeg-buap>

2. Elizondo, J.J.E., Beristain, L.J., Martínez, R.A.R., Perales, A.C., Quevedo, E.M., Sandoval, J.A.R.: Metodología para el análisis de señales encefalográficas en actividades lúdicas. *Academia Journal* 6(5), 1348–1353 (2014)
3. Emotiv E.P.O.C.: Software Development Kit (2010)
4. Goldberger, A.L., Amaral, L.A.N., Glass, L., Hausdorff, J.M., Ivanov, P.C., Mark, R.G., Mietus, J.E., Moody, G.B., Peng, C.K., Stanley, H.E.: Physiobank, physio-toolkit, and physionet. *Circulation* 101(23), e215–e220 (2000)
5. Guevara Mosquera, S.D.: Adquisición de señales electroencefalográficas para el movimiento de un prototipo de silla de ruedas en un sistema BCI. Master's thesis, Universidad Politécnica Salesiana, España (2012)
6. Harati, A., Choi, S., Tabrizi, M., Obeid, I., Picone, J., Jacobson, M.P.: The temple university hospital eeg corpus. In: *Global Conference on Signal and Information Processing (GlobalSIP)*, 2013 IEEE. pp. 29–32 (Dec 2013)
7. Ihle, M., Feldwisch-Drentrup, H., Teixeira, C.A., Witon, A., Schelter, B., Timmer, J., Schulze-Bonhage, A.: Epilepsiae – a European epilepsy database. *Comput. Methods Prog. Biomed.* 106(3), 127–138 (Jun 2012)
8. Schalk, G., McFarland, D.J., Hinterberger, T., Birbaumer, N., Wolpaw, J.R.: Bci2000: A general-purpose brain-computer interface (bci) system. *IEEE Transaction on Biomedical Engineering* 51(6), 2004 (2004)
9. Shoeb, A.H.: Application of machine learning to epileptic seizure onset detection and treatment. Ph.D. thesis, Massachusetts Institute of Technology (2009)
10. Tatum, W.: *Handbook of EEG Interpretation*. Springer Demos Medic Series, Springer Publishing Company (2007)
11. Torres-García, A., Reyes-García, C., Villaseñor Pineda, L., Ramírez-Cortés, J.: Análisis de señales electroencefalográficas para la clasificación de habla imaginada. *Revista mexicana de ingeniería biomédica* 34, 23–39 (2013)
12. Torruella, J., Llisterri, J.: Diseño de corpus textuales y orales, chap. 2, pp. 45–77. *Filología e informática. Nuevas tecnologías en los estudios filológicos* (1999)
13. Yamada, T., Meng, E.: *Practical Guide for Clinical Neurophysiologic Testing: EEG*. M - Medicine Series, Lippincott Williams & Wilkins (2009)

Diseño de un robot móvil controlado por un agente reactivo en tiempo real

Griselda Saldaña ¹, Jorge Cerezo¹, Mario Mauricio Bustillo², Apolonio Ata ²,
Beatriz Bernabé², Gerardo Martínez², Rogelio González²

¹ Universidad Tecnológica de Puebla,
Ingeniería en Tecnologías para la Automatización, México

² Benemérita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

{griselda.saldana, jorge.cerezo}@utpuebla.edu.mx,
bustillo1956@hotmail.com,
{apolonio.ata, beatriz.bernabe}@gmail.com, {gguzman,
rgonzalez}@cs.buap.mx

Resumen. En este trabajo se presenta el desarrollo de un sistema de detección de objetos en un espacio bidimensional controlado, empleando técnicas de visión artificial. Los objetos detectados tienen una geometría rígida y están expuestos a iluminación real, por lo tanto, el sistema es robusto a los cambios de iluminación y sombreado. Con el fin de manejar la gran cantidad de datos a procesar en tiempo real, se utiliza un dispositivo MyRIO el cual contiene un FPGA. Dicho dispositivo permite la comunicación con el software LabVIEW donde reside la interfaz de usuario. Empleando LabVIEW se implementa un algoritmo de seguimiento por color, con el fin de asistir a un agente reactivo, que utiliza un sensor infrarrojo para detectar la distancia a un obstáculo y realizar las funciones de forrajeo y almacenamiento.

Palabras clave: Seguimiento de objetos, LabVIEW, detección por color, robot móvil.

Design of a Mobile Robot Controlled by a Real-Time Agent

Abstract. In this paper the development of an object detection system in a controlled two-dimensional space using computer vision techniques is presented. The detected objects have a rigid geometry and are exposed to real light; therefore, the system is robust to changes in lighting and shading. In order to handle the large amount of data to be processed in real time, a MyRIO device which contains an FPGA is used. This device allows communication with the

LabVIEW software where the user interface resides. Using LabVIEW a tracking by color algorithm is implemented, in order to attend to a reactive agent, which uses an infrared sensor to detect the distance to an obstacle and perform the functions of foraging and storage.

Keywords. Tracking, LabVIEW, color detection, mobile robot.

1. Introducción

Actualmente la visión juega un papel central en el área de robótica para una serie de tareas como la auto-localización, la navegación, el reconocimiento y manipulación de objetos, el seguimiento de objetivos, la interacción social entre humano y robot, imitación, entre otros [1]. En los últimos años, los robots móviles han participado en tareas más y más complejas que a menudo requieren la colaboración entre varios individuos que, en general, difieren en sus habilidades y en su forma de percibir el entorno externo [2]. Para llevar a cabo dichas tareas, los robots móviles suelen estar equipados con potentes sistemas de visión, donde la operación en tiempo real se convierte en un desafío. Por otra parte, cuando se considera un sistema de múltiples robots, el control y la coordinación son tareas complejas y demandantes.

De acuerdo con [3] los sistemas de visión para agentes móviles se pueden clasificar en seis tipos, sensado y precepción [4-8]; mapeo y auto-localización [9, 10, 11]; reconocimiento y localización [12, 13, 14, 15, 16]; navegación y planeación [17, 18]; seguimiento [19, 20, 21]; y control por visión (servoing) [22, 23, 24].

En este trabajo, se desarrolla un sistema de visión montado en un robot móvil, el cual es capaz de navegar en un ambiente controlado empleando una cámara. El sistema está inspirado en el fenómeno de división de tareas, que se lleva a cabo en algunas colonias de insectos para realizar una tarea compleja común a un grupo como es el caso del forrajeo [25]. La tarea del robot consiste en observar su entorno para identificar un objeto de interés de un color específico, navegar hacia él y asirlo para transportarlo a un área de almacenaje. El agente controla de forma independiente una cámara de manera que puede dirigir su atención a diferentes objetos del entorno. Por otro lado, el robot cuenta con sus propios recursos de cómputo que le permiten analizar la imagen al emplear un algoritmo de seguimiento en base a color. Para procesar las imágenes en tiempo real se utiliza una tarjeta MyRIO 1900 que incluye un FPGA, el cual permite procesamiento en paralelo. El robot cuenta además con los recursos necesarios para su comunicación con otros agentes vía una red Wi-Fi. El sistema se programó empleando LabVIEW, que es un entorno de desarrollo para medición y automatización, desarrollado por National Instruments (NI) [26, 27], el cual cuenta con un conjunto de herramientas de Visión [28] y de gestión de gráficos e interfaces de usuario.

El documento está organizado de la siguiente manera, en la sección 2 se presenta una descripción del sistema, destacando la construcción del agente así como su programación en LabVIEW, en la sección 3 se muestran las pruebas realizadas al sistema así como algunos resultados. Finalmente, en la sección 4 se presentan las conclusiones y trabajo futuro.

2. Descripción del sistema

La configuración del sistema está compuesta por un robot móvil tipo VEX Clawbot [29], el cual tiene montada en su estructura una tarjeta MyRIO que se comunica de forma inalámbrica con una computadora central donde reside la interface de usuario. En dicha tarjeta se encuentra conectada una cámara de alta definición montada en un servomotor que le permite movimiento en la dirección horizontal (paneo). El robot cuenta además con dos servomotores para su locomoción y uno más para apertura/cierre de la pinza. En la parte frontal se encuentra un sensor infrarrojo que le permite determinar la distancia a los objetos de interés.

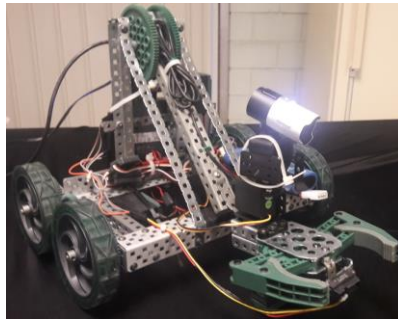


Fig. 1. Prototipo del Robot Móvil.

La tarjeta MyRIO cuenta con un Router Inalámbrico que gestiona las comunicaciones vía WiFi, utilizando variables globales mapeadas a la red que permiten utilizar la información disponible. Adicionalmente esta tarjeta permite la operación STAND-ALONE del sistema ya que el programa desarrollado se ejecuta dentro de los módulos de la misma.

2.1. Seguimiento en base a color

Para realizar el seguimiento de los objetos detectados, las imágenes obtenidas se convierten al formato HSI y a continuación se determina la escala para filtrado del color deseado, verde en este caso.

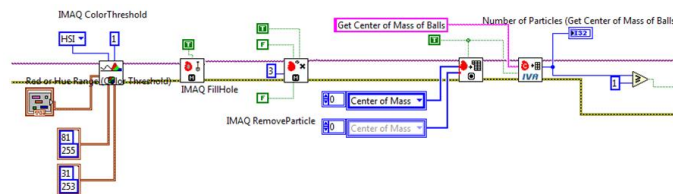


Fig. 2. Procesamiento de la Imagen para seguimiento por color.

Dentro de los límites de la imagen, el sistema de visión discrimina toda la gama de colores que no correspondan con el filtro programado (se elimina el fondo), dejando visibles solamente los elementos de color especificados. En la figura 3 se observan dos pelotas detectadas, donde se discrimina el resto de la imagen.

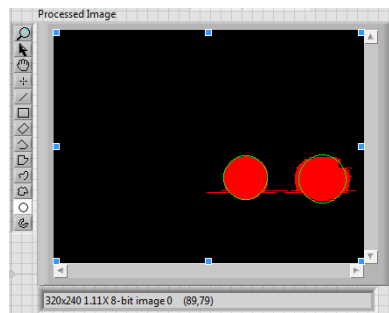


Fig. 3. Seguimiento de objetos de color verde, (el segundo en movimiento).

2.2. Interface gráfica de usuario

La interface desarrollada cuenta con un alto grado de usabilidad y funcionalidad. Cuenta con elementos que permiten configurar los parámetros para la máquina de visión, el ancho de pulso para los servomotores del robot, permite modificar en tiempo de ejecución, seleccionar la búsqueda de otra gama de color, así como visualizar la adquisición en tiempo real de la imagen obtenida por la cámara y la imagen procesada. Es posible observar el número de objetos detectados, así como las coordenadas (x , y) de sus centros de masa. También cuenta con dos indicadores de la actividad del agente, uno para el punto en que el robot se encuentra alineado con el objeto de interés (en rango) y otro para el momento en que la pinza puede tomar el objeto (pre-ajuste). La figura 4 muestra el aspecto de la interface de usuario completa.

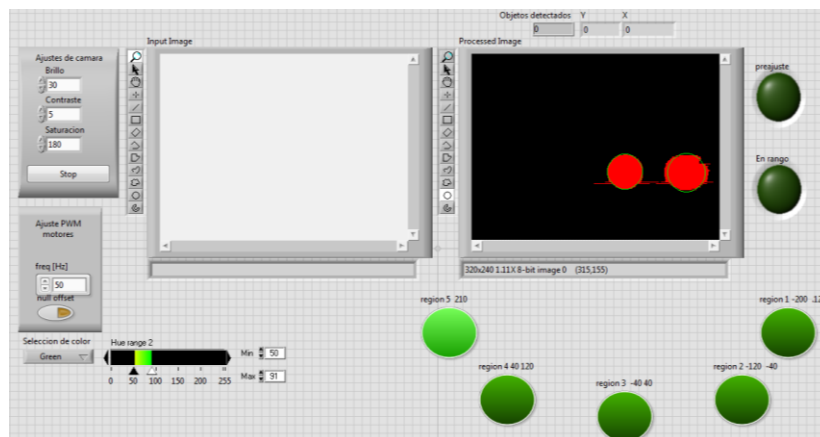


Fig. 4. Interface de usuario.

2.3. Diseño del agente reactivo

Para que el robot pueda interactuar con el entorno, se diseñó una máquina de estados que se inspira en un agente reactivo [30], que da solución a un problema local, dentro de una cancha de pruebas con un área específica de 1.68×2.24 metros bajo un ambiente de luz real; se considera además un rango de visión para la cámara de 0 a 180 grados. El sistema está diseñado modularmente para agregar nuevas acciones, patrones de comportamiento o agentes adicionales.

Durante la percepción, la cámara realiza un proceso de paneo, que es interrumpido cuando existe un objeto de interés en la pantalla y el centro del mismo se encuentra alineado con el eje X de la cámara.

Una vez localizado el centro del objeto, el plan programado para el agente consiste en determinar en cuál de las 5 zonas de visión posibles en que se ha dividido la imagen se encuentra el objeto. A continuación, el agente define las acciones de control para que el móvil alinee la cámara con el objeto a sujetar.

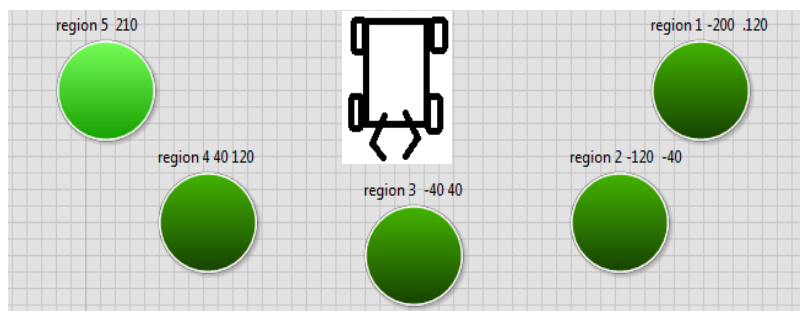


Fig. 5. Regiones del campo de visión del robot en grados.

Cuando la cámara montada en el robot se encuentra en la zona central (zona 3), se procede a realizar un movimiento fino de ajuste para centrar el objeto con la pinza y determinar qué tan lejos se encuentra el móvil del objeto de interés empleando el sensor infrarrojo para finalmente avanzar y sujetarlo.

Es importante recalcar que si por efectos de la mecánica de locomoción del robot o movimiento del objeto de interés, el móvil y la pelota se desalinean, el sistema de percepción del agente permite corregir en tiempo de ejecución la posición del robot y realizar nuevamente un paneo y acciones de alineación en dependencia de los cambios del entorno.

En la figura 6 se muestra la implementación del agente en una máquina de estados desarrollada en LabVIEW.

Este comportamiento realizado por el agente, corresponde a la función de forrajeo, ya que al asir el objeto de interés, el robot regresa a un punto de origen, donde recopila los objetos para que posteriormente otro agente se encargue de llevar a cabo el almacenamiento.

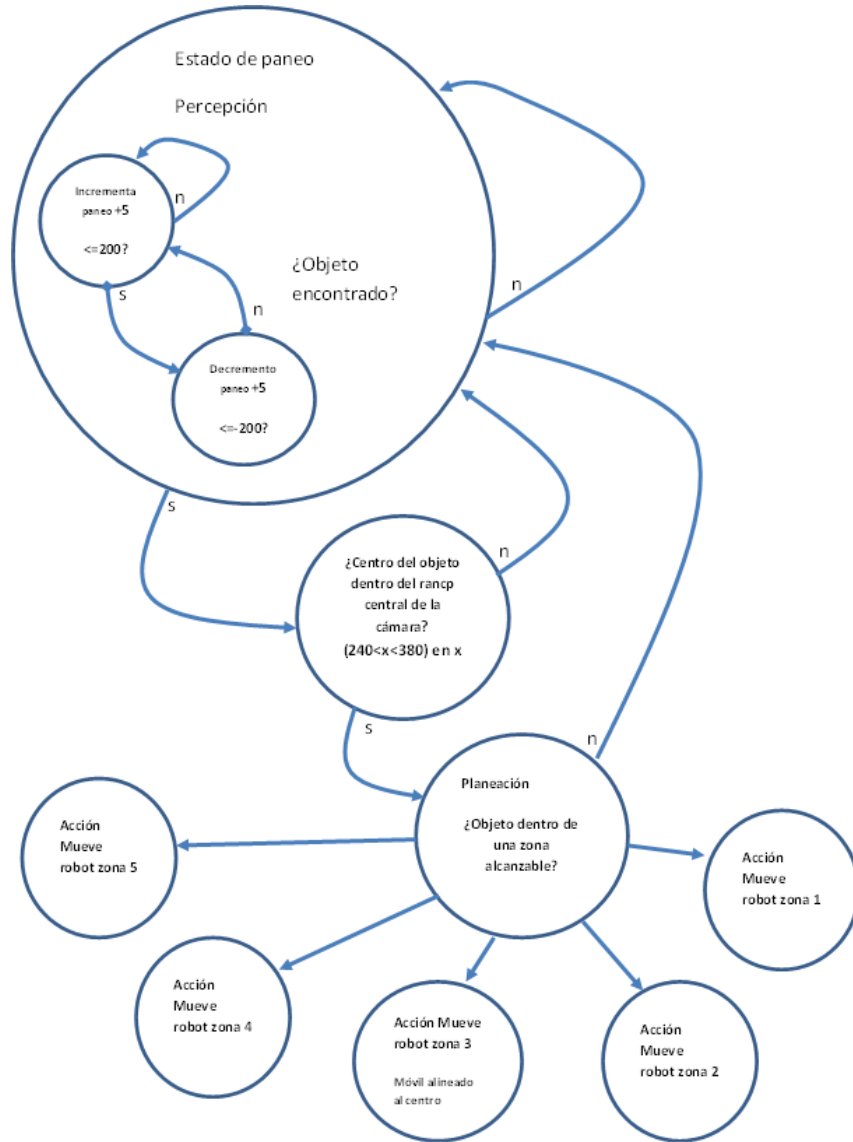


Fig. 6. Diseño del agente reactivo para el sistema de visión.

3. Pruebas y resultados

Se realizaron múltiples experimentos considerando el espacio que abarca la cancha de pruebas.

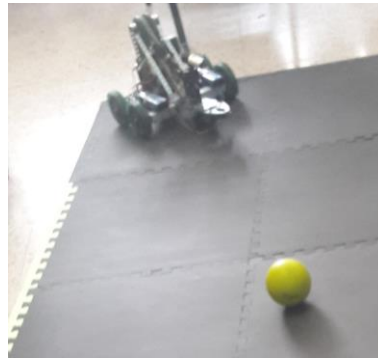


Fig. 7. Cancha de pruebas.

En este espacio se colocaron el robot y una pelota de color verde, y se verificó que desde cualquier punto el robot pudiera identificar la pelota, alinearse, tomarla y dejarla en una posición especificada. El sistema se probó bajo luz artificial y se comprobó que es robusto ante ruido y cambios de iluminación y sombreado.

Para el correcto funcionamiento del sistema, un paso muy importante fue la calibración del color, ya que dependiendo de los pesos que se dan a la saturación y la tonalidad (HSI) la discriminación puede o no mejorar. El valor de la luminosidad hace que el sistema sea más o menos inmune a los cambios de luz del ambiente.

Previamente a la utilización de este algoritmo, se probaron, *Traditional mean shift* y *Shape adapted mean*, en ambos casos se intentó detectar 4 pelotas de color verde tomando como parámetros Brillo: 30, Contraste: 5, Saturación: 180 y se obtuvieron los resultados mostrados en la Tabla 1.

Tabla 1. Objetos detectados por algoritmo.

Algoritmo	Objetos detectados con luz controlada	Objetos detectados bajo condiciones variables de luz
Traditional mean shift	3 a 4 pelotas	0 a 3 pelotas
Shape adapted mean	3 a 4 pelotas	0 a 3 pelotas
Basado en calibración de color	4 pelotas	4 pelotas

Para cada algoritmo se realizaron aproximadamente 50 pruebas con condiciones de luz controladas y a diferentes distancias del robot, obteniendo un 97% de detección para el algoritmo propuesto. Para el caso en que las condiciones de luz son variables, también se realizaron 50 pruebas para cada algoritmo y en este caso se obtuvo un 93% de detección con el algoritmo propuesto, que supera a los probados anteriormente cuyo porcentaje varía entre 86% y 90% de detección.

Para hacer que el robot se alineara con el objeto de interés, con la finalidad de asirlo, se decidió dividir el campo de visión de la cámara de paneo en 5 zonas, para proceder a hacer un ajuste fino de la pinza. Esto como resultado de múltiples pruebas

de los movimientos del robot para lograr que reaccionara a modificaciones del entorno.

Para conseguir que el robot pudiera asir los objetos de interés, se realizó un ajuste “fino” de la pinza por lo que fue necesario regular la velocidad de avance del robot, manipulando los anchos de pulso para los motores.

Con la misma finalidad, se calibró el telémetro (infrarrojo) para ajustarlo a una medida en unidades del mundo real.

A pesar de que el sistema es robusto, aún no se ha solucionado el caso en que exista oclusión de los objetos. Cada pelota dentro del sistema de visión recibe una etiqueta. Si dos pelotas dentro del área de visión del robot se cruzan, puede ocurrir que las etiquetas se inviertan. Si este caso se presenta, el robot aun será capaz de asir alguna de las dos y para el comportamiento deseado no es relevante cual tome primero. Por otro lado si dos pelotas se ubican una al lado de la otra el sistema de visión puede reconocerlas como un solo objeto, lo que no afecta en la alineación del robot, pero si en el momento de sujetar el objeto. Si algún objeto se ubica detrás del robot, el sistema de visión continuaría de manera indefinida en el estado de paneo.

4. Conclusiones

En este trabajo se ha demostrado que es posible utilizar un robot que pueda realizar una tarea común a un grupo tal como lo es el forrajeo. El uso del entorno LabVIEW permitió reducir el tiempo de implementación y realizar el análisis en tiempo real del sistema de visión y del agente. El agente implementado ejecuta las acciones de control del robot y es capaz de hacer correcciones en tiempo de ejecución. El sistema implementado es modular y configurable. Finalmente, a diferencia de otros trabajos similares, este sistema tiene la ventaja de que permite la comunicación entre agentes o hacia otros dispositivos a través de Wi-Fi gracias a la utilización del dispositivo MyRIO.

Como trabajo futuro se propone implementar un segundo robot que se encargue de la función de almacenamiento y hacer pruebas de comunicación entre agentes. Por otro lado, se implementará un algoritmo de evasión de obstáculos que permita realizar una tarea de mayor complejidad común al conjunto de agentes. Finalmente, para eliminar los problemas de oclusión, se utilizará un dispositivo Kinect que proporcione una vista superior del área de pruebas para identificar en todo momento la posición de los objetos de interés incluyendo el caso en que dichos objetos se encuentren detrás del robot.

Agradecimientos. Los autores agradecen al PRODEP por el apoyo para la realización de este trabajo.

Referencias

1. Fiack, L., Cuperlier, N., Miramond, B. J.: Embedded and real-time architecture for bio-inspired vision-based robot navigation. *Journal of Real-Time Image Processing*, 10 (4), 699–722 (2015)

2. Infantino, I., Cossentino, M., Chella, A.: An Agent Based Multilevel Architecture for robotics vision systems. *Proceedings of the International Conference on Artificial Intelligence*, 386–390 (2002)
3. Martínez-Gómez, J. et. al.: A taxonomy of Vision Systems for Ground Mobile Robots. *International Journal of Advanced Robotic Systems*, 11 (7) (2004)
4. Zhao, Z., Weng, Y.: A flexible method combining camera calibration and hand-eye calibration. *Robótica*, 31 (5), 747–756 (2013)
5. Peynot, T., Scheduling, S. and Terho, S.: The marulan data sets: Multi-sensor perception in a natural environment with challenging conditions. *International Journal of Robotics Research*, 29 (13), 1602–1607 (2010)
6. Fernández-Caballero, A., Castillo, J. C., Martínez-Cantos, J., Martínez-Tomás, R.: Optical flow or image subtraction in human detection from infrared camera on mobile robot. *Robotics and Autonomous Systems*, 58 (12), 1273–1281 (2010)
7. Ekmanis, M., Nikitenko, A.: Mobile Robot Camera Extrinsic Parameters Auto Calibration by Spiral Motion. *Engineering for Rural Development Jelgava*, 558–565 (2016)
8. Di Paola, D., et al.: Multi-sensor surveillance of indoor environments by an autonomous mobile robot. *Int. J. Intelligent Systems Technologies and Applications*, 8, 18–35 (2010)
9. Ahmmed, P., et al.: Self-localization of a mobile robot using monocular vision of a chessboard pattern. *International Conference on Electrical and Computer Engineering*, (2014)
10. Biswas, J., Veloso, M.: Depth camera based indoor mobile robot localization and navigation. *International Conference on Robotics and Automation*, 1697–1702 (2012)
11. Çelik, K., Somani, A. K.: Monocular vision SLAM for indoor aerial vehicles. *Journal of Electrical and Computer Engineering* (2013)
12. Cocias, T. T., Moldoveanu, F., Grigorescu, S. M.: Generic fitted shapes (gfs): Volumetric object segmentation in service robotics. *Robotics and Autonomous Systems*, 61 (9), 960–972 (2013)
13. Xiao, J., Adler, B., Zhang, J., Zhang, H.: Planar segment based three-dimensional point cloud registration in outdoor environments. *Journal of Field Robotics*, 30 (4), 552–582 (2013)
14. Yeom S., Woo, Y.-H.: Person-specific face detection in a scene with optimum composite filtering and colour-shape information. *International Journal of Advanced Robotic Systems*, 10 (70) (2013)
15. Zhang, J., Zhang, J., Chen, S.: Discover novel visual categories from dynamic hierarchies using multimodal attributes. *IEEE Transactions on Industrial Informatics*, 9 (3), 1688–1696 (2013)
16. Kawasue, K., Komatsu, T.: Shape measurement of a sewer pipe using a mobile robot with computer vision. *International Journal of Advanced Robotic Systems*, 10 (52) (2013)
17. Kang, T. K., Choi, I.-H., Park, G.-T., Lim, M.-T.: Local environment recognition system using modified surf-based 3d panoramic environment map for obstacle avoidance of a humanoid robot. *International Journal of Advanced Robotic Systems*, 10 (275) (2013)
18. Papadakis, P.: Terrain traversability analysis methods for unmanned ground vehicles: A survey. *Engineering Applications of Artificial Intelligence*, 26 (4), 1373–1385 (2013)

19. Park, J., Hwang, W., Kwon, H., Kim, K., Cho, D.-I.D.: A novel line of sight control system for a robot vision tracking system, using vision feedback and motion-disturbance feedforward compensation. *Robotica*, 31 (1), 99–112 (2013)
20. Ostafew, C. J., Schoellig, A. P., Barfoot, T. D.: Learning-Based Nonlinear Model Predictive Control to Improve Vision-Based Mobile Robot Path-Tracking in Challenging Outdoor Environments. *International Conference on Robotics & Automation*, 4029–4036 (2014)
21. Basso, F., Munaro, M., Michieletto, S., Pagello, E., Menegatti, E.: Fast and robust multi-people tracking from RGB-D data for a mobile robot. *Intelligent Autonomous Systems*, 265–276 (2013)
22. Siradjuddin, I., Behera, L., McGinnity, T. M., Coleman, S.: Image-Based Visual Servoing of a 7-DOF Robot Manipulator Using an Adaptive Distributed Fuzzy PD Controller. *IEEE/ASME Transactions on Mechatronics*, 19 (2) (2014)
23. Fomena, R. T., et al.: Coarsely Calibrated Visual Serving of a Mobile Robot using a Catadioptric Vision System. *IEEE/RSJ International Conference on Intelligent Robots and Systems*, 5432–5437 (2009)
24. Wang, K., Liu, Y., Li, L.: Visual Servoing Trajectory Tracking of Nonholonomic Mobile Robots Without Direct Position Measurement. *IEEE Transactions on Robotics*, 30 (4), 1026–1035 (2014)
25. Dorigo, M., Birattari, M. and Stützle, T.: Ant Colony Optimization: Artificial Ants as a Computational Intelligence Technique. *IRIDIA - Technical Report Series: TR/IRIDIA/2006-023*, 1–12 (2006)
26. Moreno Astorga, A., Moreno-Salinas, D., Chaos García, D., Aranda Almansa, J.: Simulation benchmark for autonomous marine vehicles in LabView. In: *Proceeding of OCEANS Conference*, 1–6 (2011)
27. De Asmundis R.: The Importance of a Deep Knowledge of LabVIEW Environment and Techniques in Order to Develop Effective Applications. *Practical Applications and Solutions Using LabVIEW™ Software*, Dr. Silviu Folea (Ed.), InTech (2011)
28. Relf, C. G.: *Image Acquisition and Processing with LabVIEW*. CRC Press (2004)
29. VEX Robotics, <http://www.vexrobotics.com.mx/>
30. Brooks, R. A.: A Robust Layered Control System for a Mobile Robot. *IEEE Journal on Robotics and Automation*. 2 (1), 14–23(1986)

Clasificación de polaridad con un diccionario mediante el algoritmo de Bayes

Yuridiana Alemán, Darnes Vilariño, Josefa Somodevilla

Benemérita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

yuridiana.aleman@gmail.com, darnes@cs.buap.mx, mariajsomodevilla@gmail.com
<http://www.cs.buap.mx/>

Resumen. En este artículo se presentan una serie de variantes para el algoritmo de Naïve Bayes aplicado a texto. Las variantes se enfocan en la clasificación de polaridad con el uso de un diccionario de palabras que se consideran como positivas o negativas. Entre las propuestas presentadas se utiliza una versión binaria y una ponderada de dicho diccionario. Las pruebas se realizan con dos corpus diferentes, obteniendo mejores resultados con el enfocado a opiniones de foros; además, el uso del diccionario aumenta la exactitud del clasificador y las métricas para la clase negativa.

Palabras clave: Aprendizaje automático, Naïve Bayes, polaridad, diccionario, ponderación.

Polarity Classification with a Dictionary and Bayes Algorithm

Abstract. In this article we present a set of variants for Naïve Bayes algorithm applied to text. The focus are the sentiment polarity using a polarized dictionary. The dictionary specifies if a word is a positive or negative class. Among the presented proposals we use a binary and a weighted version of the dictionary. The experiments are performed with two different corpus, getting better results with forums focused on opinions; also, dictionary increases the classifier accuracy and metrics for the negative class.

Keywords. Machine Learning, Naïve Bayes, polarity, dictionary, weighting.

1. Introducción

En la actualidad, las redes sociales son la plataforma para poder expresar opiniones, ideas y prácticamente cualquier tipo de información. La cantidad de texto que se genera es tan grande que es imposible analizarlos uno a uno; es por

esto que en la literatura se proponen varias herramientas con diferentes enfoques para poder realizar esta tarea de manera automática. Existen diferentes enfoques para el tratamiento de los textos en socialmedia, una de las más trabajadas es la detección de polaridad (positiva/negativa) de un mensaje, especialmente en la red social de Twitter.

Siguiendo con esta línea de investigación, en el presente trabajo se muestran algunas modificaciones al algoritmo de clasificación de Naïve Bayes para texto analizando dos corpus diferentes y basados en el uso de un diccionario binario.

La estructura del artículo es la siguiente: En la sección 2 se analizan algunos trabajos relacionados con el algoritmo Naïve Bayes y el análisis de polaridad de Twitter, posteriormente, en la sección 3 se explica la metodología a seguir para los experimentos, además, se explica de manera general el algoritmo utilizado. La sección 4 muestra las variantes propuestas para el algoritmo y en la sección 5 se analizan los resultados obtenidos por clasificador y corpus; finalmente, en la sección 6 se presentan las conclusiones obtenidas y el trabajo futuro.

2. Estado del arte

Entre las investigaciones más recientes relacionadas con el análisis de polaridad o con el tratamiento del algoritmo bayesiano que se revisaron para esta investigación se encuentran las siguientes:

En [6] se propone un método para analizar sentimientos en artículos de índole política publicados en socialmedia, aunque también se especifica que es pertinente para otros temas. Como entrenamiento se utilizan artículos de periódicos de Indonesia y se procesan mediante 5 módulos: Lector para determinar la clase del artículo (análisis hecho por humanos), analizador para separar cada palabra de los signos de puntuación, limpiador para eliminar palabras especiales y adverbios, análisis humano para encontrar el pronombre o nombres involucrados en la noticia, además de responder a las preguntas "¿Quién?", "¿A quién?" y "¿Qué?". Las respuestas a estas interrogantes constituirán las palabras clave a tomar en cuenta por el método propuesto. La clasificación se realiza mediante al algoritmo de Bayes, obteniendo 87% de exactitud. Entre los problemas que reportan los autores se encuentran la poca riqueza en vocabulario de los textos, la complejidad del lenguaje utilizado (Indonesio) y la dificultad al localizar con precisión el quién y el qué de cada artículo.

Los autores de [1] categorizan la polaridad en documentos multilingües mediante un nuevo enfoque basado en el meta-aprendizaje genérico. Para esta propuesta realizan la desambiguación del sentido de la palabra, así como algunas funciones basadas en la ampliación del vocabulario. Los autores se enfocan en la combinación de clasificadores y el estudio de los efectos de la desambiguación semántica en el vocabulario de los documentos y por ende, en los atributos de clasificación. Para la desambiguación y obtención de los posibles dominios de las palabras utilizan la herramienta *BabelNet*¹ la cual es un diccionario enciclopédico

¹ <http://babelnet.org/>

multilingüe. Entre los conjuntos de atributos que se analizan se encuentran: TF-IDF bolsa-de-palabras, TF-IDF con n-gramas y un recurso léxico basado en la minería de opiniones; además, se incluye un clasificador independiente para estudiar la contribución de la desambiguación y expansión de conceptos en las etiquetas POS empleadas (adjetivos, sustantivos, verbos y adverbios). Por último, bajo el supuesto de que los conceptos relacionados semánticamente tienen un pariente cercano común, se aprovecha esta posible relación entre los conceptos incluyendo un clasificador basado en la ampliación del vocabulario. Los resultados finales se extraen mediante máquinas de soporte vectorial y C4.5. Los mejores resultados se obtienen con máquinas de soporte vectorial utilizando los conjuntos de atributos que involucran desambiguación semántica.

En [7] se trabaja sobre los enfoques de ponderación de funciones para los clasificadores de Bayes enfocados a texto. La mayoría de los enfoques de ponderación analizados para clasificadores de texto Bayes, la mejora en los resultados implica más tiempo de procesamiento. Por lo tanto, los autores proponen dos enfoques nuevos: Uso de la ganancia de información para la ponderación del texto y ponderar en función de un árbol de decisión. Los resultados experimentales basados en datos de referencia y del mundo real muestran que los enfoques propuestos muestran mayor precisión y mantienen el tiempo de ejecución sin grandes cambios en los modelos finales.

[2] propone un sistema de clasificación automática para el idioma español basada en el análisis del contexto en donde fueron publicados llamado SCOPT. Para la metodología que proponen, se genera un diccionario ponderado con elementos en español con 6 emociones básicas (alegría, sorpresa, repulsión, miedo, enojo y tristeza). El método de clasificación de un tuit es determinado tomando en cuenta la ocurrencia de las palabras del diccionario afectivo así como su ponderación afectiva. La polaridad de un tuit se determina por medio de realizar la combinación lineal de los pesos asignados a cada una de las palabras que aparecen dentro de un tuit y que ocurren dentro del diccionario afectivo. Los resultados alcanzan hasta el 74% de precisión sobre algunos temas.

Finalmente, en [3] se presenta la metodología propuesta para la tarea 4 del SemEval2016, propone el uso de un recurso léxico para el preprocesamiento de datos y entrenar con un modelo de redes neuronales y una representación tipo Doc2Vec. En la metodología se incluyen diccionarios de palabras de argot, contracciones, abreviaturas y emoticones más populares en los medios sociales. Para el proceso de clasificación, se utilizan las características obtenidas sin supervisión en un clasificador SVM obteniendo resultados superiores al *baseline* de la competencia.

3. Metodología

Para la realización de los experimentos se utilizaron dos conjuntos con documentos de diferentes longitudes. El número de documentos del conjunto de entrenamiento y de prueba por corpus se muestran en la Tabla 1.

Tabla 1. Conjuntos utilizados para el análisis.

Corpus	Entrenamiento	Evaluación
<i>Opinion</i>	3,608	1,548
<i>Textos</i>	10,000	1,762

El conjunto denominado *Opinion* contiene principalmente textos cortos obtenidos de foros en donde se emite alguna opinión (positiva o negativa) sobre algún producto o aplicación para celular. Como sus textos y el número de instancias es pequeño, el vocabulario es pobre respecto al corpus de *Textos*, el cual contiene más palabras por documento al ser textos que hablan sobre noticias o eventos recientes, pero siempre manteniendo la opinión positiva o negativa del autor.

Los procedimientos realizados son los siguientes (Figura 1).

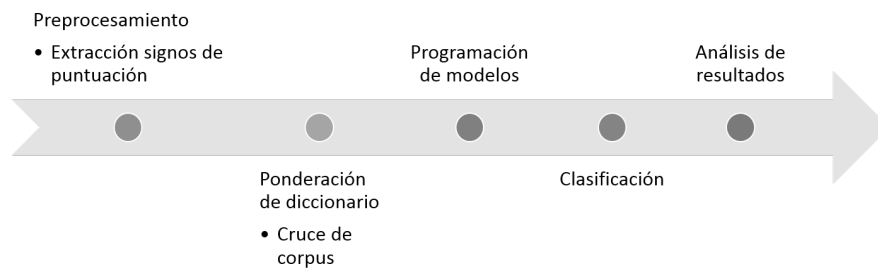


Fig. 1. Proceso de experimentación.

1. Extracción de signos de puntuación, caracteres no imprimibles de las palabras de todos los conjuntos de texto. En este caso no se hizo ningún otro tipo de eliminación ni análisis semántico de las palabras.
2. El diccionario con el que se cuenta contiene una lista de 3,841 palabras y la clase positiva o negativa para cada una de ellas (una palabra no puede ser positiva y negativa al mismo tiempo).
3. Para la programación de modelos, se proponen 4 variantes del clasificador inicial; para estos modelos los atributos son las palabras de los textos.
4. Al final se comparan los resultados obtenidos por corpus y por modelos. La evaluación se realiza con las matrices de confusión y al final un análisis de exactitud.

3.1. Clasificador Naïve Bayes

Todos los experimentos parten del clasificador de Naïve Bayes, el cual es uno de los modelos más simples para clasificación. Se basa en la suposición de que las cantidades de interés son gobernadas por distribuciones de probabilidad y que las

decisiones óptimas pueden estar hechas razonando acerca de estas probabilidades junto con la observación de datos [4]. Se basa en la aplicación del *Teorema de Bayes* para predecir la probabilidad condicional de que una instancia pertenezca a una clase $P(c_i|d_j)$ a partir de la probabilidad de las instancias dada la clase $P(d_j|c_i)$ y la probabilidad *a priori* de la clase en el conjunto de entrenamiento $P(c_i)$ (fórmula 1):

$$P(c_i|d_j) = \frac{P(c_i)P(d_j|c_i)}{P(d_j)}. \quad (1)$$

En el procesamiento de lenguaje natural, esta fórmula se adecua al texto realizando cambios como el uso del logaritmo para evitar probabilidades demasiado pequeñas y la suma de uno para evitar probabilidades de cero (suavizado). Dados estos cambios, en [5] se propone la fórmula 2 :

$$NB(D) = arg_c max log(P(c)) + \sum_{i=1}^k log(P(f_i|C)). \quad (2)$$

Todas las variantes presentadas en este artículo se basan en la extracción de $P(f_i|C)$, las cuales se explican en la siguiente sección.

4. Variantes propuestas

Aparte del clasificador usual, se presentan 4 modificaciones más. Las maneras de calcular $P(f_i|C)$ se muestran en la figura 2 y se explican a continuación:

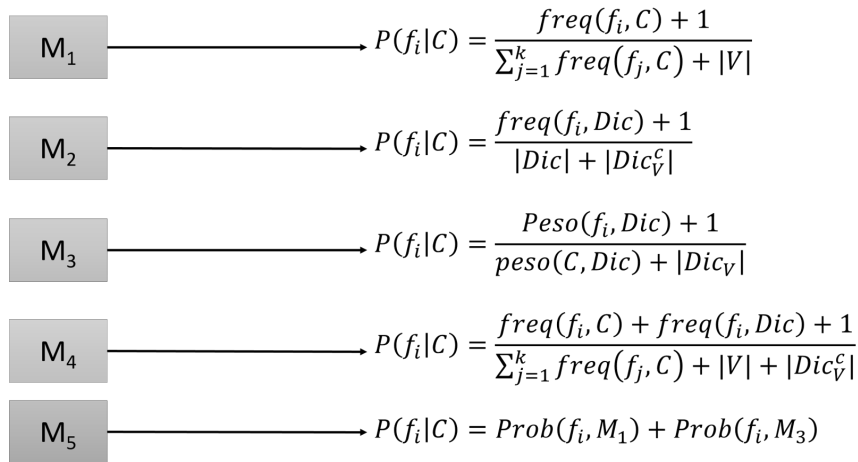


Fig. 2. Variantes propuestas del clasificador.

- M_1 : Fórmula general para la frecuencia de cada palabra por clase utilizando el suavizado. Cada palabra del conjunto de entrenamiento se convierte en atributo; este modelo fue propuesto en [5].
- M_2 : Se siguen utilizando palabras como atributos, con la variante de que sólo se utilizarán aquellas que aparezcan en el diccionario. Si una palabra del conjunto de entrenamiento se encuentra en el diccionario y esta tiene la clase del documento analizado, se contabiliza en las frecuencias. Para el caso de las frecuencias 0, se vuelve a utilizar el suavizado y la suma del total de palabras en el denominador de la función. En la fórmula:
 - $frec(f_i, Dic)$: Frecuencia de la palabra en el diccionario, en este caso al ser binario, la frecuencia será 0 si no pertenece a la clase del documento y 1 si pertenece,
 - $|Dic|$: Total de palabras del diccionario que se encuentran en el conjunto de entrenamiento.
 - $|Dic_V^c|$: Total de palabras del diccionario que se encuentran en la clase analizada en el conjunto de entrenamiento.
- M_3 : Al igual que en M_2 , sólo se utilizarán las palabras que aparezcan en el diccionario de acuerdo a la clase analizada, pero en este caso no se trata de un diccionario binario, sino que cada palabra tiene una frecuencia asignada. Esta frecuencia se obtiene mediante el *cruce de corpus*, es decir, se utiliza el corpus *Opinion* para extraer la probabilidad de cada una de las palabras del diccionario (probabilidad simple). Una vez extraída la ponderación de las palabras, ésta se utiliza para calcular el modelo del corpus *Textos*.
- M_4 : Se combina el diccionario binario con el clasificador usual, es decir, se toma la frecuencia normal de las palabras, pero si la palabra aparece en el diccionario con la misma clase del documento analizado, se le añade uno al cálculo.
- M_5 : Se combinan el diccionario ponderado con el clasificador usual (M_1 y M_3).

5. Resultados obtenidos

En la Figura 3 se muestran los resultados de la implementación del algoritmo normal con suavizado. Aunque el número de documentos bien clasificados es muy similar, por el número total de instancias los mejores resultados son para el corpus *Opinion* al tener solo un poco más de 100 instancias clasificadas incorrectamente. En ambos corpus se muestran más errores de tipo II (falsos negativos) respecto a la clase positiva.

En la Figura 4 se muestran los resultados para la clasificación utilizando únicamente el diccionario (tanto en su versión binaria como la propuesta de ponderación). Si bien el número de documentos clasificados correctamente baja considerablemente, hay que tomar en cuenta que aún se siguen clasificando bien, más del 50% de las instancias con menos atributos.

Se continúa manteniendo el comportamiento de mayor número de documentos positivos clasificados incorrectamente, además de que en el corpus *Textos*, la

	Opinion		Textos	
	Positivo	Negativo	Positivo	Negativo
Positivo	680	94	610	271
Negativo	34	741	148	734

Fig. 3. Matrices de confusión por corpus para M_1 .

	Opinion		Textos	
	Positivo	Negativo	Positivo	Negativo
M_2				
Positivo	512	262	660	221
Negativo	151	624	540	342
M_3				
Positivo	518	256	504	377
Negativo	147	628	404	478

Fig. 4. Matrices de confusión por corpus para M_2 y M_3 .

clase negativa baja considerablemente el número de documentos bien clasificados en las dos variantes del diccionario. De manera general, el diccionario binario es un poco más eficiente que el ponderado, aunque esto se puede ver seriamente afectado por la calidad del corpus que se utilice para ponderarlo.

En la Figura 5 se muestran los resultados de unir la frecuencia simple de las palabras con el diccionario binario (M_4) y el ponderado (M_5).

	Opinion		Textos	
	Positivo	Negativo	Positivo	Negativo
M_4				
Positivo	686	88	609	272
Negativo	34	741	148	734
M_5				
Positivo	688	86	610	271
Negativo	34	741	149	733

Fig. 5. Matrices de confusión por corpus para M_4 y M_5 .

Los resultados siguen manteniendo comportamientos similares, pero se observa un ligero incremento en los valores, lo cual hace que se superen los datos inicialmente obtenidos con M_1 . A diferencia del uso de los diccionarios solamente, en el corpus *Textos* mejoró considerablemente la clasificación de los textos negativos; además, se incrementaron los valores de la diagonal principal. Para visualizar de manera general las modificaciones propuestas, en la Figura 6 se muestra la exactitud obtenida por modelo propuesto y corpus.

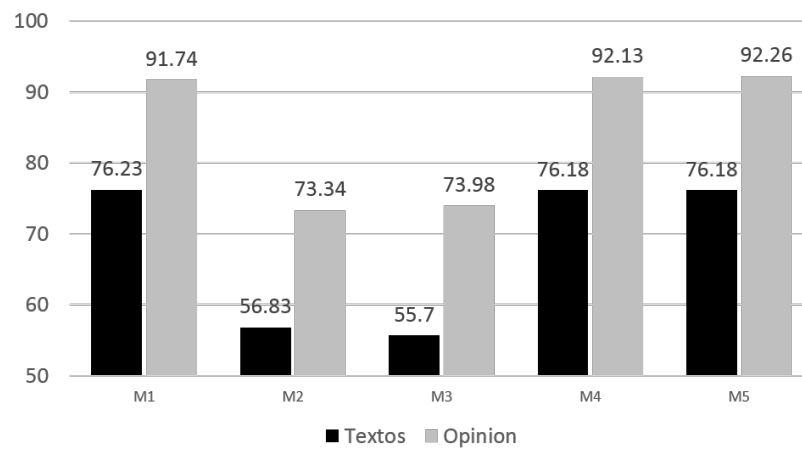


Fig. 6. Variantes propuestas del clasificador.

Si bien el algoritmo tradicional supera por ejemplo a las dos implementaciones del diccionario, en el caso de la fusión de los modelos hay un incremento que aunque muy pequeño, se vuelve importante al manejar grandes cantidades de documentos. Para el caso del corpus *Textos* no se llegó a una mejora, pero con el corpus *Opinion* se logró mejorar la calidad del clasificador.

6. Conclusiones y trabajo futuro

En el presente artículo se presentaron varias propuestas para incluir un diccionario en el clasificador de Naïve Bayes. Si bien sólo en un corpus se obtuvo un ligero incremento se pueden llegar a las siguientes conclusiones:

- Tanto en los modelos de M_4 y M_5 se encuentran mejoras en la clasificación de los documentos positivos, especialmente en el corpus *Opinion*.
- El diccionario ponderado mejora poco o nada la exactitud del clasificador, por lo que se debe analizar otras formas de ponderación que incrementen esta métrica.

- Aunque en este caso se realizaron experimentos para clases de tipo positivo-negativo, la metodología es aplicable a cualquier clasificación binaria, siempre y cuando se cuente con la lista de palabras representativas de cada clase.
- Este trabajo se puede considerar como *baseline* para la experimentación futura con otros corpus e incluso verificar la fiabilidad del diccionario por medio de técnicas estadísticas.

Referencias

1. Franco-Salvador, M., Cruz, F.L., Troyano, J.A., Rosso, P.: Cross-domain polarity classification using a knowledge-enhanced meta-classifier. *Knowledge-Based Systems* 86, 46 – 56 (2015), <http://www.sciencedirect.com/science/article/pii/S0950705115002063>
2. Gálvez-Pérez, J.R., Gómez-Torrero, B.E., Ramírez-Chávez, R.I., Sánchez-Sandoval, K.M., Castellanos-Cerda, V., García-Madrid, R., Jiménez-Salazar, H., Villatoro-Tello, E.: Sistema automático para la clasificación de la opinión pública generada en twitter. *Research in Computing Science* 95, 23–36 (2015), http://rsc.cic.ipn.mx/2015_95/Sistema%20automatico%20para%20la%20clasificacion%20de%20la%20opinion%20publica%20generada%20en%20Twitter.pdf
3. Gómez-Adorno, H., Vilariño, D., Sidorov, G., Avendaño, D.P.: Cicbuapnlp at semeval-2016 task 4-a: Discovering twitter polarity using enhanced embeddings. In: *Proceedings of the 10th International Workshop on Semantic Evaluation, SemEval@NAACL-HLT 2016*, San Diego, CA, USA, June 16-17, 2016. pp. 145–148 (2016), <http://aclweb.org/anthology/S/S16/S16-1021.pdf>
4. John, G.H., Langley, P.: Estimating continuous distributions in bayesian classifiers. In: *Proceedings of the Eleventh conference on Uncertainty in artificial intelligence*. pp. 338–345. UAI'95, Morgan Kaufmann Publishers Inc., San Francisco, CA, USA (1995)
5. Manning, C.D., Schütze, H.: *Foundations of statistical natural language processing*. MIT Press, Cambridge, MA, USA (1999)
6. Soelistio, Y.E., Surendra, M.R.S.: Simple text mining for sentiment analysis of political figure using naive bayes classifier method. *CoRR* abs/1508.05163 (2015), <http://dblp.uni-trier.de/db/journals/corr/corr1508.html#SoelistioS15>
7. Zhang, L., Jiang, L., Li, C., Kong, G.: Two feature weighting approaches for naive bayes text classifiers. *Knowledge-Based Systems* 100, 137 – 144 (2016), <http://www.sciencedirect.com/science/article/pii/S0950705116001039>

Desarrollo de un sistema aumentativo y adaptativo para apoyar el proceso de lectoescritura para alumnos con discapacidad visual de educación superior

Carmen Cerón, Etelvina Archundia, Beatriz Beltrán, Patricia Cervantes, José Galindo

Benemérita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

{mceron, etelvina, bbeltran}@cs.buap.mx,
{cervantes.patty, jgalindo2016}@gmail.com

Resumen. El propósito de este artículo es presentar el diseño y desarrollo de un Sistema Adaptativo como herramienta de apoyo para el desarrollo de las Competencias de comunicación oral y escrita en los estudiantes con discapacidad visual. La metodología utilizada fue de Prototipos y el Diseño Centrado en el Usuario, se usaron Objetos de Aprendizaje e Interfaces Naturales de Usuario como el Kinect de Microsoft. Para el desarrollo se usaron las tecnologías de OpenNI de Kinect, MySQL, Greenhouse SDK, C++, C# y se integró un conversor texto-voz (TTS) en Java. Finalmente se presentan los resultados obtenidos de una prueba piloto con un grupo focal de estudiantes con discapacidad visual bajo distintos escenarios de uso y de aprendizaje.

Palabras clave: Sistema accesible, aprendizaje adaptativo, objetos de aprendizaje, interfaces naturales de usuario, competencias.

Development of an Augmentative and Adaptive System to Support the Process of Reading and Writing for Visually Disabled Students in Higher Education

Abstract. The purpose of this article is to present the design and development of an Adaptive System like a help tool for the development the communicative oral and written Competence in the visually disabled students. The methodology used was the prototypes and the design centered in the user, it was used the learning objects and interfaces natural of the users like the Kinect of Microsoft. For the development was used the technologies of OpenNI of Kinect, MySQL, Greenhouse SDK, C++, C# and was to integrate to converter text-voice (TTS) in Java. Finally, is present the results obtained in the pilot test with focal group technique of students with visual disability under different settings of use and learning.

Keywords. Accessible system, adaptive learning, learning objects, competency.

1. Introducción

Actualmente las Tecnologías de la Información y Comunicación se consideran como herramientas de apoyo para las personas con alguna discapacidad en sus actividades diarias ya que les permiten acceder a diferentes ámbitos como son: educación, trabajo, comunicación y entretenimiento [1]. En México, la población en edad escolar que presenta problemas visuales es de un 20% y en términos mundiales los errores de refracción no corregidos (miopía, hipermetropía, astigmatismo y presbicia) constituyen la causa más importante de la discapacidad visual. De acuerdo con el Instituto Nacional de Estadística y Geografía (INEGI) la segunda discapacidad en el país es la visual, la primera es la motriz, existiendo otras discapacidades del hablar o comunicarse, escuchar, atender o aprender [2]. Por lo cual se requiere atender en las instituciones educativas a los estudiantes con alguna discapacidad visual para brindar mejores condiciones y herramientas para su vida académica universitaria. Por otra parte, lograr una mayor inclusión del uso de las Tecnologías en sus actividades para reducir la brecha digital y propiciar el desarrollo de distintas competencias genéricas y disciplinares. Una de las necesidades que se han identificado en la Educación Superior para una mayor inclusión de esta población es la de acondicionar espacios, aulas, desarrollar materiales didácticos, ayudas de software y recursos que puedan utilizarse en la modalidad presencial de forma accesible y permitan enriquecer el desarrollo de las competencias en el alumno en el proceso de enseñanza-aprendizaje.

Esta investigación tiene como propósito *el diseño y desarrollo de un Sistema Aumentativo y Adaptativo Interactivo para el Proceso de Lectura y Escritura (SIS-ADAP-LE) como herramienta de apoyo para el desarrollo de las Competencias de Comunicación oral y escrita en los estudiantes con discapacidad visual por medio de Objetos de Aprendizaje Adaptativos e Interfaces Naturales de Usuario*, integrando el uso del dispositivo del Kinect para facilitar la estimulación física y cognitiva por el reconocimiento postural del cuerpo humano, el uso gestual y reconocimiento de voz, adaptándose a las necesidades y limitaciones de los usuarios para lograr el acceso al sistema.

Su aportación permitirá apoyar el proceso aprendizaje y desarrollo de habilidades en la comunicación escrita y oral en los estudiantes universitarios con discapacidad visual para poder fortalecer sus competencias lectoras mediante actividades interactivas.

Para el diseño y desarrollo del Sistema Adaptativo se utilizó la metodología de Prototipos y el Diseño Centrado en el Usuario, el modelado en UML, las tecnologías de MySQL, Greenhouse SDK y se integró un conversor texto-voz (TTS) utilizando las herramientas de Java, C++ y C#.

El artículo está organizado como se indica a continuación:

En la sección 2 se presenta la fundamentación y la revisión teórica del trabajo de las Tecnologías de Apoyo para la Discapacidad Visual, Sistemas Adaptativos Hipermedia, la Tecnología de Interfaces Naturales de Usuario y del conversor de texto-voz.

En la sección 3 se define el análisis y el diseño del sistema en UML, la arquitectura del sistema, la implementación del sistema con las tecnologías OpenNI de Kinect, MySQL, Greenhouse SDK, C++ y C# y Java para el reconocimiento de gestos y

voz. En la sección 4 se muestran los resultados de la prueba piloto y las pruebas de funcionalidad y de valoración por los estudiantes con discapacidad del grupo focal desde un enfoque metodológico cualitativo. Finalmente se presentan las conclusiones y el trabajo a futuro de esta investigación.

2. Estado del arte

En esta sección se revisan los tópicos de Discapacidad Visual, Tecnologías de Apoyo, Sistemas Hipermedia Adaptativos, y las Interfaces Naturales de Usuario, los cuales sustentan a esta investigación.

2.1. Discapacidad visual

En el 2001, de acuerdo con la Clasificación Internacional del Funcionamiento, de la Discapacidad y de la Salud, presentada por la Organización Mundial de la Salud las personas con discapacidad “son aquellas que tienen una o más deficiencias físicas, mentales, intelectuales o sensoriales y que al interactuar con distintos ambientes del entorno social pueden impedir su participación plena y efectiva en igualdad de condiciones a las demás” [3]. La Discapacidad visual (DV) está relacionada con una deficiencia del sistema de la visión que afecta la agudeza visual, campo visual, motilidad ocular, visión de los colores o profundidad, afectando la capacidad de una persona para ver. Al hablar de DV podemos referirnos a la persona que presenta ceguera o baja visión. Algunos autores han definido términos como “dificultad visual severa”, “deficiencia visual grave”, “visión subnormal”, “visión parcial”, “visión residual”, entre otros, para definir el espacio intermedio entre la visión normal y la ausencia total o casi total de visión, caracterizado por un sistema visual con alteraciones irreversibles y una pérdida en la capacidad visual que constituye un obstáculo para el desarrollo de la vida de las personas [3]. Por lo cual tenemos dos clasificaciones:

- *Baja visión (hasta agudeza visual de 6/18)*. Es una condición de vida que disminuye la agudeza o el campo visual de la persona; es decir, que quienes presentan una baja visión ven significativamente menos que aquellas que tienen una visión normal. Siendo esta una condición que no limita a quien la padece en su capacidad para desplazarse y conducirse de la forma que lo hace una persona con una visión optima, impidiendo que las personas que le rodean comprendan las dificultades que esta condición representa para realizar todas aquellas actividades que exigen una agudeza visual mayor, considerándolos apáticos, lentos, descuidados, incómodos o latosos. El docente debe estar abierto a la posibilidad de que dentro de su grupo haya estudiantes con Baja Visión, por lo que es recomendable que identifique a este tipo de personas para ofrecer apoyos.
- *Ceguera (agudeza visual menor a 20/200)*. Es una condición de vida que afecta la percepción de imágenes en forma total reduciéndose en ocasiones a una mínima percepción de luz, impidiendo que la persona ciega reciba información visual del mundo que le rodea.

2.2. Tecnologías de apoyo y sistemas aumentativos y alternativos

La Tecnología y de las Ciencias Computacionales en la última década se han inclinado por apoyar el área de las ciencias de la salud y la educación. La tendencia es aportar aplicaciones para apoyar a las personas con discapacidad y lograr mayor inclusión en su entorno. A partir de lo anterior ha surgido el término de Tecnología de Apoyo también conocida como Tecnología de Acceso o Rehabilitación.

Para Alcantud la define como “Todos aquellos aparatos, utensilios, herramientas, programas de ordenador o servicios de apoyo que tienen como objetivo incrementar las capacidades de las personas que, por cualquier circunstancia, no alcanzan los niveles medios de ejecución que por su edad y sexo le corresponderían en relación con la población normal” [4]. Por otra parte, la Norma UNE 139802:2003 le llamo “*ayudas técnicas de software*”. Sin embargo, en 2007, dentro de ISO 9999:2007, estableció que ya no se tienen que llamar ayudas técnicas, sino “Tecnologías de Apoyo” [5], las cuales incluyeron: lectores de pantalla, magnificadores de pantalla, navegadores accesibles y navegadores alternativos.

Por otra parte de acuerdo a la ayuda que proporcionan y funcionalidad son conocidos como *Sistemas Alternativos y Aumentativos de Acceso a la Información del Entorno*: Son sistemas que engloban las ayudas o apoyos para personas con discapacidad visual y/o auditiva, *modifican la señal, aumentándola o cambiando su modalidad para poder ser percibido por ellos*.

Cuando existe un problema sensorial, la opción es el aumento de la señal que el contexto envía al sujeto. Los sistemas aumentativos se dirigen hacia la población con déficits visuales y auditivos mientras que los sistemas alternativos son medios que permiten, a quienes presentan la imposibilidad de alcanzar la información mediante una determinada modalidad sensorial, cambiar la naturaleza de la misma de modo que pueda aprehenderse mediante una modalidad que la persona mantiene funcional. *Algunas de estas son los medios tecnológicos para escritura en Braille, Tecnologías del Habla: El reconocimiento de voz y la conversión texto-voz*. Actualmente se están incorporando otras tecnologías inalámbricas para la manipulación y el control del entorno como son las tabletas, iPad y las consolas (Wii, Play Station, Kinect) las cuales promueven una mayor interactividad y autonomía a las personas con discapacidad, las cuales están siendo aplicadas en la rehabilitación y para apoyar procesos de aprendizaje en la educación.

2.3. Sistemas hipermedias adaptativos (SHA) y aprendizaje adaptativo

El Aprendizaje Adaptativo recolecta información sobre los hábitos de aprendizaje, conocimientos, debilidades y fortalezas de cada usuario para crear un camino de aprendizaje a la medida de los alumnos teniendo en consideración sus dificultades y se adapta a la forma y ritmo de aprendizaje de cada uno, logrando tener un aprendizaje diferenciado y adaptativo único para cada estudiante.

Para Gaudioso define los SHA como “*aquellos sistemas de hipermedia capaces de ajustar su presentación y navegación a las diferencias de los usuarios, reduciendo así los problemas de desorientación y falta de comprensión, propios de los sistemas hi-*

permedia no adaptativos” [6]. Para Brusilovsky “Un sistema se considera adaptativo cuando se adapta de forma automática y de forma personalizada a las necesidades del usuario [7]. Por lo que el SHA permite personalizar la información almacenada y la presenta a los usuarios según sus preferencias, conocimientos e intereses [8].

El SHA aplicado a la educación permite guiar a los estudiantes con discapacidad adaptando los contenidos y la navegación de acuerdo a las características y necesidades del usuario logrando incluir interfaces naturales de usuario para permitir una interacción adaptativa de acuerdo al dispositivo con el cual puede utilizar más el usuario teniendo en consideración su discapacidad, como en la Figura 1 se muestra.

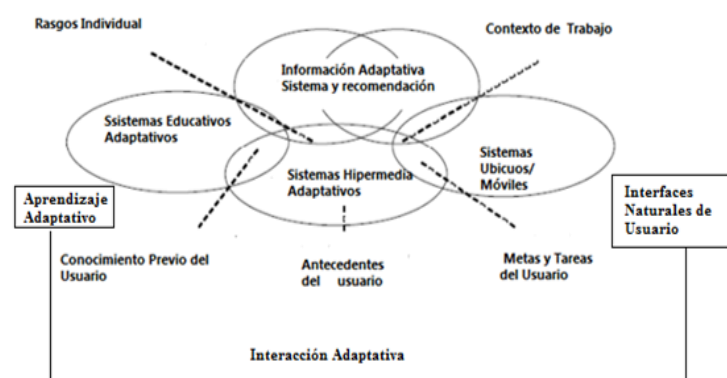


Fig. 1. Adaptación del sistema hipermedia de Brusilovsky con interfaces naturales.

2.4. Interfaces naturales de usuario

Las interfaces naturales de usuario (NUI) son aquellas en las que se interactúa con un sistema sin utilizar sistemas de mando o dispositivos de entrada de las GUI (Interfaces de Usuario Gráficas) como sería el ratón, teclado alfanumérico, joystick, etc., y en su lugar, se hace uso de movimientos gestuales tales como las manos o el cuerpo que funcionan como el mismo mando de control, algunos ejemplos son el reconocimiento a través de sensores de voz, de mirada, de movimiento y multitáctiles [9].

Existen varios dispositivos que son utilizados para implementar las NUI como son:

1. Microsoft Kinect Sensor, es un dispositivo que integra una gran cantidad de sensores que detectan el movimiento corporal del usuario que se coloca enfrente de él, al mismo tiempo captura el sonido permitiendo la interacción utilizando instrucciones de voz [10].
2. Leap Motion Controller: es una pequeña barra que se conecta a la computadora y detecta el movimiento de las manos para interactuar con la máquina [11].
3. Myo Gesture Control Armband: es una banda que se coloca en el antebrazo y por medio de los movimientos en el brazo realiza la interacción con la computadora [12].

4. Google Glass: son unos lentes creados por Google por medio de los cuales los usuarios solicitan y reciben información o servicios utilizando su voz, localización, video, entre otras variables que los lentes perciben [13].

Por lo que en esta investigación se enfoca al dispositivo del Kinect como medio de reconocimiento de gestos tanto de movimiento como de comando de voz. Según Wigdor y Wixon [14] consideran que las interfaces naturales de usuario se basan en crear interfaces que generen experiencias para poder usar la tecnología de acuerdo a las necesidades y contexto del usuario. Así mismo recomiendan ciertas pautas para el diseño de interfaces naturales tales como:

- Crear experiencias que den la sensación que es parte “de su cuerpo” como una extensión y lograr esa autonomía.
- Crear experiencias para todos usuarios desde nivel de principiante o experto de acuerdo a las características y perfil del usuario.
- Crear interfaces que considere contextos, metáforas, indicaciones visuales y propiciar escenarios con experiencias reales de su contexto.

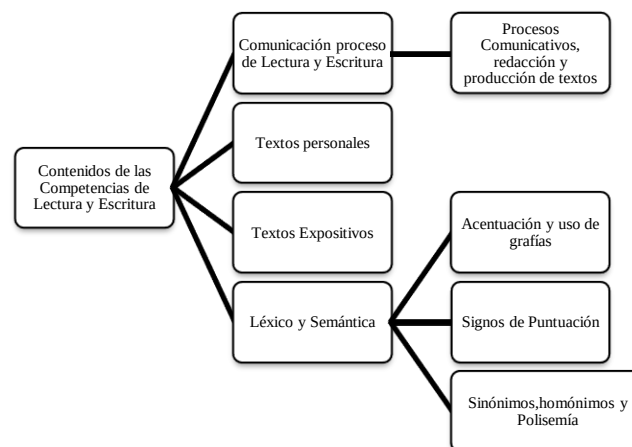


Fig. 2. Contenidos para el proceso de lectura y escritura del sistema.

2.5. Competencias del proceso de lectura y escritura

Para Tobón [15] define las competencias “como procesos complejos de desempeño integral con idoneidad en determinados contextos, que implican la articulación de diversos saberes, para realizar actividades y/o resolver problemas con sentido de reto, motivación, flexibilidad, creatividad y comprensión, dentro de una perspectiva de mejoramiento continuo y compromiso ético”. Para este sistema las competencias de lectura y escritura a considerar y los contenidos a abordar en el sistema se muestran en la Figura 2.

3. Análisis y diseño del sistema aumentativo y adaptativo para apoyar el proceso de lectura y escritura (SIS-ADAP-LE)

3.1. Diseño centrado en el usuario

Para el diseño del prototipo se utilizó el Diseño Centrado en el Usuario (DCU) que sitúa al estudiante con discapacidad visual en el núcleo del proceso de diseño de la interfaz. Para lo cual se realizó una serie de entrevistas y encuestas con el grupo focal, finalmente se plantearán los prototipos de la aplicación con balsamiq. Para el análisis y el diseño del sistema se determinaron los Casos de Uso del sistema en UML (ver Figura 3). El sistema permite identificar dos usuarios para el acceso y manipulación del sistema:

- Usuario Docente: Puede realizar consultas generales, alta de materiales y modificación con respecto a los contenidos, actividades y evaluación del tema.
- Usuario Estudiante con baja visión/Ceguera: Consulta los contenidos y materiales de información, al realizar el diagnóstico, el sistema adapta las interfaces y presenta los contenidos que requiere aprender y propone una serie de actividades lúdicas para el desarrollo de las competencias del proceso de lectura y escritura.

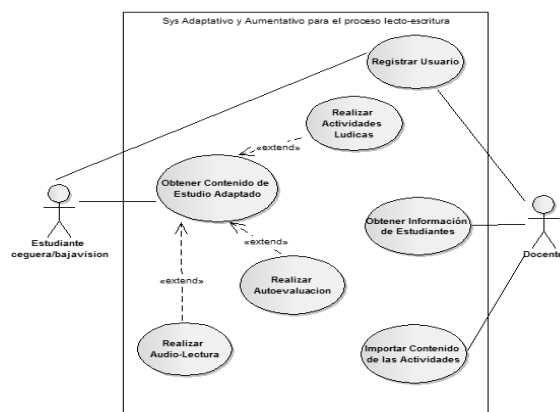


Fig. 3. Diagrama de caso de uso del sistema.

3.2. Arquitectura del sistema

El Sistema compuesto por los Modelos de Dominio, Adaptación y Usuario requiere obtener el perfil del usuario para lograr la interacción de los distintos dispositivos de Interfaces Naturales activos de acuerdo a la tecnología de apoyo para la discapacidad determinada (Baja visión/ceguera) y acceder a los contenidos mediante los Objetos de Aprendizaje, los cuales son considerados como recursos de aprendizaje compuesto por los cuatro elementos: Competencia, contenido de información, actividad de aprendizaje y la evaluación. Para el registro de la información y seguimiento académ-

mico se utilizó una base de datos en MySQL como se muestra la estructura general del sistema en la Figura 4.

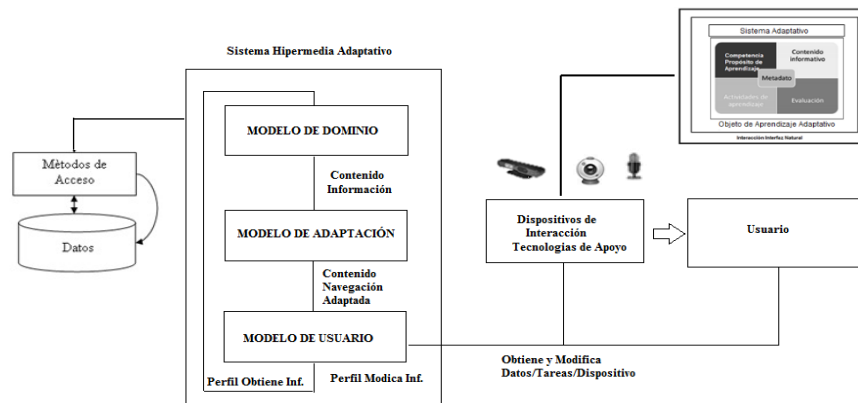


Fig. 4. Estructura general del sistema.

El sistema considera tres niveles de desempeño de las competencias son: 1) Iniciada, 2) En Proceso y 3) Desarrollada. De acuerdo al perfil pueden ser: Principiante, Intermedio y Avanzado. Para evaluar los contenidos se maneja en tres puntajes: 1) Bajo, menor a la media, 2) Regular, dentro de la media y 3) Alto, superior a la media, lo cual permite llevar el seguimiento académico de su proceso de aprendizaje.

4. Desarrollo y pruebas del sistema

Para desarrollo del sistema se usó Greenhouse SDK que permite la interacción de varios usuarios a través de la entrada gestual, voz y del dispositivo, manejo de múltiples pantallas, y en red, versión para Mac OS X. El sistema utiliza el reconocimiento de gestos y de voz, un gesto es el movimiento o posicionamiento resultante de la combinación de varios puntos del cuerpo como por ejemplo el movimiento de las manos o la producción de señas con los dedos que permite la manipulación de los objetos del sistema. Los gestos que se reconocen para el sistema son: A) Victory Pose (Posición victoria), B) ClosedFistPose (Posición puño cerrado), C) NumberonePose (Posición numero uno), D) OpenPalmFingersSpreadPose (Posición de Mano abierta y dedos separados), E) LShapePose (Posición figura de L) F) OpenPalmFingersTogetherPose (Posición de Mano abierta dedos juntos). Tal como se muestra en la figura 5.

Así también se proporciona una clasificación de los eventos del usuario en tres formas principales: dejar escapar eventos, apuntar eventos y desplazamiento de eventos GreenHouse proporciona un lenguaje en común que puede ser utilizado para definir métodos o acciones que generan los diferentes dispositivos de entrada, como son el mouse, teclado, dispositivos móviles, y gestos captados por el Kinect. La conexión entre dispositivo y el sistema para su desarrollo se realizó mediante SDK de Open NI (NITE) y/o por medio del SDK de Greenhouse, los cuales permiten la creación de una

interfaz de reconocimiento gestual y de comandos de voz. Esto permitió la construcción de archivos y/o scripts en C++ y C# y Librería Speech y Recognition en Java correspondientes para su implementación como se muestra en la Figura 6.

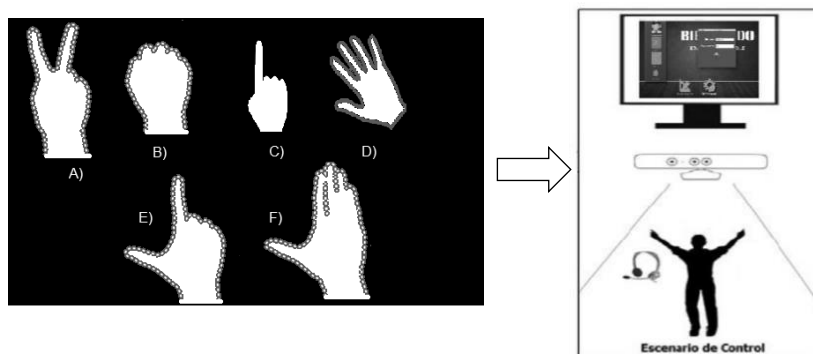


Fig. 5. Reconocimiento de gestos del sistema.

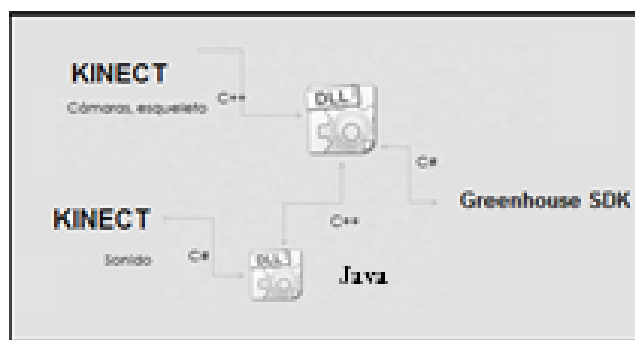


Fig. 6. Herramientas para el reconocimiento de gestos y voz del sistema.

Las personas con discapacidad visual utilizan otros sentidos para la percepción como el tacto y el oído, por lo cual el sistema permite manejar y estimular estos dos sentidos e incluso el visual dependiendo del tipo de discapacidad visual del usuario. Por lo cual el sistema se diseñó para lograr una interacción basada en gestos (movimientos o posicionamiento) y voz logrando una mayor interacción con el sistema y proporcionando una estimulación física y cognitiva para el apoyo del aprendizaje de los procesos de lectura y escritura. Como resultados de la aplicación se muestran algunas de las pantallas del sistema.

En el formulario el administrador docente podrá registrar sus datos para obtener acceso al sistema y posteriormente registrar los datos de su alumna/o para llevar un control de sus avances en las actividades y dar acceso al uso del sistema como se muestra en la Figura 7.



Fig. 7. Interfaz de registro administrador.

Un ejemplo de una actividad de nivel inicial después de recorrer el objeto de aprendizaje responde la evaluación de forma textual y auditivamente, cada respuesta correcta se registra y le permite ir avanzando en los contenidos de cada nivel para poder desarrollar las competencias de lectura y escritura, como se muestra en la Figura 8.



Fig. 8. Actividades de nivel 1 y nivel 2 respectivamente.

En las actividades de nivel 2 donde el usuario debe ir repitiendo acciones y palabras que se le indican, el objetivo es poder interactuar y para poder adquirir habilidad para ir acumulando puntos, ya que se irá reduciendo el tiempo de respuesta.

4.1. Prueba piloto

El prototipo del sistema fue piloteado en un grupo de 15 usuarios, los cuales poseen baja visión. Con respecto a las pruebas de funcionalidad se presentaron tres posibles escenarios posibles para los usuarios finales, a continuación se describen:

- Situación 1: Al usuario alumno se le dió una breve explicación del uso del sistema y se le acompañó en cada actividad.

- Situación 2: Al usuario alumno se le explicó el uso del Sistema y no se le acompañó en las actividades.
- Situación 3: Al usuario alumno no se le explicó el uso del Sistema y solo se le acompañó al inicio de una actividad.

Para cada una de las situaciones los usuarios deberán cumplir ciertas tareas para así comprobar el funcionamiento del uso de los Objetos de Aprendizaje Adaptativos y del sistema.

- Tarea 1: Poder ser detectado por el dispositivo del Kinect y registrar su perfil (Alumno con discapacidad /sin discapacidad).
- Tarea 2: Encontrar el menú de navegación y seleccionar una opción.
- Tarea 3: Utilizar los objetos de aprendizaje y recorrerlos mediante la interacción del dispositivo del Kinect y de comandos de voz.
- Tarea 4: Realizar las autoevaluaciones y reconocer su nivel de competencia.

Después de realizar las pruebas a 15 usuarios con 5 por cada situación, se muestran los resultados obtenidos en la Tabla 1, donde se observa el porcentaje de usuarios que pudieron completar la tarea.

Tabla 1. Resultados obtenidos de la prueba piloto.

Escenarios de Test			
Tareas	Situación No. 1	Situación No. 2	Situación No. 3
Tarea no.1	90%	85%	70%
Tarea no.2	100%	80%	85%
Tarea no.3	90%	85%	80%
Tarea no.4	90%	80%	90%
Promedio	92.5%	82.5%	81.2%

Para los usuarios con baja visión con una breve explicación de la Situación No. 1 el desempeño fue del 92.5% del cumplimiento de las tareas mientras que los usuarios de la Situación 2 al 82.5% y para la situación 3 lograron las tareas solo en un 81.2% esto implica que se presenta una interfaz intuitiva, sencilla y agradable para personas con discapacidad visual.

Finalmente se aplicó una encuesta de satisfacción “Valoración del Software [12], la cual evalúa siete criterios: Navegación, Interactividad, Inmersión, Usabilidad, Creatividad, Efectividad y Calidad, con una escala de 1 a 5, cuyo promedio obtenido fue de 4.5 (91.7%), a continuación, se presentan los resultados en la Figura 11.

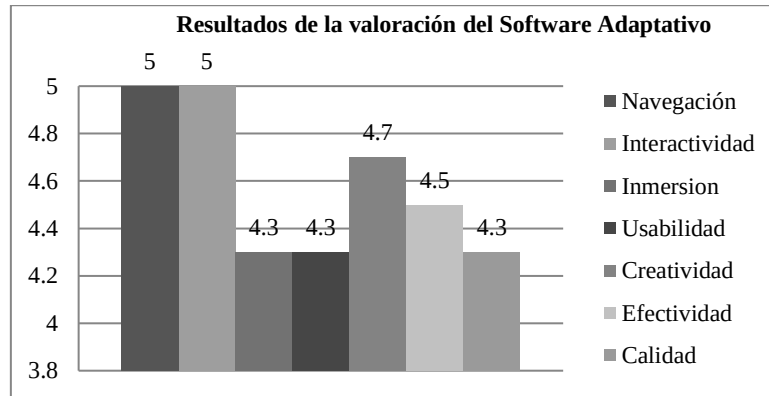


Fig. 9. Resultados de la encuesta de valoración del sistema.

5. Conclusiones y trabajo a futuro

Una de las principales contribuciones del Software es la adaptación a las necesidades del perfil del usuario con discapacidad y que mediante el uso interfaces naturales, el reconocimiento por medio de la voz de los contenidos de aprendizaje aportan una mayor interacción y ejecución, lo cual permite motivar a los estudiantes con discapacidad visual siendo una alternativa dinámica de aprendizaje para apoyar el proceso de adquisición de las competencias de lectura y escritura.

Como una perspectiva de este trabajo es realizar la evaluación de la accesibilidad y usabilidad del sistema propuesto ampliando la muestra de estudio y poder proveer a los alumnos con discapacidad materiales y herramientas digitales que puedan enriquecer el proceso de aprendizaje en distintas asignaturas, apoyar su vida académica y promover la inclusión del uso de las TIC para los estudiantes con alguna discapacidad para disminuir la brecha digital.

Referencias

1. Robles, M: Percepción visual y ceguera. Aspectos evolutivos y educativos de la deficiencia visual. Madrid, ONCE , <http://www.once.es/new> (1999)
2. Instituto Nacional de Estadística y Geografía (INEGI): Encuesta Nacional de Ingresos y Gastos de los Hogares, <http://www.inegi.org.mx/est/contenidos/proyectos/encuestas/hogares/> (2012)
3. Vázquez, J.: Clasificación Internacional del Funcionamiento, de la Discapacidad y de la Salud: CIF, Organización Mundial de la Salud. España <http://www.imsero.es/InterPresent2/groups/imsero/documents/binario/435cif.pdf> (2001)
4. Alcantud, F.: Estudiantes con Discapacidades Integrados en los Estudios Universitarios. Manual de Asesoramiento y Orientación Vocacional, Primera Edición, Síntesis, Madrid (1995)

5. AENOR (Asociación Española de Normalización y Certificación), productos de apoyo para personas en situación de discapacidad. Clasificación y terminología (UNE-EN ISO9999), <http://www.aenor.es/aenor/normas> (2017)
6. Gaudioso, Vázquez, E.: Contribuciones al Modelado del Usuario en Entornos Adaptativos de Aprendizaje y Colaboración a través de Internet mediante técnicas de aprendizaje Automático. Tesis doctoral, Facultad de Ciencias, Universidad Nacional de Educación a Distancia, Madrid, España (2002)
7. Brusilovsky, P.: Methods and techniques of adaptive hypermedia. *User Modeling and User Adapted Interaction*, (Special issue on adaptative hypertext and hypermedia), 6, 87–129 (1996)
8. Brusilovsky, P., Millan, P.: User Models for Adaptive Hypermedia y Adaptive Educational Systems. *The Adaptive, Lecture Notes in Computer Science*, 4321, 3–53 (2007)
9. Valli, A.: Notes on natural interaction. <http://www.idemployee.id.tue.nl/g.w.m.rauterberg/Movies/NotesOnNaturalInteraction.pdf> (2015)
10. Microsoft Kinect Sensor, <https://www.microsoft.com/en-us/kinectforwindows> (2015)
11. Leap Motion Controller. <https://www.leapmotion.com/> (2015)
12. Myo Gesture Control Armband. <https://www.thalmic.com/myo> (2015)
13. Google Glass. <https://developers.google.com/glass>
14. Wigdor, D., Wixon, D.: Brave NUI World, Designing Natural User Interfaces for Touch and Gesture. <http://dl.acm.org/citation.cfm?id=1995309> (2011)
15. Tobón, S., Pimienta, J., García, J.: Secuencias didácticas: Aprendizaje y Evaluación de Competencias. México, Pearson- Prentice Hall (2010)

Prototipo de una aplicación móvil accesible para apoyar a la localización de adultos mayores mediante un código compartido entre usuarios

Etelvina Archundia, Carmen Cerón, Beatriz Beltrán, Patricia Cervantes,
Fernando Rodríguez

Benemérita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

{etelvina,mceron,bbeltran}@cs.buap.mx,
{cervantes.patty,ferrgzt}@gmail.com

Resumen. El propósito de este artículo es presentar el diseño e implementación de una *aplicación móvil* para distintos dispositivos con sistema Android que permita *apoyar* la localización de adultos mayores mediante el uso de un código, que hace referencia a la longitud y latitud del usuario, estos códigos estarán almacenados en una base de datos que se consultaran mediante un Web Service. La aplicación puede compartir el código entre distintos usuarios y dispositivos para que puedan ver la localización correspondiente del adulto mayor brindando mayor protección y seguridad en el uso de la información y de la tecnología. El desarrollo de la aplicación móvil se basó en la metodología *Extreme Programación, el Diseño de Centrado en el Usuario, y las tecnologías integradas Web Service, PHP, MYSQL, Java y Apache*. Finalmente se presenta los resultados obtenidos de las pruebas de usabilidad con una muestra de usuarios adultos mayores bajo distintos escenarios de prueba.

Palabras clave: Servicios web, apps, usabilidad, tercera edad.

Prototype of Application Mobile Accessible for Bringing Support to the Localization of Elderly People through a Code that is Shared between the Users

Abstract. The purpose of this article is to introduce the design and implementation of a mobile app in gadgets that work using Android system to support the location of elderly people who have any disability using a code and share it with other apps, this code will show the user length and altitude and these data will be stored in a database that will be consulted in a Web Service. The app will be able to share the code between many users and apps and therefore the user will be able to watch the location of an elderly person providing more safety and the use of technology. The developing of this mobile app was based in

Extreme Programming methodology, Design Focused in the User and Integrated Technologies in a Web Service, PHP, MYSQL, Java and Apache. Finally, it is introduced the results that were obtained by usability testing with a sample of elderly users under different test scenarios.

Keywords. Web service, apps, usability, elderly people.

1. Introducción

Las aplicaciones tecnológicas para los adultos mayores en ocasiones representan dificultades para poder utilizarlas de forma sencilla por la falta de usabilidad de las interfaces y la carga cognitiva que requieren, lo cual ha provocado una creciente brecha digital generacional entre la tecnología y los adultos mayores, lo cual ha identificado diferentes grupos de generaciones en función de su capacidad de comunicación a través de las herramientas tecnológicas [1].

Las tecnologías deben brindar mejores alternativas de comunicación y apoyar las actividades cotidianas de los adultos mayores, para aportar una mayor calidad de vida en esta población vulnerable, la cual demanda apoyo, protección y seguridad. Es por eso que dentro de los servicios a los grupos vulnerables por parte de la Institución se lleva a cabo un proyecto de atender a personas discapacitadas, adultos mayores y población indígena. Una de las problemáticas que se han identificado en adultos mayores es la discapacidad auditiva o visual, por lo cual es necesario proporcionar herramientas o artefactos tecnológicos para apoyarlos en su entorno de trabajo, familiar y social. A finales del 2013 el organismo internacional de las Naciones Unidas ha reportado que para el año 2050 estiman más de 2000 millones de adultos mayores usarán Internet, siendo un aumento del 300% [2]. Esto conlleva que el crecimiento de esta población hará uso de distintas aplicaciones y dispositivos tecnológicos con mayor frecuencia en las actividades cotidianas y que requieren ser adaptadas a sus necesidades y condiciones personales.

En la actualidad por medio de las redes sociales y de distintas aplicaciones móviles se puede compartir información, siendo uno de los principales problemas actuales la seguridad de la información. En el caso de los adultos mayores el acceso a la tecnología no es fácil por la complejidad que muestran las interfaces y en especial las aplicaciones de localización, siendo este grupo vulnerable a extraviarse y requerir ayuda.

El presente trabajo tiene como *propósito el diseño y desarrollo de una Aplicación Móvil para dispositivos con sistema Android con la finalidad de compartir la ubicación de manera rápida y sencilla mediante un código y el uso de interfaces accesibles para el Adulto Mayor, para estar disponible en el Portal de Servicios a Grupos Vulnerables de la Institución como parte de un proyecto comunitario y social.*

La aplicación permitirá ubicar al usuario de manera exacta y en caso de moverse, el usuario podrá arrastrar el marcador para hacer coincidir con la ubicación, generar un código único, actualizado y almacenarlo en la base de datos. El Web Service permitirá la obtención de la información y mostrar la ubicación real del usuario adulto con los usuarios que se compartió la información, logrando el uso del dispositivo móvil de forma más accesible, segura mediante interfaces intuitivas y sencillas.

El artículo se organiza de la siguiente manera: En la sección 2 presenta la fundamentación teórica y el estado del arte de este trabajo. En la sección 3 se define el análisis y diseño de la aplicación móvil. En la sección 4 Se presenta la propuesta de la aplicación móvil, la implementación y pruebas de usabilidad. Finalmente, las conclusiones y el trabajo futuro se presentan en la sección 5.

2. Marco teórico

En esta sección revisaremos los tópicos de discapacidad y la tecnología, aplicaciones móviles y el diseño centrado en el usuario, los cuales son relevantes en nuestra propuesta de este trabajo.

2.1. Discapacidad y tecnologías

En general la tecnología ha tenido que proponer adaptaciones para poder ser incluyente y accesible a personas con distintas discapacidades. Por otra parte, los adultos mayores requieren ser alfabetizados en el uso de las Tecnologías de Información y Comunicación. En ese sentido ha surgido la Tecnología de Apoyo conocida también como “tecnología de adaptación o rehabilitación”. Según Alcantud la define como “Todos aquellos aparatos, utensilios, herramientas, programas de ordenador o servicios de apoyo que tienen como objetivo incrementar las capacidades de las personas que, por cualquier circunstancia, no alcanzan los niveles medios de ejecución que por su edad y sexo le corresponderían en relación con la población normal” [3]. De acuerdo al Artículo 2 de la Convención de Naciones Unidas sobre los Derechos de las Personas con Discapacidad, aprobada en 2006 [4], existe la necesidad que se diseñen aplicaciones bajo un diseño universal para todos pero que no excluye las ayudas técnicas para grupo particulares de personas con discapacidad, cuando se necesiten que se adapten a las capacidades específicas y que con respecto a la accesibilidad se enfatice tanto para contenidos y servicios como para herramientas y dispositivos.

2.2. Aplicaciones móviles accesibles

Según Gil [5], las Aplicaciones Móviles o conocidas como Apps, han establecido principios básicos para el diseño de Apps Accesibles de acuerdo con la norma “UNE 139802:2009 - Requisitos de accesibilidad del software. (ISO 9241-171:2008)”, los cuales han sido adaptados a las necesidades de los dispositivos móviles.

Las aplicaciones móviles accesibles han sido definidas como “una aplicación es accesible cuando cualquier usuario, independientemente de su diversidad funcional, puede utilizarla en su dispositivo móvil satisfactoriamente con su sistema de acceso habitual” [5].

Una aplicación móvil se identifica en la actualidad porque puede funcionar en distintos dispositivos móviles, y existen tres tipos [6]:

- Aplicaciones nativas, son aquellas que se desarrollan bajo un lenguaje y un entorno de desarrollo en específico para cada sistema operativo.
- Aplicaciones web, estas son desarrolladas usando lenguajes para el desarrollo web como lo son HTML, CSS y JavaScript.
- Aplicaciones Híbridas, como su nombre lo indica tienen un poco de cada tipo de las aplicaciones anteriores.

Las aplicaciones móviles las podemos encontrar en las diferentes tiendas como: AppStore, GooglePlay y Windows Phone Store, quienes ofrecen servicios de compras en línea, educación, finanzas, navegación. Las aplicaciones de geolocalización ofrecen servicios para facilitar las tareas como: mostrar el tráfico, optimización de rutas, lugares de interés social, recordatorios al pasar cierta zona, ver los lugares más visitados, buscar direcciones y otros. Las aplicaciones más usadas en este contexto son:

Google Maps es una aplicación móvil de mapeo desarrollado por Google para los sistemas operativos Android e iOS que utiliza la API de Google Maps para su información. Los mapas e información no están incluidos en los mapas de Google instaladas para los dispositivos, se requiere una conexión a Internet para hacer uso de ellos y permite al usuario descargar la ruta de un punto de referencia a otro, con un área de 10 millas cuadradas (26 km²) alrededor de cualquier punto [7].

Waze es la aplicación de tráfico y navegación basada en la comunidad más grande del mundo, dando información de tráfico y ruta en tiempo real por sus usuarios [7].

Foursquare consiste en una red social en la que los usuarios realizan ‘check-ins’ en los lugares que visitan mientras comparten recomendaciones y opiniones acerca de estos con sus contactos. Para poder hacerlo de manera efectiva hace uso del GPS del dispositivo para realizar check-ins de la forma más sencilla posible [7].

Una aplicación llamada Life 360, es un localizador familiar en dos versiones: básica (gratuita) y Premium (con costo), siendo esta la cual integra mejores servicios de localización, pero el costo es elevado [8].

2.3. Sistemas de posicionamiento

Los sistemas de posicionamiento se refieren al método por el cual se puede determinar la posición de un dispositivo, y utilizando servicios adicionales se puede responder preguntas como ¿Dónde estoy?, ¿Qué hay cerca?, ¿Cómo llego allá?, entre otras.

La ubicación puede ser expresada como una descripción textual o en términos espaciales [9]. Una ubicación expresada como un texto es por lo general la dirección de una calle, la ciudad, comuna, código postal, etc. Una ubicación espacial puede ser expresada en términos de coordenadas geográficas usando latitud-longitud-altitud (esta última opcional). La latitud se expresa en grados de 0-90 al norte o sur del ecuador, la longitud en grados de 0-180 al este u oeste del meridiano de Greenwich, y la altitud en metros sobre el nivel del mar. Existen varios métodos [9] para determinar la ubicación de un dispositivo, los cuales varían en las tecnologías empleadas y en la precisión con que entregan el resultado. Siendo las más utilizadas las redes telefónicas, celulares, de área local, Wi-Fi, Bluetooth y GPS entre otros.

Con respecto al *método Cell-ID*, es considerado de bajo costo que utiliza la red GSM para determinar la ubicación de un teléfono celular. El identificador de la antena o célula (Cell ID) a la que está conectado el dispositivo es usado para aproximar la ubicación del usuario. La precisión de este método depende del tamaño de la célula: unos 500 metros de exactitud dentro de una ciudad y unos 10 kilómetros en zonas rurales.

Por otra parte, el *método Bluetooth*, es una red de corto alcance, este provee un posicionamiento de buena precisión, con máximo de error de unos 10 metros, además de poder entregar posicionamiento vertical. Es muy adecuado para obtener la ubicación dentro de edificios o zonas urbanas pequeñas.

2.4. Diseño centrado en el usuario

El diseño centrado en el usuario (DCU) es aquel que, desde los inicios del desarrollo del producto, hace participar al usuario y lo involucra durante todo el proceso [10]. La Usability Professionals' Association (UPA) lo define como un enfoque de diseño cuyo proceso está dirigido por información sobre las personas que van a hacer uso del producto, siendo una guía para facilitar el logro de una mayor calidad en el uso de sistemas interactivos que consta de cuatro fases, como se muestra en la Figura 1.



Fig. 1. Proceso del diseño centrado en el usuario.

3. Análisis y diseño de la aplicación móvil de localización “TENAMIKI”

En esta sección se presenta el análisis, diseño, la arquitectura de la aplicación y el diseño de la base de datos, bajo el enfoque de la metodología ágil de la Programación Extrema y el diseño centrado en el usuario.

3.1. Metodología ágil de la programación extrema

Para el análisis y el diseño de la aplicación se utilizó la programación Extrema y sus fases [11], ya que permite de forma más sencilla el desarrollo de aplicaciones móviles con la fase de exploración, se determinaron los requerimientos por medio de las *Historias de usuarios* del grupo focal: “Adultos mayores discapacitados” y se integraron los casos de uso en UML para la aplicación [11].

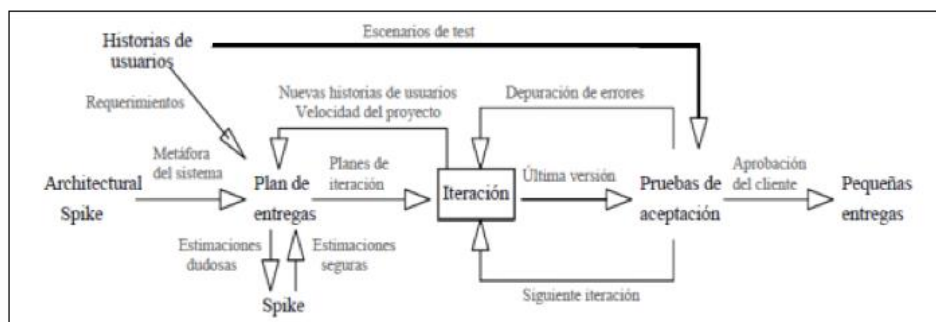


Fig. 2. Ciclo de vida de un proyecto con XP.

En la Figura 3, se muestra al usuario “Adulto Mayor” con las actividades a desempeñar y la interacción que tiene con la base de datos del dispositivo y el Web Service.

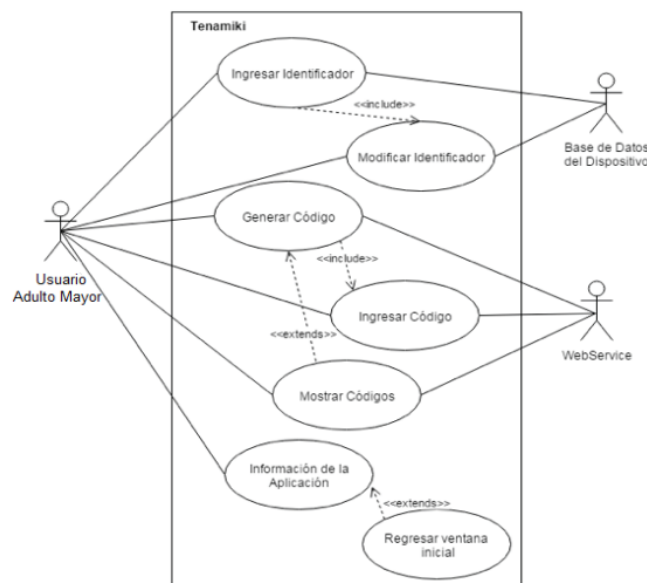


Fig. 3. Caso de uso del sistema.

3.2. Arquitectura de la aplicación

La arquitectura de la aplicación consta de dos partes: la primera parte es la aplicación para dispositivos Android, la cual estará conectada al Web Service [12], que a su vez se comunica con la base de datos que almacena e intercambia información de la aplicación. Por otra parte, se consta de un módulo Web agregado al servidor web con Apache para el portal de aplicaciones de apoyo a Adultos Mayores Accesibles (AAMA) de la institución, como se muestra en la Figura 4.



Fig. 4. Esquema de la arquitectura de la Aplicación TENAMIKI.

3.3. Diseño centrado en el usuario de la aplicación

El Diseño centrado en el usuario se adoptó a las características propias del usuario de Adulto Mayor y teniendo en cuenta que pueden tener alguna discapacidad auditiva. Así como al desarrollo de *aplicaciones móviles que requieren ser accesibles para todos*. Para lo cual se realizó una serie de entrevistas y encuestas con el grupo focal, finalmente se plantearán los prototipos de la aplicación, los cuales se realizaron con el Software Balsamic.

Los resultados obtenidos fue que la mayoría de los adultos ya poseen un teléfono inteligente, siendo el sistema operativo que más predomina el Android y de las herramientas que han escuchado y muy poco utilizado es Google Maps y Waze, pero aseguran que requieren no solo compartir su localización, sino tener certeza de que hayan recibido la localización sus familiares o conocidos y saber cuántos usuarios han consultado la ubicación que han compartido y poder llevar un registro incluso de las personas contactadas por este medio y los datos de la localización para dar un servicio de información detallado desde el portal y al usuario de la aplicación. Por lo cual se planteó las características principales con las que debe contar la aplicación:

Interface clara e intuitiva: Puesto que lo han de utilizar personas con muy pocas habilidades tecnológicas, se requiere de una interface gráfica muy clara e intuitiva, con pocos botones y que los que haya sean grandes.

Pocas opciones hasta llegar al resultado: Se quiere que la aplicación sea ágil, así que se necesita pocas pantallas y decisiones, para llegar al resultado esperado.

Práctica: La aplicación debe resultar lo más práctica posible para que el usuario la encuentre de gran utilidad.

4. Desarrollo y pruebas del prototipo de la aplicación móvil

En esta sección se muestran los resultados del prototipo de la implementación de la aplicación móvil, así como las pruebas de funcionalidad junto con las pruebas de usabilidad obtenidas del grupo focal de usuarios. Las tecnologías utilizadas PHP, MySQL, Webservices, Java, Android Studio 15, Servidor Web: Apache 2.4.10, Emulador Android: Genymotio 2.5.3. Las interfaces finales para la aplicación móvil TENAMIKI (significa en náhuatl “Encontrarse”) como se muestran en las Figura 5.



Fig. 5. Interfaces del sistema móvil de ingreso y de acceso al usuario.

4.1. Prueba piloto y de usabilidad

La aplicación móvil se le aplicó una prueba piloto un grupo de 15 usuarios adultos mayores, de los cuales el 80% tiene una disminución auditiva sin llegar a una sordera total. Para la evaluación de la aplicación, se realizaron pruebas de funcionalidad y usabilidad. Con respecto a las pruebas de Usabilidad presentaron tres posibles escenarios posibles para los usuarios finales, a continuación, se describen:

Situación 1: Al usuario se le entregó la aplicación con la WiFi, GPS y 3G apagados, además se les dio una breve explicación del uso de la aplicación.

Situación 2: Al usuario se le entregó la aplicación con la WiFi, GPS y 3G prendidos y se le explicó el uso de la aplicación.

Situación 3: Al usuario se le entregó la aplicación con la WiFi, GPS y 3G apagados y no se explicó el uso de la aplicación

Para cada una de las situaciones los usuarios deberán cumplir ciertas tareas para así comprobar el funcionamiento de la aplicación.

Tarea 1: Ingresar un código generado.

Tarea 2: Encontrar donde se muestran los códigos que se han generado.

Tarea 3: Generar un código en una localización diferente a la que se encuentra actualmente.

Tarea 4: Cambiar el identificador.

Tarea 5: Compartir el código.

Después de realizar las pruebas a 15 usuarios 5 por cada situación, se muestran los resultados obtenidos en la Tabla 1, donde se observa el porcentaje de usuarios que pudieron completar la tarea.

Tabla 1. Resultados prueba de usabilidad usando escenarios.

Tareas	Escenarios de Test		
	Situación 1	Situación 2	Situación 3
Tarea no.1	90%	95%	60%
Tarea no.2	100%	100%	100%
Tarea no.3	100%	100%	100%
Tarea no.4	100%	100%	100%
Tarea no.5	80%	100%	80%
Promedio	94%	99%	88%

Lo cual refleja que los usuarios con una breve explicación de la Situación 1 su desempeño fue del 94% del cumplimiento de las tareas mientras que los usuarios de la Situación 2 al 99% realizaron las tareas casi en su totalidad y para la situación 3 lograron las tareas en un 88%, esto implica que se presenta una interfaz intuitiva y agradable.

Finalmente se aplicó una encuesta de satisfacción “Valoración del Software [13], la cual evalúa siete criterios: Navegación, Interactividad, Inmersión, Usabilidad, Creatividad, Efectividad y Calidad, con una escala de 1 a 5, cuyo promedio obtenido fue de 4.7 (94.37%), a continuación, se presentan los resultados en la Figura 6.

5. Conclusiones y trabajo a futuro

El propósito de este artículo fue presentar una propuesta de una aplicación móvil para poder apoyar a los “Adultos Mayores” en la movilidad de sus actividades cotidianas de manera segura poder compartir su ubicación mediante el uso de un código generado y reconocer el status de los usuarios que han consultado la información, para poder ser localizados de forma más segura, logrando una aplicación intuitiva y agradable. La metodología Programación Extrema y el Diseño Centrado en el Usuario nos permitieron mostrar trabajar de forma más cercana con el usuario y generar algunas versiones previas para que el usuario probara y revisara los avances de acuerdo a sus necesidades y trabajar con un grupo focal de “Adultos Mayores” para centrarnos en una aplicación más agradable y de fácil manejo para este tipo de población.

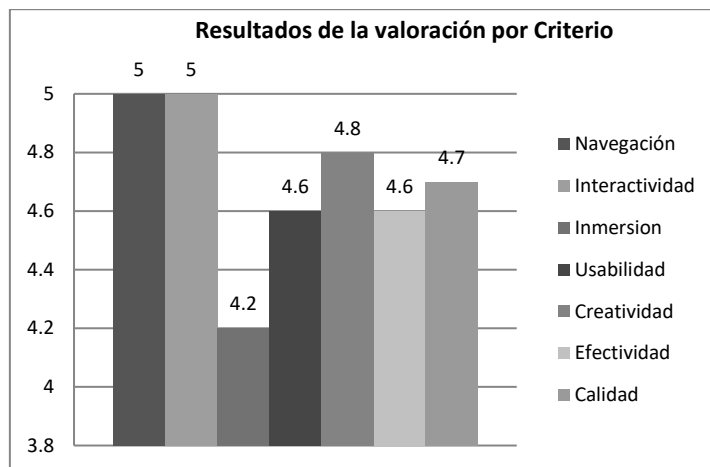


Fig. 6. Resultados de la encuesta de satisfacción.

Una de las aportaciones de esta propuesta de aplicación móvil es la facilidad de exportar la opción del código que contiene las coordenadas para otras aplicaciones y compartir entre los distintos usuarios o familiares de tal manera que los usuarios que utilicen la aplicación puedan ubicar el destino sin necesidad de buscarlo nuevamente, solo bastaría ingresar el código para poder obtener la información y de esta forma llevar un registro de los usuarios en la base de datos, quienes tienen acceso a dicha información por medio de las tecnologías de Webservice, PHP, MySQL y Java.

Como trabajo futuro se espera incorporar otras pruebas de usabilidad con una muestra mayor de usuarios y una vez que se tenga una versión estable la aplicación, se planea incorporar al “Portal de Discapacitados” a la sección de Aplicaciones Accesibles para uso de la comunidad y apoyo a grupos vulnerables.

Referencias

1. Muñoz, C.: Bienestar Subjetivo y Actividad Social con Sentido Histórico en Adultos Mayores. Revista Hacia la Promoción de la Salud, <http://google.redalyc.org/articulo.oa?id=309131077002>, 13–26
2. Naciones, U.: La sostenibilidad y la inclusión de las personas mayores en el entorno urbano. <http://www.un.org/es/events/olderpersonsday/> (2015)
3. Alcantud, F.: Estudiantes con Discapacidades Integrados en los Estudios Universitarios. Notas para su Orientación (1995)
4. Naciones, U.: Convención sobre los Derechos de las Personas con Discapacidad. <http://www.un.org/esa/socdev/enable/documents/tccconvs.pdf> (2015)
5. Gil, S.: Cómo hacer Apps Accesibles. Centro de Referencia Estatal de Autonomía Personal y Ayudas Técnica, http://www.ceapat.es/ceapat_01/index.htm, 5–85 (2013)
6. Cuello, J.: Diseñando apps para móviles. <http://appdesignbook.com/> (2015)

7. Fuentes, D., Opitz, I., Oro, N.: Análisis de Aplicación Móvil basada en Geolocalización. http://profesores.elo.ut fsm.cl/~agv/elo322/1s16/projects/reports/Informe_G14_Elo322.pdf (2015)
8. Life360. <https://play.google.com/store/apps/details?id=com.life360.android.safetymapd&hl=es> (2015)
9. Bernardos, A., Besada, J., Casar, C.: Tecnologías de localización, ETS Ingenieros de Telecomunicación. Universidad Politécnica de Madrid, http://www.upm.es/sfs/Rectorado/Organos%20de%20Gobierno/Consejo%20Social/Actividades/tecnologias_servicios_para_sociedad_informacion.pdf (2005)
10. Hassan-Montero, Y., Ortega-Santamaría, S. Informe APEI sobre Usabilidad. Gijón: Asociación Profesional de Especialistas en Información. <http://www.nosolousabilidad.com/manual/index.htm> (2009)
11. Extreme Programming. <http://www.extremeprogramming.org/> (2013)
12. Barry, D.: Service Architecture. (s.f.), <http://www.service-architecture.com/articles/web-services/soap.html>
13. Acuña, A., Romo, M.: Diseño Instruccional Multimedia. Pearson Education, México (2011)

Herramienta web para almacenar y visualizar objetos distribuidos

José M. Hernández, Carlos R. Jaimez

Universidad Autónoma Metropolitana, Unidad Cuajimalpa,
Departamento de Tecnologías de la Información, México

2123064604@alumnos.cua.uam.mx, cjaimez@correo.cua.uam.mx

Resumen. En los últimos años ha habido un incremento en el desarrollo de aplicaciones que permiten la administración de diversos tipos de recursos a través de web, debido a que sus usuarios únicamente requieren de un navegador web para acceder a la aplicación. Los objetos distribuidos son un ejemplo del tipo de recursos que se beneficiarían de una aplicación que permitiera tener acceso a ellos en línea. En este artículo se presenta una herramienta web que permite almacenar objetos, así como visualizar su estado y métodos. Esta herramienta web es un complemento de *Web Objects in XML* (WOX), el cual es un *framework* para programación de aplicaciones distribuidas basadas en objetos.

Palabras clave: Objetos remotos, objetos distribuidos, web objects in XML, visualización de objetos, almacenamiento de objetos.

Web Tool for Storing and Visualizing Distributed Objects

Abstract. In recent years there has been an increase in the development of applications that allow the administration of several types of resources through the web, because their users only require a web browser for accessing the application. Distributed objects are an example of the type of resources that would benefit of an application that allow to have online access to them. This paper presents a web tool that allows storing objects, as well as visualizing their state and methods. This web tool is a plug-in for *Web Objects in XML* (WOX), which is a *framework* for programming object-based distributed applications.

Keywords. Remote objects, distributed objects, web objects in XML, visualization of objects, object storage.

1. Introducción

La herramienta web que se presenta en este artículo es un complemento de *Web Objects in XML* (WOX) [1], el cual es un *framework* para programación de aplicacio-

nes distribuidas basadas en objetos. WOX permite construir sistemas distribuidos, utiliza XML como representación para los objetos y mensajes intercambiados [2], y proporciona comunicación síncrona y asíncrona entre clientes y servidores [3]. WOX tiene características especiales tomadas de dos paradigmas utilizados para construir sistemas distribuidos: el paradigma basado en objetos y el paradigma basado en web. En WOX es posible almacenar objetos Java y C#, los cuales pueden ser generados por aplicaciones locales o distribuidas.

El resto del artículo está organizado de la siguiente manera. El trabajo relacionado se presenta en la sección 2 con un análisis comparativo de diferentes herramientas que tienen un propósito similar al de la herramienta web que se describe en este artículo. En la sección 3 se proporciona una introducción del *framework* WOX. En la sección 4 se describe la funcionalidad de la herramienta web, su estructura e interfaz. La sección 5 se encarga de mostrar la herramienta web en funcionamiento, en particular se presenta el almacenamiento de objetos y el repositorio para visualización de objetos. Finalmente, en la sección 6 se presentan las conclusiones y el trabajo futuro.

2. Trabajo relacionado

En esta sección se describen las características y el funcionamiento de algunas herramientas que son similares a la herramienta web propuesta en este artículo. Las herramientas analizadas son las siguientes: CORBAWeb [4], SopView+ [5], PESTO [6], CORBA Object Browser [7] y Apache Axis2 [8]. Al final de la sección se presenta una comparación de estas herramientas, junto con una breve descripción de las características que se tomaron en cuenta.

2.1. CORBAWeb

Es una herramienta [4] que actúa como una puerta de enlace entre la web y CORBA (*Common Object Request Broker Architecture*, CORBA por sus siglas en inglés). A esta herramienta se le conoce como un navegador de objetos genérico, ya que su idea es lograr que los clientes puedan visualizar e invocar métodos sobre cualquier objeto CORBA local o remoto, a través de un navegador web. Con CORBAWeb un cliente puede navegar a través de enlaces de objetos CORBA utilizando URL generados dinámicamente para cada objeto remoto.

CORBAWeb funciona mediante un navegador web, en el que un usuario puede acceder e invocar métodos de objetos remotos que residan en un servidor, para ello se generan automáticamente formularios HTML a partir del lenguaje de definición de interfaces (*Interface Definition Language*, IDL por sus siglas en inglés), las cuales permiten la invocación de métodos de cualquier objeto CORBA. CORBAWeb recibe las acciones del usuario y las traduce, accede al objeto remoto requerido para invocar el método del objeto, obtiene el resultado y finalmente devuelve un documento HTML que contiene los resultados de la invocación.

2.2. SOPView+

Es un proyecto [5] de la Universidad Nacional de Seúl, Corea del Sur, desarrollado en un entorno UNIX y utiliza la herramienta *Motif widget* para poder proporcionar una interfaz gráfica. Este proyecto busca crear un navegador y visualizador de objetos, principalmente para la consulta y gestión de bases de datos orientadas a objetos. Con esta herramienta se puede explorar dentro de la base de datos con el fin de localizar algún objeto deseado, recuperar la información y visualizarla de forma gráfica. Además, esta herramienta visualiza los objetos de una forma jerárquica y permite la navegación en grandes bases de datos cambiando el objeto base, el cual es un objeto que se vuelve el nodo principal por el que se inicia la navegación.

SOPView+ permite a los usuarios cambiar el objeto base a lo largo de la búsqueda de objetos en las bases de datos; para ello se pone un *ancla* sobre el objeto. Esto hace que sea posible para los usuarios explorar objetos en una gran base de datos de una forma más sencilla.

2.3. PESTO

Esta herramienta [6] se crea a partir del proyecto GARLIC [9], cuyo objetivo es construir un sistema de información capaz de integrar datos que residen en diferentes sistemas de bases de datos. Los datos pueden consultarse a través de un lenguaje similar a SQL, el cual se ha ampliado para incluir características orientadas a objetos. El proyecto GARLIC busca proporcionar una interfaz novedosa que ofrezca la consulta y navegación de objetos, llamada PESTO, que es el trabajo conjunto entre IBM y la Universidad de Wisconsin - Madison, en donde trabajan en el desarrollo de una interfaz para usuarios.

El explorador portable de objetos estructurados (*Portable Explorer of Structured Objects*, PESTO por sus siglas en inglés) ofrece una interfaz propia en la que se puede navegar sobre objetos, y al igual que SOPView+, está diseñado para explorar bases de datos de objetos y ofrece la posibilidad de hacer consultas complejas de objetos en bases de datos, a través de un lenguaje similar a SQL gracias al paradigma *query in place*. Esta herramienta proporciona a los usuarios una interfaz que permite la navegación interactiva y la consulta de los contenidos de las bases de datos de GARLIC. Al igual que SOPView+, PESTO permite a los usuarios moverse hacia delante y hacia atrás en la base de datos, para consultar y navegar por los objetos, los cuales se muestran en una jerarquía de referencia que puede representarse por la conexión de nodos de objetos a través de enlaces.

2.4. CORBA object browser

Esta herramienta [7] tiene como finalidad acceder a objetos CORBA directamente desde un navegador web utilizando un esquema URI; permite la navegación e invocación sobre objetos CORBA de la misma forma en que los usuarios navegan en Internet. En esta herramienta es posible visualizar y ejecutar los métodos de un objeto determinado desde un navegador web; sin embargo, para realizar lo anterior, se debe

acceder mediante un prototipo de navegador llamado HotJava Web Browser, el cual ya no está disponible actualmente.

La ventaja del navegador HotJava Web Browser es que contaba con un mecanismo para acceder a objetos CORBA que se ejecutan en un ORB (*Object Request Broker*, ORB por sus siglas en inglés) seguro, a través del CORBA Object Browser. El acceso a objetos seguros a través del navegador requería la autenticación con el ORB remoto y también tener una comunicación segura.

2.5. Apache axis2

Es un motor de servicios web [8] desarrollado por *Apache Software Foundation*, el cual es interoperable, implementado en C++ y en Java, en el que se pueden crear aplicaciones interoperables y distribuidas. Al igual que WOX, Apache Axis2 es un *framework* de código abierto, basado en XML, que utiliza SOAP para el intercambio de mensajes. A pesar de que Apache Axis2 trabaja con objetos, no guarda el estado de los objetos, por lo que sus métodos solo se invocan como si fueran métodos estáticos. Una característica particular de esta herramienta es que a pesar de que ofrece la ejecución de métodos sobre objetos, no cuenta con una interfaz para sus usuarios, ya que es necesario llamar a los objetos por medio de su URL.

2.6. Tabla comparativa

En la Tabla 1 se presenta una comparación de las herramientas analizadas en el estado del arte: H1) CORBAWeb, H2) SOPView+, H3) PESTO, H4) CORBA Object Browser y H5) Apache Axis2. Se muestra un símbolo de verificación si la herramienta tiene la característica, o una x si no la tiene. A continuación se proporciona una breve descripción de las características que se tomaron en cuenta.

Tabla 1. Comparación de características de las herramientas analizadas.

Características	H1	H2	H3	H4	H5
Código libre	✓	✗	✗	✓	✓
Interoperabilidad	✓	✗	✗	✓	✓
Basado en objetos	✓	✓	✓	✓	✗
Servicio web	✓	✗	✗	✗	✓
Uso de XML	✗	✗	✗	✗	✓
Visualización de objetos	✓	✓	✓	✓	✗
Visualización de atributos	✗	✓	✓	✓	✗
Ejecución de métodos	✓	✗	✗	✓	✓
Interfaz web	✓	✗	✗	✓	✗
Manejo de bases de objetos	✗	✓	✓	✗	✗

Código libre. Se refiere al software distribuido y desarrollado libremente, es gratuito y cualquier usuario puede tener acceso al código fuente.

Interoperabilidad. Esta característica se refiere a que la herramienta es capaz de comunicarse en diferentes plataformas o lenguajes de programación.

Basado en objetos. Se refiere a que la herramienta utiliza objetos remotos.

Servicio web. Se refiere a que la herramienta está enfocada a servicios web.

Uso de XML. Esto significa que la herramienta utiliza XML como medio de comunicación entre el cliente y el servidor.

Visualización de objetos. Esta característica se refiere a que la herramienta analizada permite una forma de visualización de los objetos de forma gráfica.

Visualización de atributos de los objetos. Esta característica se refiere a que la herramienta analizada puede mostrar los atributos que contiene cada objeto.

Ejecución métodos. Se refiere a que se permite la ejecución de métodos pertenecientes a un objeto a través de un navegador web.

Interfaz web. Esta característica se refiere a que la herramienta analizada proporciona una interfaz web en la que el usuario puede visualizar los métodos de los objetos remotos.

Manejo de bases de objetos. Quiere decir que la herramienta es capaz de navegar por bases de objetos o repositorios.

3. Objetos web en XML

Esta sección proporciona una breve introducción del *framework* WOX, el cual combina características de los sistemas distribuidos basados en objetos y los sistemas distribuidos basados en web. Algunas de las características de este *framework* se presentan en los siguientes párrafos.

WOX utiliza URLs para identificar objetos remotos de manera única, siguiendo los principios de la arquitectura de transferencia de estados de representación (*Representational State Transfer*, REST por sus siglas en inglés) [10]. Esta es una característica importante porque todos los objetos son identificados de manera única por su URL y pueden ser accedidos desde cualquier lugar en la web, ya sea a través de un navegador web o mediante programación.

WOX emplea un serializador eficiente y fácil de usar, llamado serializador WOX [2], el cual es la base del *framework* para serializar objetos, solicitudes y respuestas intercambiadas entre clientes y servidores. Este serializador es una biblioteca independiente basada en XML, la cual es capaz de serializar objetos Java y C# a XML y viceversa. Una de sus principales características es la generación de XML estándar para objetos, el cual es independiente del lenguaje de programación y permite alcanzar interoperabilidad entre diferentes lenguajes de programación orientados a objetos. Aplicaciones escritas en los lenguajes Java y C# pueden interoperar.

WOX tiene un conjunto de operaciones estándar y especiales que se aplican sobre objetos locales y remotos. Estas operaciones incluyen la solicitud de referencias remotas, invocación de métodos estáticos (llamadas de servicios web), invocación de métodos de instancia, destrucción de objetos, solicitud de copias, duplicación de objetos, actualización de objetos, subida de objetos, invocación de métodos asíncronos,

entre otras. Algunas de las operaciones son descritas en [1]. El mecanismo utilizado por WOX en una invocación a un método se ilustra en la Figura 1.

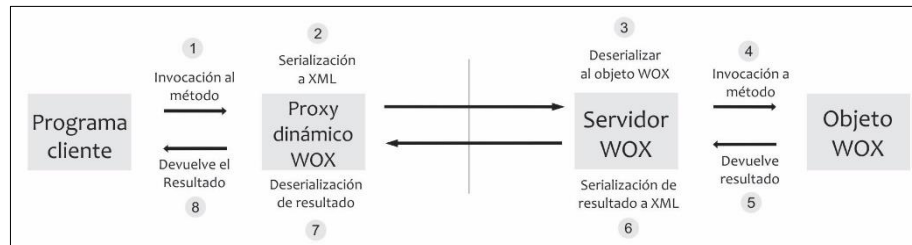


Fig. 1. Mecanismo de invocación de método remoto en WOX.

La serie de pasos llevados a cabo en la invocación de un método en WOX es la siguiente: 1) El programa cliente WOX invoca un método sobre una referencia remota (la forma en la que el cliente invoca un método sobre una referencia remota es exactamente la misma forma que la que utiliza para un objeto local); 2) El proxy dinámico WOX toma la solicitud, la serializa a XML y la manda a través de la red al servidor WOX; 3) El servidor WOX toma la solicitud y la deserializa a un objeto WOX; 4) El servidor WOX carga el objeto y ejecuta el método sobre él; 5) El resultado de la invocación del método es regresada al servidor WOX; 6) El servidor WOX serializa el resultado a XML y es regresado al cliente, ya sea el resultado real o una referencia a éste (el resultado es guardado en el servidor en caso de que una referencia haya sido enviada); 7) El proxy dinámico WOX recibe el resultado y lo deserializa al objeto apropiado (objeto real o referencia remota); 8) El proxy dinámico WOX regresa el resultado al programa cliente WOX.

Desde el punto de vista del programa cliente WOX, éste solamente realiza la invocación del método y obtiene el resultado de regreso de forma transparente. Las bibliotecas cliente WOX llevan a cabo el proceso de serializar la solicitud y mandarla al servidor WOX, así como recibir el resultado de la invocación del método y deserializarla. Las siguientes secciones presentan la herramienta web que permite almacenar y visualizar objetos WOX a través de un navegador web.

4. Herramienta web propuesta

En esta sección se describe la funcionalidad de la herramienta web desarrollada, su estructura e interfaz.

4.1. Funcionalidad

La funcionalidad de la herramienta web propuesta se describe a continuación, a través de una serie de acciones que se pueden llevar a cabo en ella.

Acceso a un objeto WOX específico a través de la red. La herramienta web permite a sus usuarios acceder a un objeto WOX almacenado en el sistema a partir de

un URL único generado por el mismo sistema. Con dicho URL los usuarios pueden acceder directamente al objeto sin necesidad de entrar al repositorio y buscar su objeto deseado entre todos los existentes en el mismo repositorio.

Visualización de objetos de forma gráfica. Los usuarios pueden ver, a través de una interfaz, los atributos de un objeto WOX determinado. Se proporcionan dos vistas: la primera es una tabla que representa al objeto y dentro de ella se visualizan todos sus atributos sin importar su tipo de dato; la segunda es una vista del código XML que representa al objeto, el cual es visualizado con un formato que permite al usuario entender cada etiqueta del código XML.

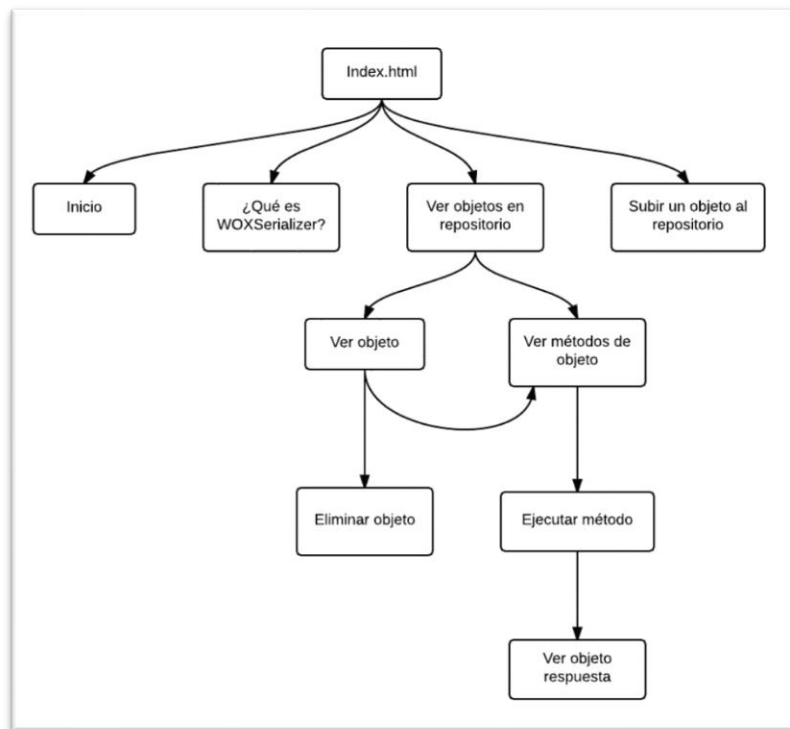


Fig. 2. Mapa de navegación de la herramienta web.

Visualización de métodos de cualquier objeto WOX. Los usuarios pueden visualizar todos los métodos que pertenecen a un objeto en particular, sin importar los parámetros que cada método requiera para ser invocado. Cada método cuenta con un espacio en el que se muestra la respuesta que devuelve el mismo al ser invocado.

Ejecución de cualquier método perteneciente a la clase de un objeto WOX. Mediante la herramienta web es posible invocar cualquier método de un objeto WOX; también se cuenta con una verificación de parámetros para evitar que el usuario introduzca valores de parámetros erróneos y así poder asegurar una correcta invocación de un método. Además, el sistema puede regresar una respuesta de la invocación de un método sin importar su tipo de dato. Es importante mencionar que para que la herra-

mienta web pueda ejecutar un método sobre un objeto, debe contar con la clase a la que pertenezca tal objeto.

Almacenamiento de objetos WOX en el servidor. Se proporciona a los usuarios la posibilidad de que ellos mismos puedan subir objetos WOX al servidor para su posterior manipulación, ya sea como parámetro para la invocación de un objeto o simplemente para almacenarlo en el repositorio para su posterior utilización.

4.2. Estructura

La herramienta web consiste de una interfaz web dinámica que permite a los usuarios el acceso a objetos WOX alojados en un repositorio. El usuario puede visualizar cada objeto existente en el repositorio de una forma gráfica, además de que se le permite ejecutar los métodos de cada objeto a través de un navegador web.

La Figura 2 muestra el mapa de navegación, en el que se representa la estructura de visualización que se emplea en la herramienta web.

4.3. Interfaz

La página de inicio de la herramienta web se muestra en la Figura 3, donde puede observarse el menú con las siguientes cuatro opciones: 1) Inicio; 2) ¿Qué es WOX serializer?; 3) Ver objetos en el servidor; y 4) Subir objeto al servidor.

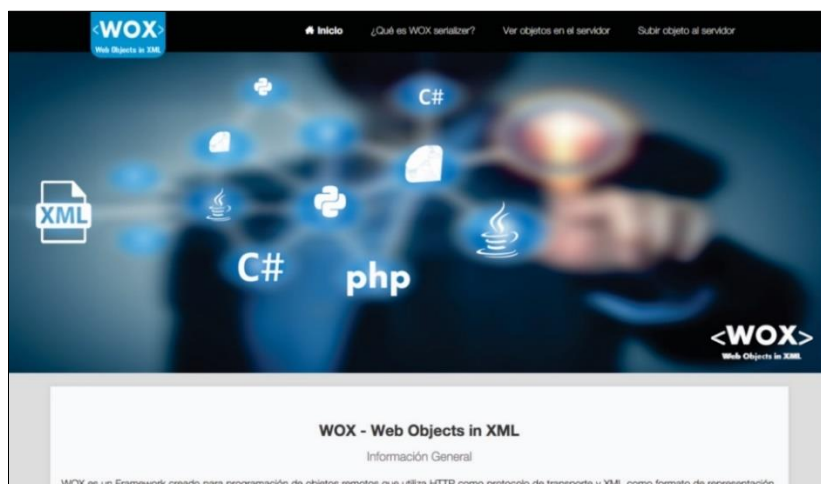


Fig. 3. Página de inicio de la herramienta web.

5. Herramienta web en funcionamiento

En esta sección se presenta la herramienta web en funcionamiento. Por restricciones de espacio, solamente se describe cómo se almacenan objetos en la herramienta web y cómo se visualizan a través del repositorio.

5.1. Almacenamiento de objetos

La Figura 4 muestra el funcionamiento de la herramienta web para almacenar objetos WOX, mediante una serie de pasos que se llevan a cabo entre cliente y servidor. El paso 1 es la serialización de un objeto con el serializador WOX, para lo cual puede utilizarse cualquiera de los lenguajes de programación soportados por WOX [11]; como segundo paso el usuario procede a subir el objeto al servidor al seleccionar la opción *Subir un objeto al servidor* del menú principal de la herramienta web; en el paso 3 el servidor se encarga del almacenamiento del objeto en el directorio llamado *UploadTemp* del repositorio de la herramienta web; en el paso 4 se procede a verificar el objeto; se deserializa para verificar que no contenga errores o que ya exista y dependiendo del resultado, se mueve el objeto a la papelera si éste tenía errores (paso 5.1), o se registra en el repositorio en caso de que sea un objeto sin errores (paso 5.2); finalmente, en el paso 6 se le comunica al usuario el resultado del almacenamiento del objeto WOX en el repositorio.

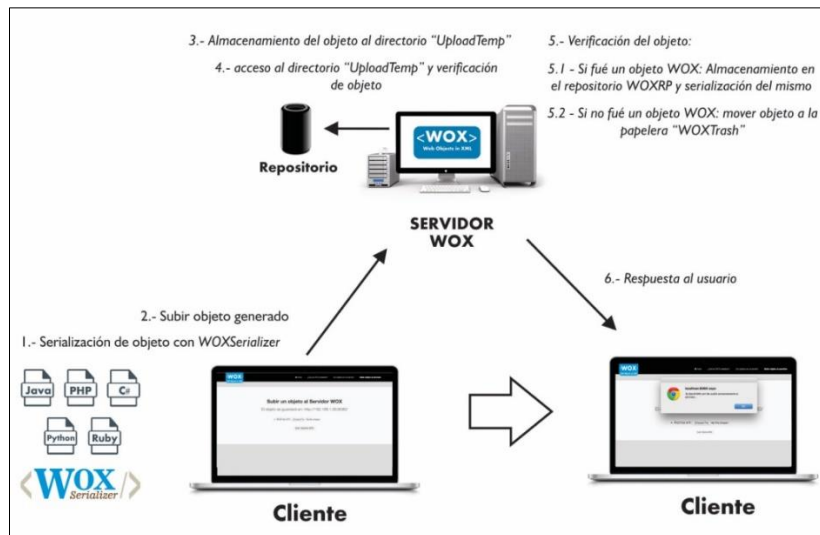


Fig. 4. Almacenamiento de un objeto WOX.

5.2. Visualización de objetos

La Figura 5 muestra el funcionamiento del repositorio para visualización de objetos WOX, mediante una serie de pasos que se llevan a cabo entre cliente y servidor. En el paso 1 el usuario accede al repositorio al seleccionar la opción *Ver objetos en el servidor* del menú principal de la herramienta web; en el segundo paso se deserializa el objeto *WOXR.xml*, el cual es una lista de objetos de la clase *WOXObject* que contiene la información de todos los objetos almacenados en el repositorio; en el paso 3, una vez deserializado el objeto *WOXR.xml*, se puede acceder a la información de los objetos registrados, por lo que se genera una tabla que contendrá el URL y la clase de

cada objeto registrado, así como dos opciones, una para visualizar el objeto y otra para acceder a sus métodos (paso 4); finalmente, en el paso 5 se muestra al usuario la tabla generada, para que éste elija el objeto deseado para su visualización.

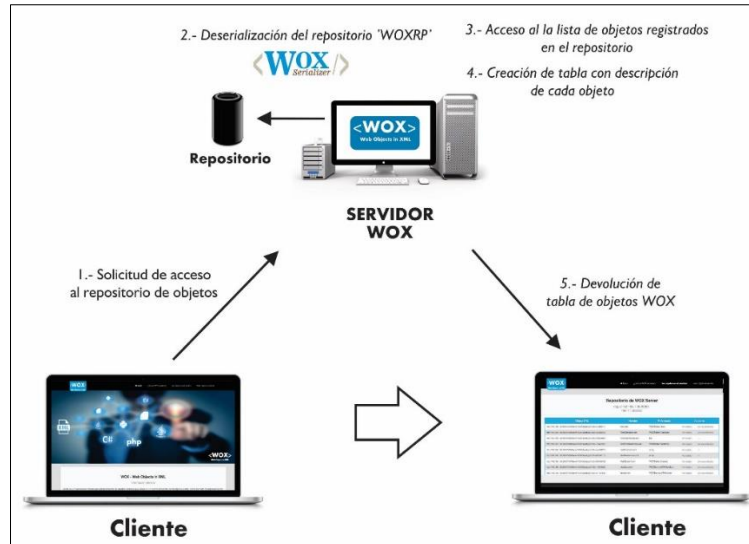


Fig. 5. Acceso al repositorio de objetos para su visualización.

Repositorio de WOX Server

http://192.168.1.90:8080/

Hay 11 objetos:

Object URL	Nombre	Referencia	Acciones	
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=312963317	libro.xml	WOXTester.Libro	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=324965425	TestComputer.xml	WOXTester.Computer	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=396534936	TestListCourses.xml	list	Ver objeto	---
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=735476957	TestPrimitiveArrays.xml	WOXTester.TestArray	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=554794375	TestProducts.xml	array	Ver objeto	---
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=662240713	TestReferences.xml	array	Ver objeto	---
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=666462496	TestStudent.xml	WOXTester.Student	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=140436062	visualizer.xml	WOXServer.WOXVisualizer	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=677257542	invoker.xml	WOXServer.WOXInvoker	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=242283885	WOXSer.xml	WOXServer.WOXServer	Ver objeto	Ver sus métodos
http://192.168.1.90:8080/WOXServer/WOXObject.jsp?objectId=057462666	B_GandhiMX.xml	WOXTester.Biblioteca	Ver objeto	Ver sus métodos

Fig. 6. Repositorio de objetos de la herramienta web.

En la interfaz de la herramienta web para acceder al repositorio de objetos el usuario debe seleccionar la opción *Ver objetos en el servidor*, la cual lo llevará a la página que se muestra en la captura de pantalla de la Figura 6. El repositorio de objetos tiene cuatro columnas: el URL del objeto en el servidor, el nombre del objeto, la referencia remota al objeto y las acciones que se pueden ejecutar sobre el objeto (*Ver objeto* y

Ver sus métodos). A través del hipervínculo *Ver objeto* se visualiza la representación gráfica del objeto y su representación en XML, como se ilustra en la Figura 7.

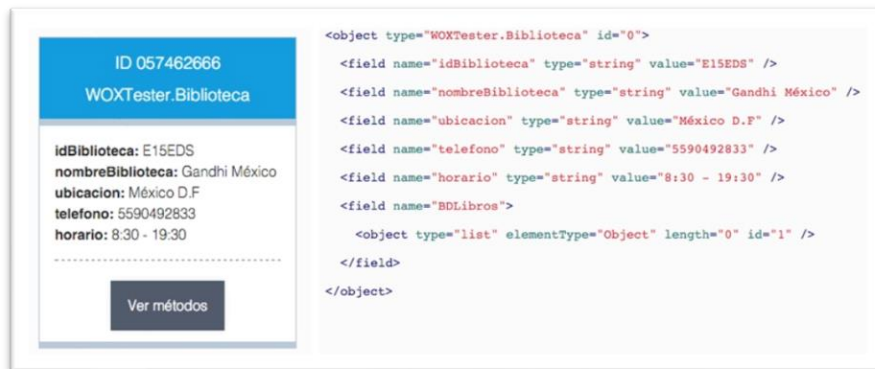


Fig. 7. Visualización de un objeto gráficamente y en XML.

6. Conclusiones y trabajo futuro

En este artículo se presentó una herramienta web que permite almacenar objetos, así como visualizar su estado y métodos. Esta herramienta web es un complemento de *Web Objects in XML (WOX)*, el cual es un *framework* para programación de aplicaciones distribuidas basadas en objetos.

La herramienta web cumplió con los objetivos planteados al inicio de su desarrollo, ya que se logró exitosamente diseñar e implementar el almacenamiento de objetos-WOX y proveer un repositorio de objetos a partir del cual es posible visualizar cada objeto que ha sido creado en éste, de manera gráfica y en su representación XML; todo ello a través de un navegador web. Asimismo, se implementó exitosamente la funcionalidad para la visualización de los métodos de cualquier objeto WOX que se encuentra en el repositorio, así como la ejecución de métodos de un objeto WOX a través de una interfaz donde se proporcionan los valores para cada uno de sus parámetros.

Como trabajo futuro es necesario afinar la implementación de la ejecución de métodos de un objeto WOX, ya que en esta versión de la herramienta web solamente es posible ejecutar métodos cuyos parámetros son de tipo primitivo. La ejecución de métodos que reciben objetos como parámetros no está finalizada.

Referencias

1. Jaimez-González, C., Lucas, S.: Implementing a State-based Application Using Web Objects in XML. In: Proceedings of the 9th International Symposium on Distributed Objects, Middleware, and Applications (DOA 2007), Lecture Notes in Computer Science, 4803, 577–594 (2007)

2. Jaimez-González, C., Lucas, S., López-Ornelas, E.: Easy XML Serialization of C# and Java Objects. In: Proceedings of the Balisage: The Markup Conference 2011, Balisage Series on Markup Technologies, 7 (2011)
3. Jaimez-González, C., Lucas, S., Luna-Ramírez, W.: A Web Tool for Monitoring HTTP Asynchronous Method Invocations. In: Proceedings of the 7th IEEE International Conference for Internet Technology and Secured Transactions, (ICITST-2012), 127–132 (2012)
4. Merle, P., Gransart, C., Geib, J.: CorbaWeb: A Generic Object Navigator. Disponible en: http://www.lifl.fr/~merle/papers/96_WWW5/paper/Overview.html
5. Chang, S., Kim, H.: SOPView+: An Object Browser Which Supports Navigating Database by Changing Base Object. In: Proceedings of the 21st International Conference on Computer Software and Applications Conference, COMPSAC 97 (1997)
6. Carey, M., Haas, L., Maganty, V., Williams, J.: PESTO: An Integrated Query/Browser for Object Databases. In: Proceedings of the 22th International Conference on Very Large Data Bases, Mumbai, India (1996)
7. Kumar, G., Jalote, P.: A Browser Front End for CORBA Objects. In: 10th International World Wide Web Conference (2001)
8. Apache Software Foundation. Web Services - Apache Axis, Disponible en: <http://ws.apache.org/axis/>
9. Tork, M., Arya, M., Haas, L., Carey, M., Cody, W., Fagin, R., Schwarz, P., Thomas, J., Wimmers, E.: The Garlic Project. In: Proceedings of the 1996 ACM SIGMOD International Conference on Management of Data (1996)
10. Fielding, R.: Architectural Styles and the Design of Network-based Software Architectures. PhD thesis, Disponible en: <http://www.ics.uci.edu/~fielding/pubs/dissertation/top.htm>(2000)
11. Web Objects in XML. Disponible en: <http://woxserializer.sourceforge.net/>

Sistema web personalizable con diseño adaptable para la administración de reservaciones en hoteles

Betzabet García-Mendoza, Carlos R. Jaimez-González

Universidad Autónoma Metropolitana, Unidad Cuajimalpa
Departamento de Tecnologías de la Información, México

209363465@alumnos.cua.uam.mx, cjaimez@correo.cua.uam.mx

Resumen. En este artículo se presenta el diseño e implementación de un sistema web personalizable con diseño adaptable para facilitar a administradores de hoteles pequeños el manejo de reservaciones de sus habitaciones, gestionando las ventas desde su propio sitio web; así como para facilitar y garantizar de inmediato las reservaciones para sus huéspedes potenciales. El sistema web que se presenta utiliza el patrón de diseño web adaptable, por lo que es posible visualizarlo automáticamente en diferentes dispositivos, tales como computadoras de escritorio, laptops, tabletas y teléfonos móviles.

Palabras clave: Sistema web personalizable, diseño web adaptable, control de reservaciones, comercio electrónico.

A Customizable Web System with Responsive Design for Administering Hotel Bookings

Abstract. This paper presents the design and implementation of a customizable web system with responsive design, for facilitating administrators of small hotels the control of room bookings, managing the sales from their own web site; and facilitates and guarantees the bookings for their potential guests. The web system presented uses the responsive web design pattern in order to be automatically visualized in different devices, such as desktop computers, laptops, tablets and mobile phones.

Keywords. Customizable web system, responsive web design, room bookings, e-Commerce.

1. Introducción

La planificación de viajes está dominada actualmente por los recursos en línea: agencias de viajes, sitios web de opiniones de viaje, sitios web de opiniones de hote-

les, redes sociales, sitios web de hoteles, etc. Por ejemplo, el 93% de los viajeros a nivel mundial dice que sus decisiones de reserva se ven influenciadas por las opiniones en línea [1]. Cuando una persona ha decidido el lugar a visitar, el siguiente paso es buscar dónde hospedarse y dado que el primer medio de búsqueda es el Internet, no hay mejor oportunidad para los hoteles que promocionar sus servicios por Internet. En la actualidad existe una amplia variedad de agencias de viajes, buscadores de hoteles y otras opciones, que permiten localizar hoteles en todo el mundo; sin embargo, la información del hotel que se muestra en estos sitios web es muy escasa y poco útil.

Por otro lado, el 72% de los viajeros a nivel mundial afirman que la posibilidad de efectuar reservaciones a través de un dispositivo móvil es muy útil, y sólo el 25% de los alojamientos permite esta funcionalidad [1]. Es por ello que incluir un motor de reservaciones en el sitio web de un hotel es la mejor opción para los hoteles que desean vender de manera fácil y directa, ofreciendo comodidad a sus clientes y abarcando un amplio mercado de ventas: el Internet.

Los sitios web de hoteles tienen como objetivo principal ofrecer el servicio de alojamiento, además de otros servicios generales. Estos sitios se dividen en secciones, donde se describen los diferentes servicios con los que cuenta el hotel, mediante galerías fotográficas, así como secciones de contacto y solicitudes de reservación.

En este artículo se presenta el diseño e implementación de un sistema web personalizable con diseño adaptable para facilitar a administradores de hoteles pequeños el manejo de reservaciones de sus habitaciones, gestionando las ventas desde su propio sitio web; así como para facilitar y garantizar de inmediato las reservaciones para sus huéspedes potenciales. El sistema web que se presenta utiliza el patrón de diseño web adaptable, por lo que es posible visualizarlo automáticamente en diferentes dispositivos, tales como computadoras de escritorio, laptops, tabletas y teléfonos móviles.

El resto del artículo está organizado de la siguiente manera. La sección 2 analiza tres herramientas existentes para personalización de sitios web para hoteles. El análisis y diseño del sistema se presenta en la sección 3, donde se analizan algunos sitios web para hoteles, se realiza un estudio de usuarios, se describe la metodología de desarrollo utilizada para la creación del sistema y se presentan los módulos principales que lo componen. La sección 4 proporciona detalles de la implementación del sistema, con un mapa de navegación, una descripción de las tecnologías y herramientas utilizadas y la implementación del patrón de diseño web adaptable. Las pruebas del sistema y sus resultados se presentan en la sección 5. Finalmente, en la sección 6 se presentan las conclusiones y el trabajo futuro.

2. Herramientas existentes

En esta sección se describen tres herramientas que permiten personalizar sitios web para hoteles con motor de reservaciones en línea: *PRO Internacional* [2], *Creowebs* [3] y *Obehotel* [4]. Al final de la sección se presenta una tabla comparativa con características de las herramientas analizadas, la cual se explica detalladamente en [5].

PRO Internacional [2]. De acuerdo a su sitio web, *PRO Internacional* es una empresa dedicada a brindar servicios web integrales, que van desde la creación de sitios

web con alto grado de funcionalidad e interactividad, hasta acciones en redes sociales y campañas de posicionamiento en buscadores. Entre los sitios web con alto grado de funcionalidad que crea *PRO Internacional*, se encuentran los sitios web para hoteles con módulo integrado de reservaciones de habitaciones en línea. Estos sitios web cuentan con las siguientes funcionalidades: reserva de habitaciones en línea, tarifarios por temporadas, galería de fotos, suscripción a boletín electrónico, comentarios, encuestas, buzón de sugerencias. Cabe señalar que *PRO Internacional* no es una herramienta gratuita, y no cuenta con demostración.

Creoweb [3]. Es una herramienta web que permite crear sitios web a partir de una variedad de plantillas predeterminadas, las cuales están diseñadas para diferentes tipos de sitios web, tales como blogs, sistemas de reservaciones para hoteles o restaurantes, etc. *Creoweb* proporciona los siguientes módulos para creación de sitios web: plantillas, importación de contenido, editor de imágenes, blog y noticias, redes sociales, videos, música, galerías de fotos, formulario de contacto, compartir archivos, pagos por *PayPal*, uso de *Google AdSense*, reservaciones, uso de *Google Maps*, encuestas, visualización de ruta de navegación y versión para dispositivos móviles.

ObeHotel [4]. De acuerdo a sus creadores, *ObeHotel* ofrece cuatro diferentes servicios: pasarela de reservaciones; *web completa*; *channel manager*; y *máxima visibilidad*. El servicio de *web completa* de *ObeHotel* es el que está relacionado con el sistema web presentado en este artículo, por lo cual fue el que se analizó. Las principales características de *web completan* son las siguientes: diseño personalizado, creación y configuración de cualquier tipo de habitación, reservación de múltiples habitaciones, gestión de ofertas y paquetes, diferentes tarifas para el sitio web y para los canales de venta, código de rastreo para monitorear visitas, integración con *Google Analytics* y *Google Adwords* y módulo web de opiniones. Debido a que el sistema no es gratuito y no cuenta con una demostración para la parte de administración del sistema, no se pudieron probar los módulos correspondientes al administrador.

En la Tabla 1 se muestra una comparación de funcionalidades entre las herramientas analizadas: *PRO Internacional*, *Creoweb* y *ObeHotel*. El símbolo de verificación indica que la herramienta tiene la funcionalidad, mientras que la *x* que no la tiene. Se proporciona también una descripción de las funcionalidades mostradas en la Tabla 1.

Gestión de habitaciones. Funcionalidad para permitir al administrador la creación, modificación y eliminación de habitaciones.

Reservación múltiple. Funcionalidad para permitir la reservación múltiple en una sola transacción por el mismo huésped potencial.

Tarifas por temporadas. Es la funcionalidad para registrar tarifas por habitación según la temporada del hotel; dichas temporadas son definidas por el administrador con fechas de inicio y fin sin posibilidad de traslapes entre ellas.

Gestión de ofertas y paquetes. Funcionalidad que permite la creación, modificación y eliminación de ofertas y paquetes. Las ofertas son descuentos o promociones en el precio de los servicios; y los paquetes son un grupo de servicios en conjunto.

Mapa de disponibilidad. Esta funcionalidad permite visualizar el mapa de disponibilidad con la relación de habitaciones disponibles y ocupadas. En algunos casos se muestra gráficamente y en otros es una tabla del total de las habitaciones y su estado.

Comentarios. Funcionalidad para incluir una sección que permita a los huéspedes publicar comentarios sobre su estancia en el hotel.

Encuestas. Funcionalidad que permite al administrador del hotel realizar encuestas de satisfacción de la estancia de los huéspedes.

Buzón de sugerencias. Esta funcionalidad permite incluir un buzón de sugerencias con un formulario, donde el huésped puede emitir una sugerencia al hotel y ésta se envía al correo electrónico del administrador.

Plantillas prediseñadas. Funcionalidad que permite al administrador del hotel elegir una plantilla prediseñada para su sitio web.

Formularios personalizables. Es la funcionalidad que permite al administrador del hotel personalizar formularios; añadiendo y eliminando campos.

Estadísticas. Funcionalidad que permite al administrador conocer información relevante sobre las visitas al sitio web.

Versión móvil. Esta funcionalidad permite visualizar el sitio web desde un dispositivo móvil, ajustando el diseño del sitio web según la resolución del dispositivo.

Tabla 1. Características para evaluar las herramientas analizadas.

<i>Características</i>	<i>PRO Internacional</i>	<i>Creowebs</i>	<i>ObeHotel</i>
Gestión de habitaciones	✓	✓	✓
Reservación múltiple	✗	✗	✓
Tarifas por temporadas	✓	✓	✗
Gestión de ofertas y paquetes	✗	✗	✓
Mapa de disponibilidad	✓	✓	✓
Comentarios	✓	✓	✓
Encuestas	✓	✓	✗
Buzón de sugerencias	✓	✗	✗
Plantillas prediseñadas	✓	✓	✗
Formularios personalizables	✗	✓	✗
Estadísticas	✗	✗	✓
Versión móvil	✗	✓	✗

3. Análisis y diseño

Esta sección presenta el análisis y diseño llevado a cabo previo a la implementación del sistema web. Se describe la estructura e información proporcionada en sitios web para hoteles; se presenta un estudio de usuarios, donde tres perfiles fueron identificados para el sistema; se muestra la metodología de desarrollo que se utilizó para la creación del sistema; y finalmente, se proporcionan detalles de los módulos principales que componen al sistema web.

3.1. Estructura de sitios web de hoteles

Con el objetivo de conocer la estructura, tipo de información y funcionalidades con las que cuentan los sitios web de hoteles, se visitaron seis sitios web de hoteles ubicados en lugares turísticos de México: 1) Las cúpulas en Malinalco, Estado de México [6]; 2) La casa del laurel en Taxco, Guerrero [7]; 3) El encanto en Bernal, Querétaro [8]; 4) Hotel Malinalco en Malinalco, Estado de México [9]; 5) Casa abierta en Valle de Bravo, Estado de México [10]; y 6) Hotel Cristal en Chignahuapan, Puebla [11]. Un análisis comparativo de la información y características que se ofrecen en cada uno de los sitios web de los hoteles visitados se presenta en [5].

3.2. Estudio de usuarios

Los perfiles de usuario describen las características de los usuarios del sistema y permiten diseñar un sistema centrado en las necesidades de los usuarios. El estudio de usuarios que se llevó a cabo consistió en identificar a los usuarios del sistema web, mediante el diseño de una encuesta para cada uno, con preguntas que permitieron obtener información sobre sus intereses, necesidades, y características, las cuales se tomaron también en cuenta para el diseño del sistema web resultante.

Se identificaron dos tipos de usuarios: administrador y huésped potencial. Además, se definió un perfil de hotel, debido a que el sistema web se planteó para hoteles con características específicas. Se realizaron un total de 16 encuestas a administradores de hoteles ubicados en los pueblos mágicos de Huasca de Ocampo y Real del Monte, ambos pertenecientes al estado de Hidalgo.

También se realizaron un total de 84 encuestas a turistas que se encontraron en estos pueblos mágicos. Una vez realizadas las encuestas, se analizaron los resultados obtenidos y se procedió a definir los perfiles de usuario, los cuales se utilizaron para el diseño del sistema.

3.3. Metodología de desarrollo

Para la realización del sistema web se utilizó una metodología iterativa e incremental, la cual consiste en realizar varias iteraciones, con las siguientes etapas en cada una de ellas: análisis, diseño, implementación y pruebas. Esta metodología es incremental ya que en cada iteración se agregan nuevas funcionalidades al sistema. Con esta metodología fue posible obtener de manera inicial un prototipo funcional desde la primera iteración, y posteriormente fue posible evaluar con el cliente dicho prototipo para obtener retroalimentación y resolver problemas de manera temprana.

A través del desarrollo iterativo e incremental se logra una versión más estable del sistema al final de cada iteración, la cual tiene mayor calidad y posee nuevas funcionalidades con respecto a versiones anteriores. Para este proyecto se realizaron cuatro iteraciones, agregando nuevos módulos al sistema hasta completar la funcionalidad total del mismo; se hicieron adecuaciones de acuerdo a la retroalimentación recibida del cliente, y se corrigieron los errores detectados en las pruebas de las iteraciones anteriores.

3.4. Módulos del sistema

Los requerimientos generales para el sistema web fueron obtenidos del hotel *Bosque Mágico*, el cual fue uno de los hoteles con los que se tuvo contacto en la etapa de estudio de usuarios. Se eligió este hotel debido a la variedad de opciones y servicios que ofrece, lo cual permitió tener un sistema más general y completo. A partir de los requerimientos y del análisis previo realizado, se determinó que el sistema web estaría compuesto de cuatro módulos principalmente.

Módulo de Edición. Permite al administrador personalizar el portal web con las características de su hotel. El administrador proporciona información al sistema acerca de habitaciones, paquetes, servicios, reservaciones, entre otros.

Módulo de Administración. Permite al administrador manejar las reservaciones del hotel, ya que se tiene acceso a la información de cada una. El administrador es capaz de visualizar a través de un calendario, la disponibilidad del hotel.

Motor de Reservaciones. Este módulo permite al huésped potencial y al administrador verificar la disponibilidad del hotel, proporcionando fecha de ingreso, fecha de salida y número de personas. También es posible reservar una o más de las habitaciones mostradas en la lista de disponibilidad.

Portal Web. Permite al huésped potencial navegar a través del portal que fue personalizado por el administrador a través del módulo de edición. En este portal web el huésped potencial puede realizar reservaciones y encontrar información acerca de paquetes, habitaciones, servicios, atracciones, ubicación del hotel, entre otros.

Para cada uno de estos módulos se llevaron a cabo los correspondientes casos de uso, modelo de interfaz, diagramas de clases y diagramas de secuencia, los cuales fueron utilizados en cada iteración con el administrador del hotel, para obtener retroalimentación y llegar a acuerdos en el desarrollo del sistema.

4. Implementación

La implementación del sistema web se presenta en esta sección. Se muestra el mapa de navegación del sistema con los dos usuarios del sistema: administrador y huésped potencial; se proporcionan las tecnologías y herramientas utilizadas para el desarrollo del sistema; y finalmente, se describe la implementación del patrón de diseño web adaptable, el cual permitió adaptar la interfaz del sistema web a diferentes dispositivos para mejorar su visualización y navegación.

4.1. Mapa de navegación

La Figura 1 muestra el mapa de navegación del sistema, con los dos usuarios del sistema: el administrador del hotel y el huésped potencial. El mapa de navegación muestra el conjunto de páginas web a las cuales tiene acceso cada usuario.

El administrador puede acceder dos secciones principalmente: la edición del portal y la administración. En la edición del portal el administrador puede personalizar el portal web mediante la creación, modificación y eliminación de habitaciones, temporadas, servicios, paquetes, promociones, información de contacto, galerías fotográficas.

cas y atracciones. En la sección de administración, el administrador puede manejar todas las reservaciones, realizar reservaciones manualmente en el sistema, modificar reservaciones existentes, visualizar reservaciones en el mapa de disponibilidad de habitaciones y desplegar estadísticas del sistema.

El huésped potencial puede acceder a todas las páginas web personalizadas por el administrador: servicios, atracciones, galerías fotográficas, información de contacto, promociones, habitaciones, paquetes y reservaciones de habitaciones y paquetes. Para realizar una reservación de habitación o paquete, el huésped potencial puede verificar la disponibilidad de habitaciones en la página web de reservaciones.

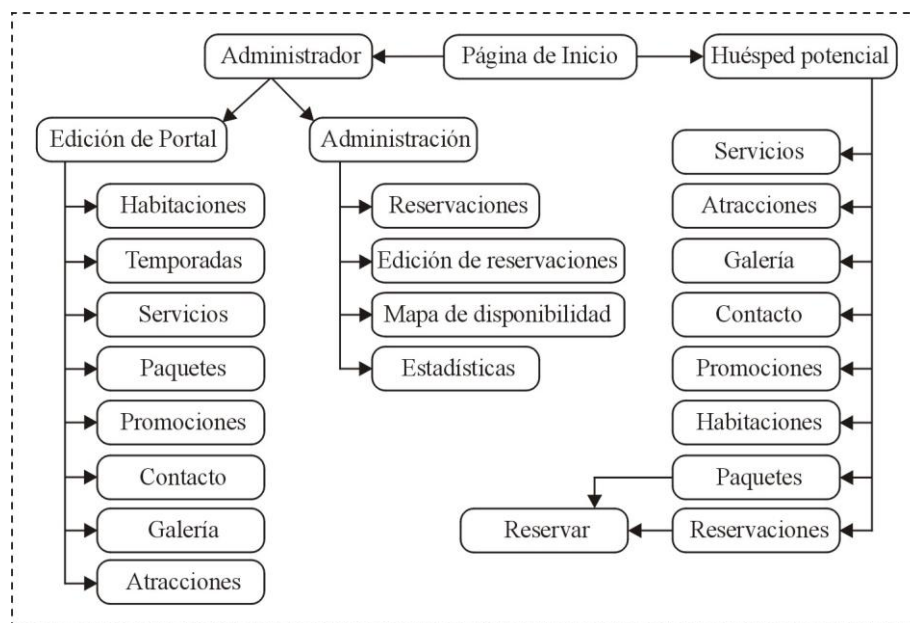


Fig. 1. Mapa de navegación del sistema.

4.2. Tecnologías y herramientas utilizadas

Para el desarrollo del sistema web se utilizaron una serie de tecnologías y herramientas, algunas de las cuales se mencionan a continuación.

Java Server Pages (JSP). Es una tecnología Java que permite el desarrollo de páginas web dinámicas, combinando HTML estático con HTML generado dinámicamente. Los archivos JSP fueron utilizados para la creación de páginas web dinámicas, donde la información se extrae de una base de datos y se procesa.

JavaScript. Es un lenguaje de programación interpretado, diseñado para hacer páginas web interactivas y dinámicas. JavaScript fue principalmente utilizado para la creación dinámica de páginas web, especialmente para la estructura dinámica de la información que se despliega en los diferentes dispositivos.

Cascading Style Sheets (CSS). Es un lenguaje de estilos utilizado para presentar información en una página web. El mecanismo para cambiar la posición de elementos HTML en una página web es ampliamente utilizado para organizar la información cuando se accede al sistema web desde diferentes dispositivos, debido a que el espacio de visualización disponible en cada dispositivo puede ser diferente.

NetBeans. Es un ambiente de desarrollo integrado, el cual permite escribir, compilar, depurar y ejecutar programas. NetBeans permitió concentrar todos los archivos del proyecto, tales como clases Java, archivos JSP y JavaScript, entre otros.

Apache Tomcat. Es un servidor web y contenedor de *servlets*, el cual permite compilar archivos JSP que son convertidos en *servlets*. Apache Tomcat fue utilizado para el procesamiento de los archivos JSP del sistema web.

MySQL. Es un sistema manejador de bases de datos relacionales, basado en el lenguaje de consulta estructurado (*Structured Query Language*, SQL por sus siglas en inglés). MySQL fue utilizado para la creación de la base de datos del sistema web, así como para ejecutar sentencias tales como *SELECT*, *INSERT*, *UPDATE*, entre otras.

4.3. Diseño web adaptable

El diseño web adaptable (*Responsive Web Design*, RWD por sus siglas en inglés) [12, 13] es un patrón de diseño que permite adaptar la interfaz de un portal web a diferentes dispositivos con el objetivo de mejorar la visualización y la experiencia de navegación de los usuarios. Actualmente, hay una gran variedad de dispositivos móviles con diferentes resoluciones y tamaños de pantalla, lo cual es la principal razón por la que el diseño de interfaces web requiere adaptarse a esos dispositivos. RWD está basado en *Media Queries*, los cuales son una serie de instrucciones en CSS que definen ciertas características de anchura, altura, orientación, color y otros parámetros del dispositivo desde el cual el portal web es accedido. Esto permite aplicar diferentes estilos para cada dispositivo mediante la determinación de su tamaño de pantalla o resolución, de tal manera que sea posible tener diferentes vistas del mismo portal web.

Debido a que actualmente hay diferentes dispositivos desde los que un sitio web puede ser accedido, el sistema web presentado en este artículo define tres vistas para ser visualizado: una vista para teléfonos móviles, una vista para tabletas y una vista para computadoras de escritorio y laptops. La forma más frecuente de distribuir la información es a través de un número específico de columnas de acuerdo a la vista; por ejemplo, una columna para la vista de teléfonos móviles, dos columnas para la vista de tabletas y tres columnas para la vista de computadoras de escritorio y laptops. La Figura 2 muestra un ejemplo de cómo puede ser arreglada la distribución de información en columnas de acuerdo al dispositivo.

Cabe señalar que en algunos casos los *Media Queries* no son lo suficientemente flexibles para organizar información. Este es el caso del sistema web presentado en este artículo, en el cual cierta información es creada dinámicamente y la cantidad de elementos a ser visualizados en el portal web es variable. Por esta razón, en este sistema fue necesario adicionalmente la implementación de varias funciones JavaScript, las cuales permitieron la organización de la información de acuerdo al dispositivo desde el cual el portal web es accedido. Las funciones JavaScript creadas fueron con-

troladas utilizando los atributos *window.innerWidth* y *window.innerHeight*, los cuales permiten conocer el tamaño de la ventana del navegador web, de tal forma que se puede determinar el número de columnas que será utilizado para dividir la información y ser visualizada. Las funciones JavaScript creadas son ejecutadas cuando se disparan los eventos *onresize* y *onload* del navegador web.

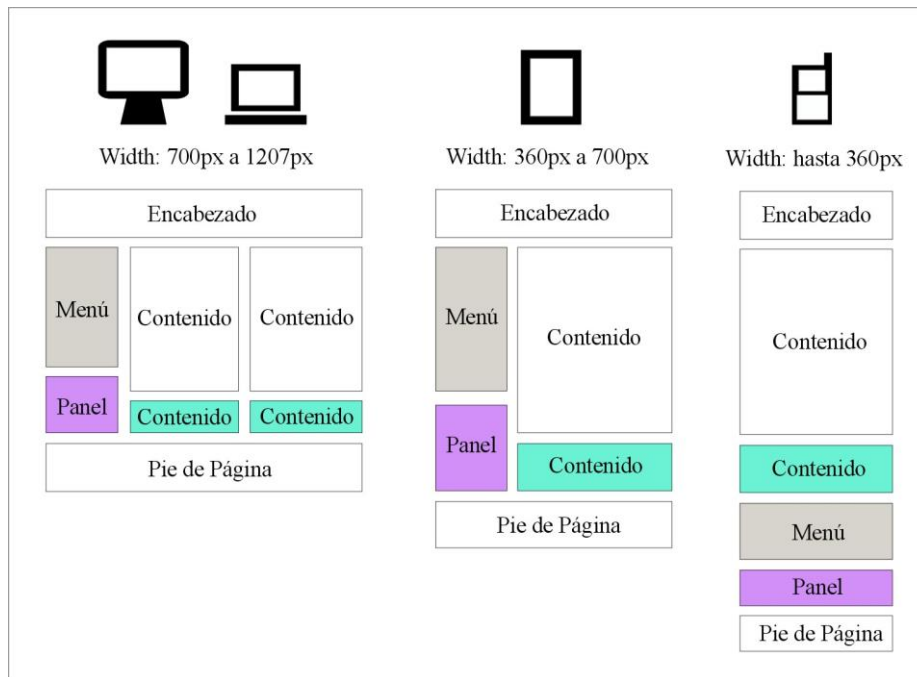


Fig. 2. Distribución de información en columnas de acuerdo al dispositivo.

5. Pruebas y resultados

Esta sección presenta las pruebas llevadas a cabo en el sistema y sus resultados. Se realizaron dos tipos de pruebas: de funcionalidad y de interfaz.

Las pruebas de funcionalidad fueron realizadas a través de controles de prueba para determinar el éxito o fracaso del sistema en tareas específicas, las cuales fueron parte del levantamiento de requerimientos con el usuario y plasmadas como casos de uso que se realizaron en la etapa de análisis y diseño de la metodología de desarrollo. Los resultados de estas pruebas fueron los siguientes: 92% de los controles de prueba fueron exitosos, mientras que el 8% de ellos falló. Cabe señalar que las pruebas fallidas estuvieron relacionadas con tareas donde las fechas de reservaciones se traslapaban y con precios que estaban equivocados de acuerdo a las temporadas.

Las pruebas de interfaz se llevaron a cabo para verificar la correcta visualización del portal web en los diferentes dispositivos para los cuales fue diseñado: teléfono móvil, tableta, laptop y computadora de escritorio. La Tabla 2 muestra los dispositi-

vos en los que el portal web fue probado. Las columnas de la tabla indican la siguiente información acerca del dispositivo: tipo, marca, modelo, sistema operativo, navegador web, resolución de la pantalla en pixeles. También se muestran las capturas de pantalla de la visualización del sistema web con las habitaciones disponibles en un hotel en tres diferentes dispositivos de los que se utilizaron en las pruebas (1, 3 y 5).

Tabla 2. Dispositivos utilizados en las pruebas de interfaz.

<i>Tipo</i>	<i>#</i>	<i>Marca</i>	<i>Modelo</i>	<i>S.O.</i>	<i>Navegador</i>	<i>Resolución</i>
Laptop	1	Dell	XPS - M1330	Windows	Chrome	1280 x 800
Laptop	2	HP	Pavilion 15p213ns	Windows	Chrome	1366 x 768
Tableta	3	Apple	iPad mini 2	iOS 8	Safari	768 x 1024
Tableta	4	Samsung	Galaxy Tab 4	Android	Chrome	800 x 1280
Teléfono	5	Alcatel	One touch	Android	Android	320 x 568
Teléfono	6	Samsung	Galaxy S4	Android	Android	360 x 640

La Figura 3 muestra la interfaz web desplegada en una laptop con una resolución de 1280 por 800 pixeles. La Figura 4 muestra del lado izquierdo la interfaz web desplegada en una tableta con una resolución de 768 por 1024 pixeles, y en el lado derecho muestra la misma interfaz web desplegada en un teléfono móvil con una resolución de 320 por 568 pixeles.



Fig. 3. Sistema web desplegado en una laptop (1280x800).

La página web de la Figura 3 está dividida en dos paneles: el panel superior, el cual tiene el menú con diferentes opciones para el huésped potencial; y el panel inferior, el cual es utilizado para desplegar el contenido de la página y contiene las habitaciones disponibles organizadas en tres columnas.

La captura de pantalla de la Figura 4 (izquierda) muestra la misma página web desplegada en una tableta, con información reducida y menú simplificado. La información de las habitaciones disponibles es organizada en dos columnas, el menú está oculto en la parte superior y se visualiza del lado izquierdo cuando el usuario da click sobre el icono ubicado en la esquina superior derecha. La Figura 4 (derecha) es la misma página web, pero desplegada en un teléfono móvil, en la que las habitaciones están organizadas en una sola columna y el menú está también oculto.



Fig. 4. Sistema web desplegado en una tableta (768x1024) y en un teléfono móvil (320x568).

6. Conclusiones y trabajo futuro

En este artículo se presentó el diseño e implementación de un sistema web personalizable con diseño adaptable, el cual cumplió con los requisitos planteados al inicio de su desarrollo, así como con las necesidades de los clientes, las cuales se tomaron a partir del levantamiento de requerimientos y de las encuestas realizadas. El sistema web facilita a administradores de hoteles pequeños el manejo de reservaciones de sus habitaciones, gestionando las ventas desde su propio sitio web; así como facilita y garantiza de inmediato las reservaciones para sus huéspedes potenciales.

Se logró exitosamente que el sistema web se adaptara automáticamente a diferentes dispositivos, tales como computadoras de escritorio, laptops, tabletas y teléfonos móviles, utilizando el patrón de diseño web adaptable, mediante el uso de *CSS*, *Media Queries* y *JavaScript*. Esta característica de adaptación se verificó a través de las pruebas de interfaz que se realizaron sobre diferentes dispositivos. En particular, los resultados se mostraron mediante capturas de pantalla de las interfaces desplegadas en tres dispositivos: en una laptop con una resolución de 1280 por 800 píxeles, en una tableta con una resolución de 768 por 1024 píxeles, y en un teléfono móvil con una resolución de 320 por 568 píxeles.

Es necesario llevar a cabo pruebas de usabilidad con administradores de hoteles, así como con huéspedes potenciales. El sistema será ajustado después de las pruebas realizadas y se pondrá en funcionamiento en los hoteles que cumplan con el perfil de hotel definido en el estudio de usuarios.

Referencias

1. Trip Barometer: Encuesta realizada a hoteleros y viajeros. Disponible en: http://www.tecnohotelnews.com/wp-content/uploads/2013/03/Infografia_TripBarometer.jpg
2. Sitio web de PRO Internacional. Disponible en: <http://www.prointernacional.com/>
3. Sitio web Creoweb. Disponible en: <https://creoweb.com/>
4. Sitio web ObeHotel. Disponible en: <http://www.obehotel.com/>
5. García-Mendoza, B., Jaimez-González, C. R.: Propuesta de Sistema Web Personalizable para el Control de Reservaciones en Hoteles. *Research in Computing Science: Advances in Intelligent Information Technologies*, 79, 135–145 (2014)
6. Sitio web del hotel *Las cúpulas*. Disponible en: <http://lascupulas.com.mx/>
7. Sitio web del hotel *La casa del laurel*. Disponible en: <http://www.hotelentaxco.com/es/index.html>
8. Sitio web del hotel *El encanto*. Disponible en: <http://elencantobernal.com/acerca>
9. Sitio web del hotel *Hotel Malinalco*. Disponible en: <http://www.hotelmalinalco.com.mx/>
10. Sitio web del hotel *Casa abierta*. Disponible en: <http://www.casabierta.com.mx/index.html>
11. Sitio web del hotel *Hotel Cristal*. Disponible en: <http://www.hotelcristalchignahuapan.com.mx/>
12. Tendencias en diseño web. Disponible en: <http://www.tecnohotelnews.com/2013/07/las-tendencias-de-futuro-en-diseno-web-para-hoteles/>
13. Responsive Web Design. Disponible en: <http://www.designtribe.ie/news/24/61/New-Responsive-B-B-Hotel-website>

Serialización de objetos PHP a XML

Lidia N. Hernández-Piña, Carlos R. Jaimez-González

Universidad Autónoma Metropolitana, Unidad Cuajimalpa,
Departamento de Tecnologías de la Información, México

210368177@alumnos.cua.uam.mx, cjaimez@correo.cua.uam.mx

Resumen. La serialización es el proceso de escribir un objeto en un medio de almacenamiento como un archivo o un buffer de memoria, con el objetivo de transmitirlo a través de una red; la deserialización es el proceso inverso. En este artículo se presenta un serializador y deserializador de objetos PHP a XML, el cual es una biblioteca independiente escrita en el lenguaje de programación PHP. Una de sus principales características es la generación de representaciones XML estándar, las cuales son independientes del lenguaje de programación e interoperables con los serializadores *Web Objects in XML* (WOX) existentes.

Palabras clave: Serialización de objetos PHP, web objects in XML, PHP a XML, interoperabilidad, WOX.

Serialization of PHP Objects to XML

Abstract. Serialization is the process of writing an object to a storage medium as a file or as a buffer in memory, with the aim of transmitting it through a network; de-serialization is the inverse process. This paper presents a serializer and de-serializer of PHP objects to XML, which is an independent library written in the PHP programming language. One of its main features is the generation of standard XML representations, which are independent of the programming language and interoperable with the existing *Web Objects in XML* (WOX) serializers.

Keywords. Serialization of PHP objects, web objects in XML, PHP to XML, interoperability, WOX.

1. Introducción

La interoperabilidad es una característica importante en los sistemas distribuidos basados en objetos, ya que permite la comunicación de programas (clientes y servidores), escritos en diferentes lenguajes de programación orientados a objetos. Existen algunos problemas fundamentales que tienen que ser resueltos por los diferentes len-

guajes para alcanzar la interoperabilidad. Algunos de estos problemas se exponen a continuación.

Mapeo de tipos de datos. Los tipos de datos son uno de los principales problemas para alcanzar la interoperabilidad entre diferentes lenguajes de programación. Debe de existir un acuerdo de mapeo entre los tipos de datos de los lenguajes de programación que se desean comunicar. Una forma de resolver este problema es mediante tablas de mapeo con los diferentes tipos de datos soportados por los lenguajes de programación que desean comunicarse.

Representación de objetos. Se debe establecer una forma estándar de representar objetos, ya sea que estén escritos en Java, C#, C++, o algún otro lenguaje de programación orientado a objetos. El formato estándar debe considerar la representación de las estructuras y tipos de los diferentes lenguajes de programación, tales como clases, tipos de datos primitivos, arreglos, y clases definidas por el usuario.

Mensajes. Los mensajes que se intercambian también deben de estar escritos en una forma estándar para que puedan ser entendidos por ambas partes.

Serialización y deserialización. En el contexto de almacenamiento y transmisión de datos, la serialización es el proceso de transformar un objeto a un estado en el que pueda ser almacenado permanentemente a un medio tal como un archivo, una base de datos, o un flujo para ser transmitido a través de la red. Deserialización es el proceso inverso, en el cual se transforma la versión serializada en XML del objeto en un objeto en memoria.

Este artículo presenta un serializador y deserializador de objetos PHP a XML, llamado *PHP Web Objects in XML* (PHPWOX) [1]. El XML generado por PHPWOX es independiente del lenguaje de programación, y puede ser utilizado por otros serializadores y deserializadores *Web Objects in XML* (WOX) [2], los cuales permiten interoperabilidad entre aplicaciones escritas en PHP y aplicaciones escritas en los lenguajes de programación soportados por WOX: Java, C# [3] y Python [4].

El resto del artículo está organizado de la siguiente manera. La sección 2 proporciona algunos conceptos importantes para entender los procesos de serialización y deserialización. En la sección 3 se presentan algunos parsers y serializadores PHP que fueron analizados. Una tabla comparativa de parsers y serializadores se presenta en la sección 4. La sección 5 describe las clases que son parte de PHPWOX, se presenta su arquitectura y una descripción de la funcionalidad de los módulos que lo componen. Finalmente, en la sección 6 se presentan las conclusiones y el trabajo a futuro.

2. Conceptos

La serialización es un proceso que consiste en la codificación de un objeto en un medio de almacenamiento como puede ser un archivo o un buffer de memoria, con el objetivo de transmitirlo a través de una conexión en red como una serie de bytes o en otro formato como XML. La serie de bytes o el formato pueden ser usados para crear un nuevo objeto que es idéntico en todo al original, incluido su estado interno.

Al programa capaz de serializar un objeto directamente en un archivo de forma automática se le denomina *serializador*; mientras que un *parser* es un programa que

puede leer y escribir XML en un archivo, en éste la serialización es de forma personalizada, es decir, la realiza el usuario del *parser*.

Para serializar objetos en XML básicamente se sigue el proceso que se observa gráficamente en la Figura 1, el cual se describe a continuación.

1. Se obtiene el nombre, tipo y valor de cada uno de los atributos del objeto. Esto se realiza mediante el proceso de reflexión, el cual es la capacidad que tiene un programa para observar y opcionalmente modificar su estructura de alto nivel. Por medio de esta capacidad es posible acceder a la información de los objetos, conociendo y/o ejecutando sus atributos y métodos públicos, todo ello en tiempo de ejecución; también se utiliza la introspección, la cual es la obtención del tipo de dato en particular.
2. Una vez que se obtiene el nombre y valor de cada atributo del objeto a serializar, se escriben en un archivo XML. En caso de que el valor no sea tipo primitivo sino un objeto, habrá que representar también en el documento XML todos los atributos de dicho objeto.

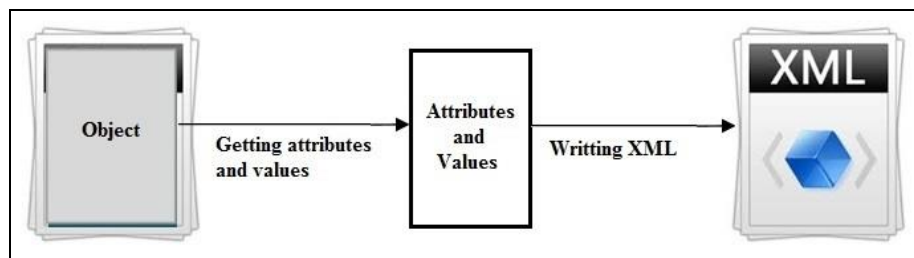


Fig. 1. Serialización de un objeto a XML.

El proceso inverso es la deserialización, en la cual se siguen los siguientes pasos y se puede observar en la Figura 2.

1. Se extrae la información acerca de un objeto del archivo XML.
2. Un objeto es una instancia de una clase, por lo cual para crear el objeto es necesario primero crear una clase basándose en la información obtenida del archivo XML, por lo que se crea la clase.
3. Se instancia el objeto con la información obtenida del documento XML.

3. Parsers y serializadores existentes

En esta sección se presenta el análisis y comparación de algunos *parsers* y serializadores PHP existentes. Los *parsers* que se analizaron fueron DOM y *SimpleXML*; los serializadores que no utilizan XML que se revisaron fueron la función *serialize()* de PHP, la función *json_encode()* de PHP y la biblioteca *Igbinary*; y el único serializador XML que se analizó fue *Pear*. Aunque existen varios *parsers* y serializadores de objetos PHP incorporados como funciones PHP o como complementos externos, ninguno de ellos permite interoperar con otros lenguajes de programación. En los

siguientes párrafos se proporciona una breve descripción de cada *parser* o serializador.

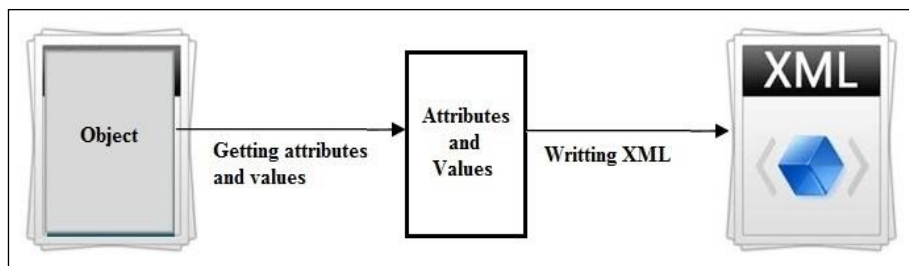


Fig. 2. Deserialización de un objeto a partir de un documento XML.

Document Object Model (DOM) [5]. El modelo de objetos de documento es una interfaz de programación de aplicaciones (*Application Programming Interface*, API por sus siglas en inglés) que proporciona un conjunto estándar de objetos para representar documentos HTML y XML, es un modelo para describir cómo combinar tales objetos, y es una interfaz estándar para accederlos y manipularlos. A través de DOM es posible acceder y modificar el contenido, estructura y estilo de documentos HTML y XML. PHP tiene una extensión de DOM incluida en la distribución de XAMPP [6], la cual permite crear y manipular documentos XML. Cabe señalar que con DOM el programador debe especificar cómo debe serializarse cada objeto, lo cual significa que no hay forma para serializar automáticamente los objetos de cualquier clase; este proceso es llevado a cabo manualmente por el programador.

SimpleXML [7]. Es una biblioteca de PHP incluida en la distribución de XAMPP, la cual permite manipular archivos XML. Comparado con *DOM*, *SimpleXML* requiere la escritura de menos código para leer elementos de un documento XML. Al igual que con *DOM*, una desventaja de *SimpleXML* es que el programador debe especificar cómo debe serializarse cada objeto, este proceso es llevado a cabo manualmente por el programador.

Función *serialize()* [8]. Es una función de PHP que regresa una cadena que contiene un flujo de *bytes* que representan cualquier valor que puede ser almacenado en PHP, lo cual significa que la serialización no es una representación en XML. La función *unserialize()* puede restaurar los valores originales a partir de la cadena mencionada. Utilizando la función *serialize()* para almacenar un objeto se guardarán todos los valores de las variables del objeto. Para este tipo de serialización es necesario tener la clase a partir de la cual el objeto fue creado.

Función *json_encode()* [9]. Esta función regresa un valor serializado en una cadena. PHP implementa un conjunto más amplio de JSON y también codifica y decodifica valores escalares y valores nulos. El estándar de JSON solamente acepta estos valores cuando están anidados en un arreglo o en un objeto. La función *json_encode()* puede serializar objetos directamente solo si sus atributos son públicos.

IgBinary [10]. Es un serializador que sustituye al serializador predeterminado de PHP. Este serializador no utiliza una representación de texto, sino una representación

binaria. Una ventaja de este serializador es que utiliza muy poca memoria para la representación de un objeto; sin embargo, una desventaja es que la representación del objeto no es visible al usuario.

Pear XML_Serialize [11]. Es un *framework* de componentes PHP reutilizables. *XML Serialize* es parte de este *framework* y permite serializar datos complejos a XML, tales como arreglos y objetos. *Pear* permite serialización directa de un objeto a XML, lo cual significa que el usuario no necesita especificar cómo se debe realizar la serialización del objeto.

4. Comparación

La Tabla 1 muestra una comparación de las características que se tomaron en cuenta en los *parsers* y serializadores analizados: H1) *DOM*, H2) *SimpleXML*, H3) Función *serialize()*, H4) Función *json_encode()*, H5) *IgBinary*, H6) *Pear XML_Serializer*. El símbolo de verificación indica que el *parser* o serializador tiene la característica que se indica, mientras que la *x* indica que no la tiene. Una breve descripción de las características consideradas para evaluar los *parsers* y serializadores se presenta en los siguientes párrafos.

Tabla 1. Comparación de características de los parsers y serializadores analizados.

Características	H1	H2	H3	H4	H5	H6
Interoperabilidad	x	x	x	✓	x	x
Convertir objeto a XML	x	x	x	x	x	x
Convertir XML a objeto	x	x	x	x	x	x
Licencia gratuita	✓	✓	✓	✓	✓	✓
XML bien formado	✓	✓	x	x	x	✓
Manual usuario/documentación	✓	✓	✓	✓	✓	✓
Ejemplos de uso	✓	✓	✓	✓	✓	✓

Interoperabilidad. Es la posibilidad de serializar un objeto en un lenguaje de programación y deserializarlo en un lenguaje distinto, y viceversa. Por ejemplo, serializar un objeto en el lenguaje PHP y deserializarlo en el lenguaje Java.

Convertir directamente objeto a XML. Es la posibilidad de serializar un objeto a XML automáticamente por medio de un método o procedimiento existente, sin necesidad de que el programador lo tenga que hacer manualmente.

Convertir directamente XML a objeto. Es la posibilidad de obtener uno o más objetos automáticamente a partir de un método o procedimiento existente y un archivo XML, sin necesidad de que el programador lo tenga que hacer manualmente.

Licencia gratuita. Se refiere a que el serializador se distribuye sin costo, disponible para su uso y por tiempo ilimitado.

XML bien formado. Se refiere a un documento XML en el cual las etiquetas están correctamente anidadas, además de tener sus correspondientes etiquetas iniciales y finales.

Manual de usuario y documentación. Se refiere a la existencia de un manual o guía de uso del *parser* o serializador.

Ejemplos de uso. Es la presencia de ejemplos de uso dentro de la documentación y/o manual del serializador o *parser* a analizar.

5. Implementación del serializador

Esta sección proporciona una explicación de las clases que son parte del serializador y del deserializador que se desarrolló, también se presenta su arquitectura y una descripción de la funcionalidad de los módulos principales que la componen.

5.1. Clases

La Figura 3 muestra un diagrama con las clases que fueron desarrolladas como parte del serializador y del deserializador, las cuales son las siguientes: *Easy*, *SimpleWriter*, *SimpleReader*, *Encode*, *Serial*, *CreateClass* y *TypeMapping*.

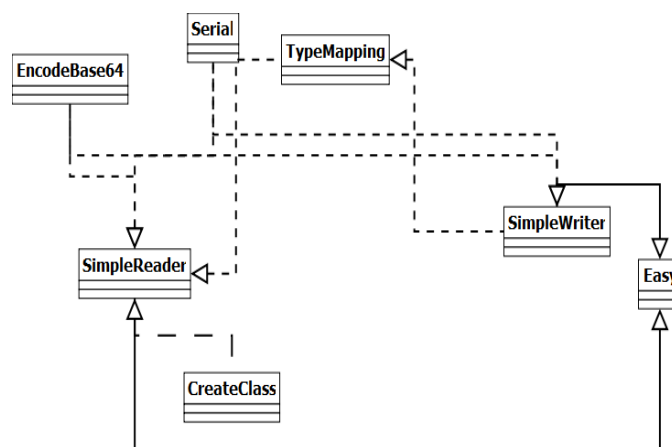


Fig. 3. Clases implementadas para el serializador y para el deserializador.

Easy. Esta clase es utilizada por los usuarios finales para serializar y deserializar objetos a y desde XML. Tiene dos métodos: *save (ob:object, filename:file)*, el cual serializa un objeto a XML y lo almacena a un archivo XML, el primer parámetro es el objeto a ser serializado, y el segundo parámetro es el archivo para almacenar el XML generado; y el método *load (filename:file):Object*, el cual deserializa un objeto desde el archivo XML que es especificado en el parámetro de entrada.

SimpleWriter. Esta clase toma un objeto PHP y lo representa como XML, utilizando el *parser* XML DOM. Utiliza un método *write()* que recibe un objeto PHP, lo analiza para determinar su tipo de dato, extrae sus campos y valores, y lo serializa.

SimpleReader. Esta clase representa el deserializador, lee un objeto de un archivo XML para representarlo como un objeto PHP. Esta clase utiliza el método *read()* para

tomar como punto de partida un elemento DOM del archivo XML, determina su tipo de dato, y extrae su información para crear el objeto específico requerido. Cabe señalar que si la clase del objeto no existe, utiliza la clase *createClass* para crearla.

Encode. Esta clase permite codificar y decodificar matrices de *bytes* base64; tiene varios métodos para realizar las operaciones *encode()* y *decode()*, tales como los siguientes: *encode (byte[] source)*, y *decode (byte[] source)*.

TypeMapping. Esta clase proporciona el mapeo de tipos de datos del lenguaje de programación PHP a WOX y viceversa.

Serial. Interfaz que define constantes para la representación XML de objetos PHP.

CreateClass. Esta clase genera clases; permite crear un documento PHP con una clase, tomando como entrada su nombre y sus atributos. La clase generada tendrá un constructor, métodos *set* y *get* para cada uno de sus atributos.

6. Arquitectura y módulos

Los módulos principales de PHPWOX son los siguientes: 1) módulo serializador, el cual serializa objetos PHP a XML a través de la clase *SimpleWriter*; 2) módulo deserializador, el cual deserializa objetos PHP desde un documento XML a través de la clase *SimpleReader*; y 3) módulo generador de clases, el cual crea clases PHP a partir del XML mediante la clase *CreateClass*. La Figura 4 muestra un diagrama con la funcionalidad de los módulos del serializador y deserializador desarrollados, en donde se muestran los objetos y sus representaciones XML después de haber sido serializados y deserializados. Hay algunas notas indicadas como A1, A2, A3, A4 y A5, las cuales son explicadas en los siguientes párrafos.

A1. En esta sección se observa el objeto *ProductJava* de la clase *Product*, el cual está escrito en el lenguaje de programación Java, con sus atributos y valores. Este objeto es serializado a través del serializador WOX [2, 3], el cual genera el archivo *ProductJava.xml* mostrado a la derecha.

A2. El archivo *ProductJava.xml* es deserializado con el deserializador PHPWOX [1], el cual no tiene la clase *Product*, de tal forma que crea la clase *Product.php*, y crea una instancia de esa clase con los valores indicados en el archivo *JavaProduct.xml*, dando como resultado el objeto *ProductPhp* en el lenguaje de programación PHP, el cual es idéntico al objeto *ProductJava*, pero en PHP.

A3. El objeto *ProductPhp* obtiene nuevos valores en sus atributos.

A4. El objeto *ProductPhp* es serializado por el serializador PHPWOX, obteniendo el archivo *ProductPhp.xml*.

A5. El archivo *ProductPhp.xml* es deserializado por el deserializador WOX en Java [2], obteniendo como resultado el objeto modificado en PHP. En Java este objeto es llamado *ProductJavaMod*, el cual es mostrado en la Figura 4.

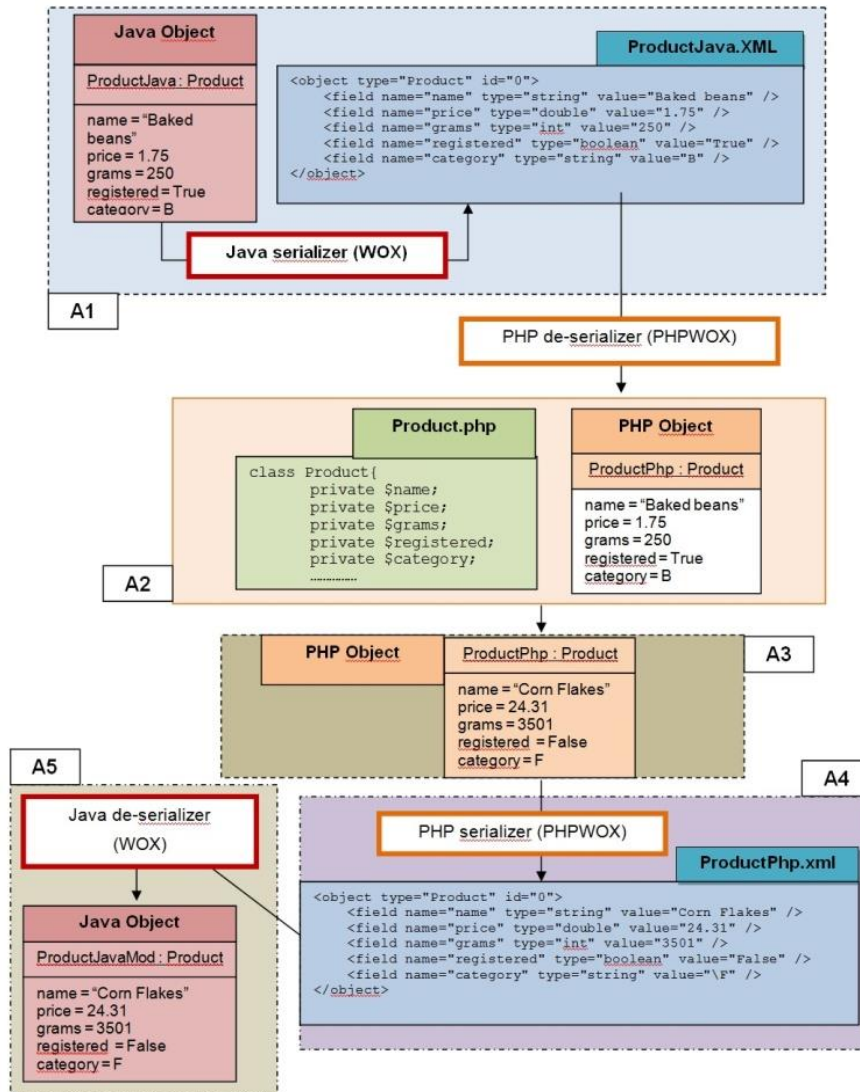


Fig. 4. Diagrama con los módulos para el serializador y deserializador.

7. Conclusiones y trabajo futuro

Este artículo presentó un serializador y deserializador de objetos PHP a XML, llamado PHPWOX. Se puede concluir que el XML generado por PHPWOX es independiente del lenguaje de programación, y puede ser utilizado por otros serializadores y deserializadores WOX existentes, lo cual permite interoperabilidad con los lenguajes de programación Java, C# y Python.

El desarrollo de PHPWOX fue descrito, el conjunto de clases implementadas, la arquitectura y los módulos, con un claro ejemplo de la funcionalidad de los módulos, y mostrando las capacidades de PHPWOX. En el ejemplo presentado se evidencia la interoperabilidad lograda entre PHPWOX y los demás serializadores WOX, ya que se muestran los objetos y sus correspondientes representaciones XML, las cuales fueron obtenidas al utilizar el serializador y deserializador PHPWOX.

Como parte del trabajo futuro es necesario implementar un *framework* por encima de PHPWOX, el cual permita manejar objetos distribuidos, para poder comunicarse con otros *frameworks* WOX que ya están implementados.

Referencias

1. PHPWOX. Disponible en: <http://phpwoxserializer.sourceforge.net/>
2. Serializadores Web Objects in XML (WOX). Disponible en: <http://woxserializer.sourceforge.net/>
3. Jaimez-González, C., Lucas, S., López-Ornelas, E.: Easy XML Serialization of C# and Java Objects. In: Proceedings of the Balisage. The Markup Conference 2011, Balisage Series on Markup Technologies, 7 (2011)
4. Rodríguez-Martínez, A., Jaimez-González, C. R.: Serializador de Objetos a XML en el Lenguaje de Programación Python. Avances de Ingeniería Electrónica 2013, 444–451 (2013)
5. Document Object Model (DOM) para PHP. Disponible en: <http://php.net/manual/es/book.dom.php>
6. XAMPP. Disponible en: <http://www.apachefriends.org/es/download.html>
7. SimpleXML. Disponible en: <http://www.php.net/manual/en/book.simplexml.php>
8. Serialización de objetos (función serialize). Disponible en: <http://php.net/manual/es/language.oop5.serialization.php>
9. json_encode. Disponible en: <http://us2.php.net/manual/es/function.json-encode.php>
10. Igbinary. Disponible en: <http://pecl.php.net/package/igbinary>
11. Pear XML_Serializer. Disponible en: http://pear.php.net/package/XML_Serializer/redirected

Sistema web para registro y administración de cursos de educación continua

Carlos R. Jaimez-González, Wulfrano A. Luna-Ramírez

Universidad Autónoma Metropolitana, Unidad Cuajimalpa,
Departamento de Tecnologías de la Información, México

{cjaimez, wluna}@correo.cua.uam.mx

Resumen. En este artículo se presenta el análisis, diseño e implementación de un sistema web para el registro y administración de los cursos de educación continua que se imparten en la Universidad Autónoma Metropolitana, Unidad Cuajimalpa (UAM-C). El objetivo de este sistema web es apoyar en la automatización de las tareas relacionadas con el manejo de los cursos, talleres y diplomados que se imparten en la UAM-C. Se presentan las etapas que se llevaron a cabo para el desarrollo del sistema web y se proporcionan algunos ejemplos de los artefactos generados, entre los cuales destacan los casos de uso, el diagrama de clases, los diagramas de secuencia, los prototipos de interfaz, entre otros.

Palabras clave: Educación continua, sistema web, metodología de desarrollo, proceso unificado.

Web System for Registration and Administration of Continuing Education Courses

Abstract. This paper presents the analysis, design and implementation of a web system for registration and administration of continuing education courses, which are taught at the Metropolitan Autonomous University, Campus Cuajimalpa (UAM-C). The aim of this web system is to support the automation of tasks related to the administration of courses, workshops and diplomas. The paper also presents the stages carried out for developing the web system, as well as some examples of the artifacts generated, which are mainly use cases, class diagrams, sequence diagrams, interface prototypes, among others.

Keywords. Continuing education, web system, development methodology, unified process.

1. Introducción

En la Universidad Autónoma Metropolitana, Unidad Cuajimalpa (UAM-C) se imparten talleres, diplomados y cursos de educación continua de una amplia variedad de temas, con duración, cupos, requisitos y precios variables. Actualmente, el manejo de estos talleres, diplomados y cursos de educación continua se realiza de manera manual a través de su responsable; desde la apertura de un curso con toda su información relevante, el registro de sus participantes, el pago de las inscripciones respectivas para cada curso, y la búsqueda de aula para su impartición. Un sistema web para automatizar estas tareas permitiría tener un mayor control y manejo del registro y administración de los talleres, diplomados y cursos, por lo que se planteó el diseño e implementación de tal sistema. En este artículo se describen las etapas [1] que se siguieron para el levantamiento de requerimientos, análisis, diseño e implementación de un sistema web para el registro y administración de talleres, diplomados y cursos de educación continua que se imparten en la UAM-C.

El resto del artículo está organizado de la siguiente manera. La sección 2 presenta los requerimientos del sistema web, en donde se identifican dos usuarios principales: el administrador de los cursos y el usuario general del sistema. En la sección 3 se describe la etapa de análisis con la identificación y redacción de algunos casos de uso, así como el diagrama de clases resultante. El diseño del sistema web se discute en la sección 4, en la cual se presenta un mapa de navegación del sistema, un prototipo de interfaz y algunos diagramas de secuencia. La sección 5 presenta la implementación del sistema web, donde se describen las tecnologías y herramientas utilizadas para el desarrollo del sistema, así como algunas capturas de pantalla que muestran el sistema web en funcionamiento. Finalmente, en la sección 6 se proporcionan las conclusiones y el trabajo futuro.

2. Requerimientos

La Figura 1 muestra un ejemplo de la información que actualmente se muestra al anunciar un curso de educación continua en la página web de la UAM-C.

Los requerimientos que debía cumplir el sistema web surgieron a partir de entrevistas con la responsable del área de educación continua, en donde se identificaron dos tipos de usuarios principalmente: los participantes (usuarios generales) y el administrador del sistema. Cabe señalar que los talleres, diplomados y cursos que se imparten son conocidos como servicios en el área de educación continua. A continuación, se muestran los requerimientos funcionales para cada uno de los usuarios.

2.1. Requerimientos para los participantes

El sistema debe proporcionar la siguiente funcionalidad a los participantes que desean inscribirse en un servicio:

Curso de Fundamentos de Java Impartido por: Dr. Carlos González Fecha de Inicio: 24 de Enero 2015 Sesiones: Lunes y Miércoles Horario: 15:00-18:00 Duración: 24 Horas Costos: \$700.00 Comunidad UAM \$1,200.00 Comunidad Externa Requisitos de inscripción y pago: Para completar el proceso de inscripción, se debe realizar el pago correspondiente en Banamex a la Cuenta 464121803, o si lo hace vía Internet con la clave 002180464100218036 a nombre de la Universidad Autónoma Metropolitana. El comprobante original de pago debe ser presentado en el área de Educación Continua, en horarios de 10:00 a 17:00 horas. Nota: En caso de no cubrir con el mínimo de 10 participantes requerido, el curso será pospuesto. Contacto: Área de Educación Continua Teléfono 5257 1489 ext.122
--

Fig. 1. Información de un curso de educación continua.

- Mostrar los servicios que cumplan con ciertos criterios de selección determinados por el usuario, tales como fechas en las que se llevará a cabo el servicio, horarios, tipo de servicio, nombre del servicio, nombre del instructor, etc.
- Una vez que el usuario se haya decidido por un servicio específico, podrá elegirlo y registrarse si el cupo del servicio lo permite. Para el registro de un usuario al servicio será suficiente que proporcione su nombre completo, email, contraseña (para acceso al sistema de inscripciones), teléfono de contacto e institución de procedencia.
- El usuario procederá a elegir el tipo de pago para su inscripción: tarjeta de crédito, cheque, o depósito bancario. El sistema proporcionará confirmación de su registro y emitirá un número de registro único.
- Una vez confirmada su inscripción, el usuario deberá ser capaz de verificar en el sistema web, mediante su email y contraseña, los servicios en los que está inscrito.
- El sistema web deberá proporcionar a sus usuarios la posibilidad de registrarse para recibir notificaciones por email de los servicios que se abrirán posteriormente y que sean de su interés.
- El usuario deberá ser capaz de proponer e impartir un taller, diplomado o curso de educación continua y subir archivos con los temarios.

2.2. Requerimientos para el administrador

El sistema web también deberá proporcionar la siguiente funcionalidad a quienes administran los talleres, diplomados y cursos de educación continua:

- Registro y mantenimiento de servicios. El administrador deberá ser capaz de registrar nuevos servicios, modificar su información y eliminarlos. Cada servicio debe incluir toda la información necesaria, tal como nombre del servicio, tipo del servicio, breve resumen del servicio, temario del servicio, fechas en las que se llevará a cabo el servicio y su horario, precios del servicio, cupo del servicio, descuentos posibles, nombre del instructor o instructores, lugar donde se llevará a cabo el servicio.
- Envío automático de email a alguna cuenta proporcionada explícitamente por el administrador. El envío de emails al administrador será automático cuando sucedan en el sistema eventos tales como: cuando se registre un nuevo participante a un servicio, cuando el cupo de un servicio ha llegado a su límite, cuando explícitamente un participante desea mayor información del servicio, entre otros.
- Algunos reportes útiles para la administración y seguimiento de cada servicio, tales como los siguientes: próximos servicios, número de participantes inscritos en un servicio, servicios impartidos en un determinado periodo de tiempo, servicios impartidos por un instructor determinado, entre otros.

Tabla 1. Caso de Uso Agregar servicio de educación continua.

<i>Caso de Uso</i>	<i>Agregar servicio de educación continua</i>
Descripción	Permite agregar un elemento a la oferta de servicios de educación continua: curso, taller o diplomado.
Actor	Administrador
Precondición	Haber ingresado al sistema.
Antecedente	El sistema mostrará un menú con las siguientes opciones: Agregar servicio, Modificar servicio y Revisar propuestas.
Disparador	Seleccionar la opción Agregar servicio.
Secuencia normal	Paso / Acción 1. El sistema muestra un formulario para ingresar la información del servicio: tipo de servicio, nombre de servicio, nombre del instructor, a quien está dirigido, descripción, fecha de inicio, fecha de fin, horas de duración, días de sesiones, horario, cupo, botón de subir imagen, botón de subir temario en archivo PDF. 2. El usuario llena los campos con la información solicitada. 3. El usuario presiona el botón Listo.
Postcondición	El sistema muestra un mensaje de confirmación de la petición y visualiza todos los servicios almacenados hasta el momento.
Secuencia alternativa	Paso / Acción 3a. El sistema muestra un mensaje de error y regresará al paso 1. Esto es en caso de que la operación en la base de datos no se haya llevado a cabo exitosamente. 3b. El sistema muestra una alerta en caso de que algún campo tenga un formato o tipo de dato incorrecto y regresará al paso 1.

3. Análisis

En esta sección se describe la etapa de análisis del sistema web, en la cual se obtuvieron una serie de casos de uso y un diagrama de clases, pero debido a cuestiones de espacio solamente se presentan las descripciones de dos casos de uso.

3.1. Casos de uso

Los casos de uso identificados para el administrador fueron los siguientes: 1) ingresar al sistema; 2) agregar servicio de educación continua; 3) modificar servicio de educación continua; 4) eliminar servicio de educación continua; 5) enviar notificaciones; y 6) revisar propuestas. En la Tabla 1 se muestra el caso de uso *agregar servicio de educación continua*.

Para el participante en un servicio se identificaron los siguientes casos de uso: 1) mandar propuesta de servicio de educación continua; 2) consultar servicio de educación continua; 3) preinscribirse a un servicio de educación continua; y 4) mandar propuesta de servicio de educación continua. Por restricciones de espacio únicamente se muestra el caso de uso de agregar servicio de educación continua.

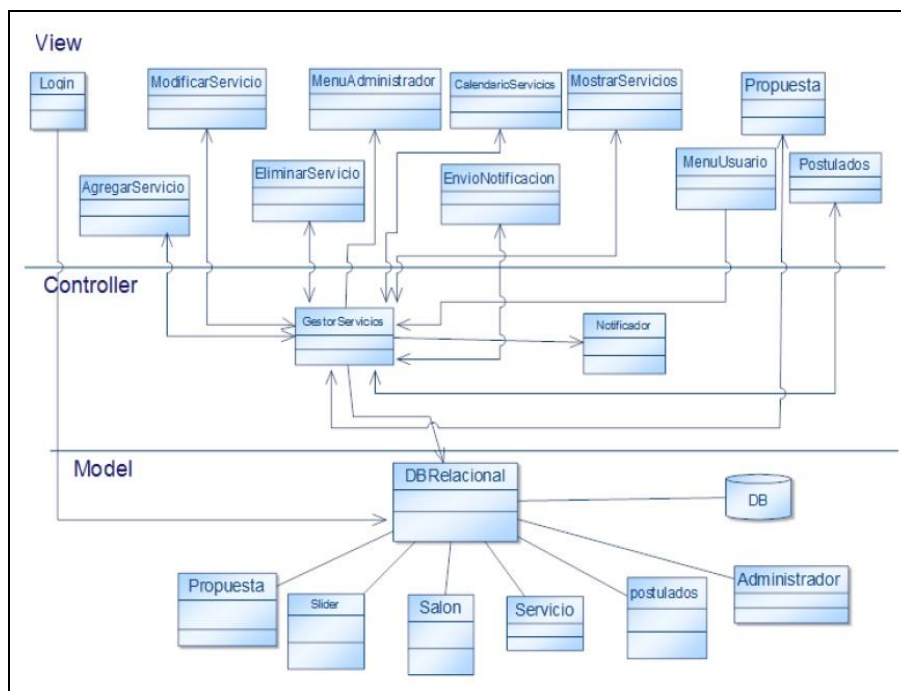


Fig. 2. Diagrama de clases y patrón de diseño MVC.

3.2. Diagrama de clases

En la Figura 2 se muestra el diagrama de clases del sistema web, donde se observa la dependencia de clases y el uso del patrón de diseño [2] Modelo-Vista-Controlador (*Model-View-Controller*, MVC por sus siglas en inglés).

En la parte de *View* se muestran las clases correspondientes a las interfaces gráficas que el administrador y el usuario general pueden acceder. *Controller* comprende las clases que llevan el control del sistema y redireccionan a diversas interfaces de acuerdo a las acciones realizadas por los usuarios del sistema. Finalmente, en el *Model* se maneja el acceso a la fuente de datos del sistema web, principalmente a la clase DBRelacional, la cual es la encargada de la manipulación de datos del sistema.

4. Diseño

Esta sección proporciona el diseño del sistema web a través de varios artefactos: el mapa de navegación, el prototipo de interfaz y los diagramas de secuencia.

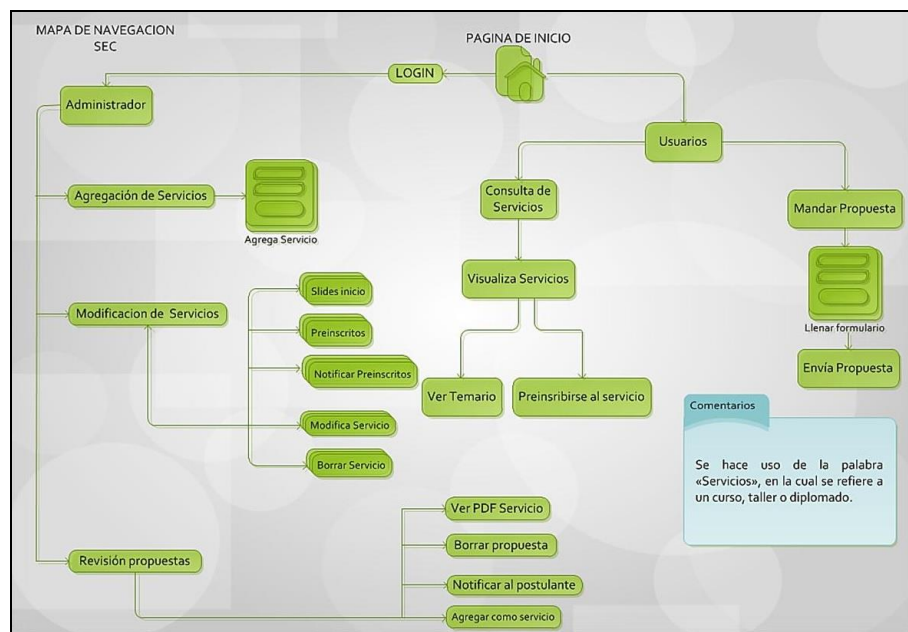


Fig. 3. Mapa de navegación del sistema web.

4.1. Mapa de navegación

El mapa de navegación que se muestra en la Figura 3 es la representación gráfica de la organización del sistema web, en el cual se expresan todas las relaciones de jerarquía y secuencia de las páginas web que componen al sistema web, además de que permite elaborar escenarios de comportamiento de los usuarios. El mapa de nave-

gación refleja los dos usuarios que tiene el sistema: el administrador y el participante (usuario general).

El administrador tiene tres opciones de navegación: agregar servicio, modificar servicio y revisar propuesta. La modificación de un servicio incluye la modificación de la lista de alumnos preinscritos en el servicio, la notificación de preinscripción a los alumnos, la modificación de los datos del servicio y la eliminación del servicio. La revisión de propuestas tiene cuatro opciones: revisar el archivo PDF con la propuesta, eliminar propuesta, notificar al postulante de la propuesta y agregar propuesta como servicio. El usuario general puede consultar los servicios disponibles, preinscribirse a un servicio y enviar propuestas de servicios.

A partir del mapa de navegación un diseñador puede realizar una vista general de cada una de las páginas que componen al sistema web, ya sea mediante el uso de prototipos de interfaz o a través de páginas web HTML estáticas.

4.2. Prototipo de interfaz

La Figura 4 muestra el prototipo de interfaz de la página de inicio del sistema web, la cual puede ser visualizada por cualquier usuario. Como se muestra en el mapa de navegación, un usuario general solamente puede realizar tres operaciones dentro del sistema: consulta de servicio, mandar propuesta y preinscripción a un servicio. Al seleccionar cualquiera de estas opciones, se mostrará la página web correspondiente que se especifica en el mapa de navegación.



Fig. 4. Prototipo de interfaz de la página de inicio del sistema web.

4.3. Diagramas de secuencia

En esta sección se proporciona un ejemplo de diagrama de secuencia: la Figura 5 muestra el diagrama de secuencia *Agregar servicio de educación continua*, el cual es llevado a cabo por un usuario administrador. El diagrama de secuencia mostrado en esta sección corresponde al caso de uso del mismo nombre, el cual fue presentado previamente en la sección de análisis.

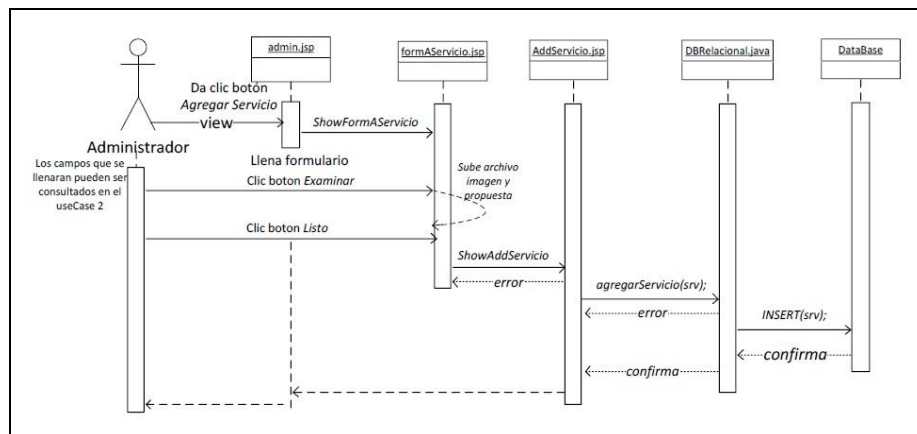


Fig. 5. Diagrama de Secuencia *Agregar servicio de educación continua*.

5. Implementación

La implementación del sistema web se llevó a cabo una vez que se concluyeron las etapas de análisis y diseño, a partir de las que se obtuvieron los siguientes artefactos: casos de uso, diagrama de clases, mapa de navegación, prototipos de interfaz, diagramas de secuencia y diagrama de componentes. En esta sección se describen las herramientas y tecnologías que fueron utilizadas para el desarrollo del sistema, así como algunas capturas de pantalla del sistema web terminado.

5.1. Tecnologías y herramientas utilizadas

Las tecnologías y herramientas utilizadas para la implementación del sistema web se describen a continuación.

HTML [3]. El lenguaje de marcado de hipertexto (*Hypertext Markup Language*, HTML por sus siglas en inglés) se utilizó para crear el contenido de las páginas web que componen el sistema.

JSP [4]. Las páginas Java del lado del servidor (*Java Server Pages*, JSP por sus siglas en inglés) se utilizaron para la creación de contenido dinámico que se ejecuta en un servidor web. Las JSP contienen una mezcla de HTML y segmentos de código Java (scripts) que se ejecutan en un servidor web.

MySQL [5]. Es un manejador de bases de datos relacionales (*Database Management System*, DBMS por sus siglas en inglés) que se utilizó para crear la base de datos del sistema web y para permitir la ejecución de sentencias del lenguaje de consulta estructurado (*Structured Query Language*, SQL por sus siglas en inglés).

CSS [6]. Es un lenguaje para crear hojas de estilo en cascada (*Cascading Style Sheet*, CSS por sus siglas en inglés) que se utilizó para proporcionar la presentación de las páginas que componen el sistema web.

JavaScript [7]. Es un lenguaje script basado en objetos, diseñado específicamente para hacer que las páginas web sean dinámicas e interactivas. JavaScript es un lenguaje para hacer programación web dinámica del lado del cliente, el cual se utilizó para proporcionar interactividad a las páginas del sistema web.

AJAX y JQuery [8]. *AJAX* se utilizó para dar mayor interactividad a las páginas web sin la necesidad de refrescar la página web completa, ya que funciona a través de llamadas asíncronas al servidor web. Por otro lado, *JQuery* es una biblioteca que contiene funciones para realizar ciertas acciones en una página web.

Notepad++ [9] y VIM [10]. La mayor parte de la codificación se realizó en Notepad++ y VIM, ya que permiten la indentación automática, coloreado y resaltado de código, así como varias facilidades disponibles para la programación.

Kompozer [11]. Es un editor de HTML que permite la creación de páginas web de manera gráfica, mediante la colocación de elementos sobre un panel.

Neatbeans [12]. Es un ambiente de desarrollo integrado (*Integrated Development Environment*, IDE por sus siglas en inglés) que permite la creación y edición de archivos HTML, CSS, JavaScript, Java, entre otros.



Fig. 6. Página de inicio del sistema.

5.2. Interfaz del sistema web

A continuación se proporcionan algunas capturas de pantalla del sistema web terminado. La Figura 6 muestra la página de inicio del sistema, donde se observa la oferta de servicios de educación continua y dos cajas de texto para introducir nombre de usuario y contraseña para poder acceder al sistema. En la Figura 7 se observa la pági-

na para agregar un nuevo servicio (taller, diplomado o curso de educación continua) al sistema.

Fig. 7. Página para agregar un servicio al sistema.

En la Figura 8 se muestra la página que permite editar los detalles de un servicio de educación continua ya registrado. En la Figura 9 se observa la página que proporciona la lista de propuestas de servicios de educación continua, en la que se incluye el tipo de servicio (curso, taller o diplomado), el instructor que impartirá el servicio, el título de la propuesta y un menú de operaciones para manipular la propuesta.

Fig. 8. Página para ver y editar un servicio.

6. Conclusiones y trabajo futuro

En este artículo se presentó el análisis, diseño e implementación de un sistema web para el registro y administración de los servicios de educación continua que se imparten en la UAM-C. El objetivo de este sistema web es apoyar en la automatización de

las tareas relacionadas con el manejo de los cursos, talleres y diplomados que se imparten en la UAM-C, tales como el registro, la preinscripción e inscripción de participantes a los servicios ofertados, notificación de nuevos servicios, envío de propuestas de servicios de educación continua, revisión y aprobación de servicios, entre otras.



Tipos de servicio	Imparte	Título	Menú de Operaciones
Curso	Carlos Jaimez	JavaEE	[Iconos de operaciones]
Taller	Arturo Wulfrano	PHP	[Iconos de operaciones]
Diplomado	Rafael PerÁz	IA	[Iconos de operaciones]
Taller	Benedicto XVI	Curso de hay mama	[Iconos de operaciones]
Diplomado	Erick Lopez	Programación logica	[Iconos de operaciones]

Fig. 9. Página para visualizar las propuestas de servicios.

Se presentaron las etapas que se llevaron a cabo para el desarrollo del sistema web y se proporcionaron algunos ejemplos de los artefactos generados, entre los cuales destacan los casos de uso, el diagrama de clases, los diagramas de secuencia, los prototipos de interfaz, entre otros.

Dentro del trabajo futuro se planean realizar pruebas de funcionalidad y usabilidad del sistema con la responsable del área de educación continua (usuario administrador) y con participantes potenciales de los servicios de educación continua que se ofrecen.

Referencias

1. Jacobson, I., Booch, G., Rumbaugh, J.: The Unified Software Development Process. Boston, MA, USA, Addison-Wesley Longman Publishing Co., Inc., 1ra edición (1999)
2. Gamma, E., Helm, R., Johnson, R., Vlissides, J.: Design Patterns: Elements of Reusable Object-oriented Software. Boston, MA, USA: Addison-Wesley Longman Publishing Co (1995)
3. Jaimez-González, C.: Programación de Web Estático. Colección Una Década, México: Universidad Autónoma Metropolitana, Unidad Cuajimalpa, Disponible en: http://www.cua.uam.mx/pdfs/biblioteca/colecciondelibros-uamc/pdfs/05programacion_web_estatico.pdf (2015)
4. H. Bergsten H.: Java Server Pages Pocket Reference. Sebastopol, CA, USA. O'Reilly Media (2009)
5. MySQL Documentation. Disponible en: <https://dev.mysql.com/doc>
6. CSS Reference. W3C. Disponible en: <https://www.w3.org/Style/CSS/>
7. Jaimez-González, C.: Programación de Web Dinámico. Colección Una Década, México: Universidad Autónoma Metropolitana, Unidad Cuajimalpa, Disponible en: http://www.cua.uam.mx/pdfs/biblioteca/colecciondelibros-uamc/pdfs/04Programacion_web_dinamico.pdf (2015)
8. Ajax y JQuery. Disponible en: <https://api.jquery.com/category/ajax/>

9. Notepad++. Editor de código fuente. Disponible en: <https://notepad-plus-plus.org/>
10. VIM. Editor de texto. Disponible en: <http://www.vim.org/>.
11. Kompozer. Sistema web de autoría. Disponible en: <http://kompozer.net/>
12. Netbeans. Ambiente de desarrollo integrado, Disponible en: <https://netbeans.org/>

Hacia una herramienta web para visualizar la ejecución de programas escritos en el lenguaje Java

Miguel Castillo-Cortes, Carlos R. Jaimez-González

Universidad Autónoma Metropolitana, Unidad Cuajimalpa,
Departamento de Tecnologías de la Información, México

210368127@alumnos.cua.uam.mx, cjaimez@correo.cua.uam.mx

Resumen. Este artículo presenta una propuesta de herramienta web de apoyo para el aprendizaje de la programación a nivel universitario, la cual permitirá visualizar una animación que representa la ejecución de programas escritos en el lenguaje de programación Java. La herramienta propuesta apoyará en la comprensión de programas y entendimiento de los conceptos básicos de la programación, tales como creación de variables, asignación de valores a variables, uso de estructuras de control de flujo de programa y llamadas a funciones con paso de parámetros. En este artículo también se presenta un análisis comparativo de herramientas con funcionalidad similar a la propuesta.

Palabras clave: Tecnología educativa, computación en educación, comprensión de software, visualización de software.

Towards a Web Tool to Visualize the Execution of Programs Written in the Java Programming Language

Abstract. This paper presents a proposal of web tool to support the learning of programming at university level, which will allow to visualize an animation that represents the execution of programs written in the Java programming language. The proposed tool will help to understand programs and programming basic concepts, such as the creation of variables, setting values to variables, use of control structures and function calls with parameters. This paper also presents a comparative analysis of tools with similar functionality to the proposal.

Keywords. Educational technology, computer science in education, understanding software, software visualization.

1. Introducción

Entender un programa es tener la habilidad de comprender de qué forma el código fuente realiza la tarea para la cual fue creado. La comprensión de programas es un área de la Ingeniería de Software enfocada a elaborar técnicas y herramientas, basadas en procesos cognitivos y de ingeniería, que son utilizadas para tareas de reutilización, inspección, mantenimiento, migración, extensión de software, entre otras aplicaciones; pero también se pueden emplear en áreas como la educación o la capacitación. Esto con el objetivo de lograr un mejor entendimiento de las aplicaciones informáticas [1].

Uno de los principales retos en el área de comprensión de programas es poder determinar la relación existente entre el dominio del problema y el dominio del programa [2]. El primero es el resultado de la ejecución del programa, por ejemplo, imprimir en pantalla el resultado de una operación; el segundo son todos los componentes del programa, que producen el comportamiento del sistema, por ejemplo, los métodos utilizados para poder resolver alguna operación [2]. Para poder representar estas relaciones, entre dominio del problema y el dominio del programa, se recurre al uso de la visualización de software (VS), la cual es una rama de la Ingeniería de Software cuyo objetivo es generar, a partir de ciertos aspectos del sistema, una o más representaciones multimedia, con la finalidad de facilitar el entendimiento del mismo [3,4].

En este artículo se presenta una propuesta de herramienta web destinada a facilitar el proceso de comprensión de programas escritos en el lenguaje de programación Java. Esta herramienta facilitará el aprendizaje debido a que fomentará la interconexión de conceptos básicos de la programación. De esta manera, se reforzarán los conocimientos del estudiante mediante representaciones visuales en el sistema.

El resto del artículo está organizado de la siguiente manera. La sección 2 presenta el marco teórico. La sección 3 describe una serie de herramientas similares a la propuesta y hace un análisis comparativo. En la sección 4 se presenta la propuesta de herramienta para visualizar la ejecución de programas escritos en el lenguaje Java. Finalmente, las conclusiones y el trabajo futuro se presentan en la sección 5.

2. Marco teórico

En el desarrollo de sistemas de comprensión de software es importante apoyarse de otros campos de estudio, tales como psicología cognitiva, la psicopedagogía, la visualización de programas, entre otros. En los siguientes apartados se abordan los temas de modelos cognitivos y visualización de software, los cuales son muy relevantes para la propuesta que se presenta en este artículo.

2.1. Modelos cognitivos

El desarrollo cognitivo es un campo de investigación de la psicología cognitiva que estudia los procesos mentales involucrados en la creación de conocimiento. Aprender es combinar el conocimiento ya adquirido con nuevos conceptos mediante procesos de aprendizaje; por ejemplo, el modelo *Bottom Up* que parte de conceptos específicos

obtenidos de la lectura del programa para posteriormente generar abstracciones generales; o el modelo *Top Down*, el cual supone que ya se tiene noción de la funcionalidad del programa para proponer una hipótesis, que posteriormente se verificará al revisar el código fuente [5,6].

Los modelos cognitivos están conformados por tres elementos base: el conocimiento, el proceso de asimilación y un modelo mental. El primero plantea dos distinciones: el conocimiento interno, que se refiere a los conocimientos que ya se poseen; y el conocimiento externo, que son los nuevos conceptos que proveerá el sistema. El segundo elemento lo conforman las estrategias de aprendizaje, tales como *Bottom Up* o *Top Down*. Finalmente, el tercero es una representación que se le da al sistema, con la cual se intenta explicar el comportamiento del mismo [5,6].

2.2. Visualización de software

La visualización de programas es una rama de la Ingeniería de Software relacionada con la representación visual de la información, cuyo objetivo es generar, a partir de ciertos aspectos del sistema, una o más vistas que facilitarán el entendimiento del mismo. Una vista es una forma de representar los elementos de un sistema y las relaciones entre estos elementos [5].

Las representaciones multimedia orientadas a la comprensión de software deben aportar tres vistas básicas: la salida del sistema, los elementos del programa que generen dicha salida, y la relación entre las dos anteriores. Éstas permiten elaborar representaciones que asocian el dominio del problema con el dominio del programa; que facilitan relacionar los conocimientos que se poseen y los conceptos usados por el sistema [4].

Dichas representaciones no son fáciles de construir, dado que están fuertemente basadas en factores cognitivos y de ingeniería, por lo que se vuelve importante apoyarse de varias áreas del conocimiento como lo son el diseño gráfico, la psicología cognitiva, la psicopedagogía, la computación y otras disciplinas relacionadas con la creación de efectos multimedia y el aprendizaje [3].

Los aspectos del software necesarios para elaborar una representación se obtienen usando técnicas de extracción de la información que se clasifican por el tipo de información que extraen, la cual puede ser estática o dinámica [2]. Las técnicas de extracción de la información estática son utilizadas para verificar que no existan errores sintácticos, lexicográficos o incluso semánticos, haciendo uso de herramientas de análisis sintáctico para examinar el código fuente [2]. Las técnicas de extracción de la información dinámica se encargan de obtener información sobre el comportamiento del programa y sus componentes (métodos, variables, invocaciones, etc.) en tiempo de ejecución. Una de las técnicas para obtener estos datos, es el uso de la instrumentación de código (IC), la cual consiste en colocar instrucciones en el código fuente con el fin de monitorear su comportamiento durante la ejecución [2].

La visualización de software incluye dos áreas fundamentales: la visualización de programas y la visualización de algoritmos, dependiendo de lo que se desea comprender. La visualización de programas permite tener vistas del código fuente y su estructura. Una representación estática es de utilidad para verificar la validez del código

fuelle, usualmente se utilizan editores de código para dar una mejor presentación y legibilidad del código fuente mediante indentado, diferencias de colores entre palabras reservadas y otros identificadores. La representación dinámica muestra información del programa en tiempo de ejecución, por ejemplo, resaltando las instrucciones del código cuando éstas están siendo ejecutadas [7]. La visualización de algoritmos se enfoca en representar la semántica del programa, esto significa que genera vistas del comportamiento del programa en ejecución sin mostrar las instrucciones y operaciones que hacen posible dicho comportamiento [7]. Estos conceptos influyen en el desarrollo de herramientas de comprensión de software.

3. Estado del arte

En esta sección se hace una revisión del estado del arte de herramientas para la comprensión de programas. También se proporciona un análisis comparativo al final de la sección.

3.1. BlueJ

BlueJ [8] es un entorno de desarrollo diseñado para aprender programación orientada a objetos con el lenguaje Java. Utiliza técnicas de visualización e interacción para mejorar la experiencia del usuario. Es un entorno en constante desarrollo, el cual fue creado por la Universidad de Kent en Canterbury, Reino Unido, y la Universidad de La Trobe, en Melbourne. BlueJ colorea el fondo de cada bloque de código para identificar visualmente las secciones del programa, ayuda en la detección de llaves fuera de lugar y muestra detalles de los objetos instanciados.

3.2. Jeliot 3

Jeliot 3 [9] es una herramienta de visualización de software que muestra cómo se interpreta un programa en lenguaje Java; las llamadas a métodos, variables y operaciones se visualizan mediante una representación multimedia que permite seguir paso a paso el proceso de ejecución del programa. En la Figura 1 se observa una captura de pantalla de Jeliot3, donde se muestran 3 paneles: el panel superior izquierdo se utiliza para ingresar el código que será ejecutado; el panel inferior izquierdo muestra la salida de la ejecución del programa; mientras que el panel de la derecha visualiza la representación del programa en ejecución, donde se tienen cuatro áreas, que identifican la inicialización de variables, la representación de operaciones y la de arreglos.

3.3. jGrasp

jGrasp [10] es un entorno de desarrollo que genera de forma automática visualizaciones animadas para mejorar la comprensión del software. jGrasp es capaz de generar diagramas de estructuras de control para los lenguajes Java, C, C++, entre otros. También crea diagramas UML. El visualizador representa estructuras de datos como

pilas, colas, listas ligadas, etc. En la Figura 2 se observa una captura de pantalla de jGrasp en la cual se muestra el área donde se escribe el código fuente, al ejecutar el código fuente muestra en el panel izquierdo inferior las instancias de los elementos creados en el programa como lo son las variables, arreglos, etc. En una ventana independiente se puede visualizar de forma gráfica la representación de estas instancias.

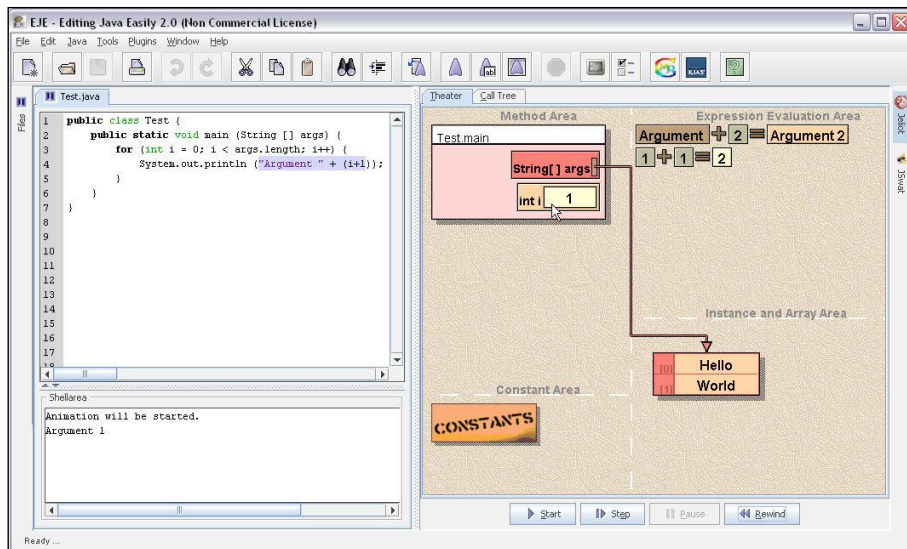


Fig. 1. Ejecución y representación de un segmento de código en Jeliot 3.

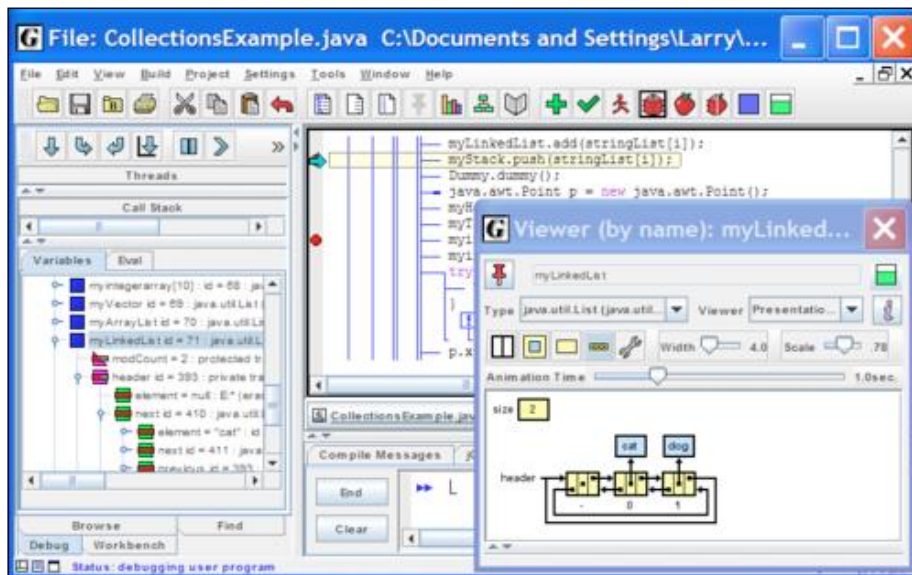


Fig. 2. Ejecución y representación de un segmento de código en jGrasp.

3.4. Scratch

Scratch [11] es una aplicación en la cual se pueden programar historias interactivas, juegos y animaciones, con tan solo definir reglas mediante bloques de instrucciones. Los proyectos desarrollados en esta aplicación pueden ser compartidos dentro de su catálogo de proyectos, de esta forma se tiene acceso a las aplicaciones creadas por terceros. En la Figura 3 se puede observar un programa realizado en la plataforma Scratch, que muestra una ventana con tres secciones, en la parte de la izquierda se tiene el área donde se visualizará la animación del programa y la parte central hay un conjunto de bloques de instrucciones que pueden ser usadas por el usuario; instrucciones tales como avanzar, girar, cambiar dirección, cambiar valor de una variable entre otras. En la tercera sección se encuentra el área donde el usuario arrastrará los bloques de instrucción que desea usar para su animación y los podrá acomodar como si de un rompecabezas se tratase. Para ejecutar la animación se tiene que dar clic sobre los bloques que se unieron.

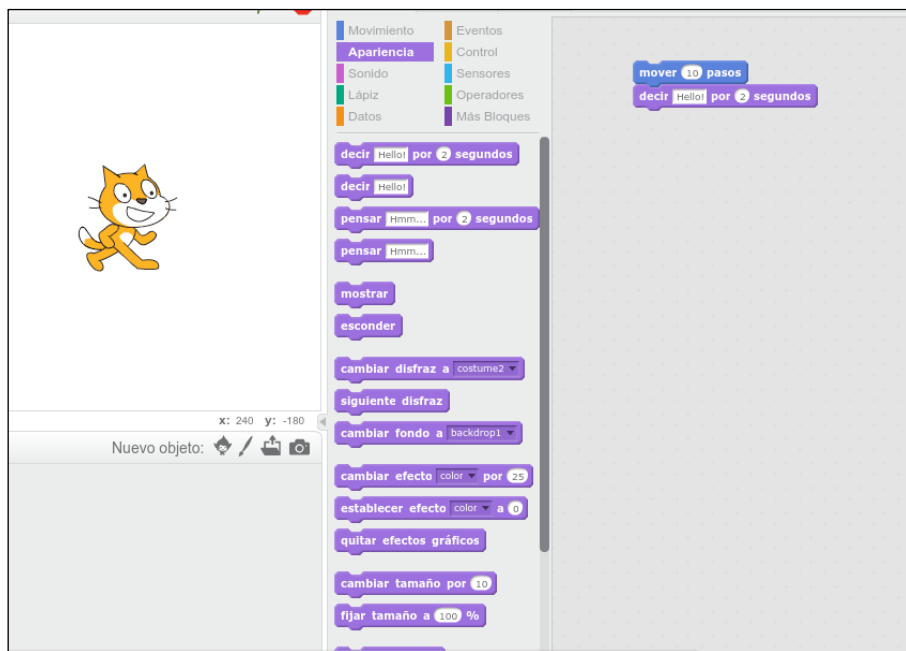


Fig. 3. Ejecución de un programa en la plataforma Scratch.

3.5. Otras herramientas

Existen otras herramientas para la comprensión de software que no se analizaron a profundidad, como las que se mostraron anteriormente. Esto fue debido a que estas herramientas no se consideraron tan relevantes como las que se revisaron detalladamente en las secciones anteriores. Algunas de éstas son las siguientes: Understand

[12], CodeSurfer [13], Imagix 4D [14], ShriMP [15], Alma [16], SeeSoft [17], Extra-vis [18], entre otras.

3.6. Comparación

En esta sección se presenta una tabla comparativa de características de las herramientas expuestas anteriormente y la herramienta propuesta. Las herramientas que se muestran en la Tabla 1 son las siguientes: H1) BlueJ, H2) Jeliot3, H3) jGrasp, H4) Scratch, H5) Herramienta web propuesta. El símbolo de verificación indica que la herramienta tiene la característica, mientras que la x indica que no la tiene.

Tabla 1. Características de las herramientas analizadas y la propuesta.

<i>Características</i>	<i>H1</i>	<i>H2</i>	<i>H3</i>	<i>H4</i>	<i>H5</i>
Representación del código con animaciones	✗	✓	✓	✓	✓
Aplicación web	✗	✗	✗	✓	✓
Representación de clases	✗	✓	✓	✗	✗
Iluminado de bloques de código	✓	✗	✗	✗	✓
Código editable	✓	✓	✓	✗	✓
Representación visual de métodos	✗	✗	✗	✗	✓
Representación visual de estructuras control	✗	✗	✗	✓	✓
Representación visual de paso parámetros	✗	✓	✓	✗	✓

4. Propuesta

Después de analizar las diferentes herramientas en la tabla de comparación, la herramienta web propuesta en este artículo servirá de apoyo al aprendizaje de la programación, ya que facilitará la comprensión de los programas vistos en clase o creados por los estudiantes.

El estudiante podrá entrar a la herramienta web y escribir su código en la zona correspondiente; posteriormente podrá hacer click en la opción para iniciar la animación, ya sea una animación continua o paso a paso. Para el caso de la animación continua se generará sin pausas de principio a fin; mientras que en la animación paso a paso se hará una pausa por cada instrucción de código, para que el estudiante interactúe con el sistema y decida si prosigue al siguiente paso de la animación.

El entorno de la herramienta web que se desarrollará será una combinación de las diferentes herramientas ya estudiadas anteriormente, donde se tendrá una sección para escribir el código fuente elaborado en Java, una sección donde se dibujarán automáticamente todas las animaciones, y finalmente, la sección de los controles donde se podrá elegir si se inicia la animación, se reinicia o se ejecuta paso a paso. La herramienta web constará de distintos módulos independientes, tanto para el reconocimiento del código como para la animación y sus distintos escenarios.

La base de la herramienta está pensada para estudiantes principiantes en el tema, por lo cual solo se logrará hacer que la herramienta reconozca ciertos tipos de varia-

bles tanto privadas como públicas, tales como string, int, boolean, double y arreglos de los tipos mencionados; algunas estructuras de control de flujo de programa, tales como for, if, else, elseif, while, do while; y llamadas a funciones con parámetros. Para el desarrollo de la herramienta se utilizarán las siguientes tecnologías y lenguajes: HTML, JavaScript, CSS, JQuery, JQuery-ui, y three.js.

4.1. Objetivos

El objetivo general es diseñar e implementar una herramienta web enfocada a principiantes, para visualizar animaciones que mejoren la comprensión de programas escritos en el lenguaje de programación Java. Los objetivos particulares son los siguientes: 1) diseñar e implementar un mecanismo para identificar el programa del estudiante; 2) diseñar e implementar un mecanismo para desplegar automáticamente en pantalla variables en memoria, llamadas a métodos y estructuras de control de flujo de programa; 3) diseñar e implementar un mecanismo para permitir al estudiante seguir paso a paso la ejecución de un programa; 4) diseñar e implementar el sitio web; 5) determinar el conjunto de elementos que será posible reconocer por el sistema.

4.2. Metodología

El desarrollo de la herramienta web para visualizar la ejecución de programas en Java considera realizar las siguientes actividades: 1) elaboración del estado del arte, el cual ya fue concluido; 2) creación de un editor de código, el cual ya se ha iniciado; 3) diseño e implementación de un analizador léxico, básico, con el cual se pretende distinguir a partir del código fuente, los elementos que serán representados, el cual ya está en proceso de elaboración; 4) elaboración de los mecanismos para reconocer variables tanto públicas como privadas, inicialización de variables, asignación de valores, tipos de datos, métodos, parámetros enviados a través de los métodos y algunas estructuras de control; 5) diseñar e implementar los módulos de animación para los elementos que serán representados, dichos módulos serán para representar variables, funciones y estructuras de control; 6) elaborar los mecanismos que dibujen en pantalla las representaciones del código fuente; 7) desarrollar los mecanismos de animación para las representaciones; 8) implementar la interacción de la animación con el usuario; 9) desarrollo de controladores para manipular la ejecución del código fuente y de su respectiva representación.

4.3. Interfaz preliminar de la herramienta

En esta sección se muestran algunas capturas de pantalla de la interfaz preliminar de la herramienta web para visualizar la ejecución de programas escritos en el lenguaje Java. En la Figura 4 se muestra la captura de pantalla de la interfaz de inicio de la herramienta, donde se aprecian tres paneles: el panel del lado izquierdo es para escribir el programa en el lenguaje Java. El panel central es donde se visualizará la animación de la ejecución del programa, y el panel inferior contiene dos botones, uno para ejecutar el programa de manera continua, y uno para ejecutarlo o paso a paso.

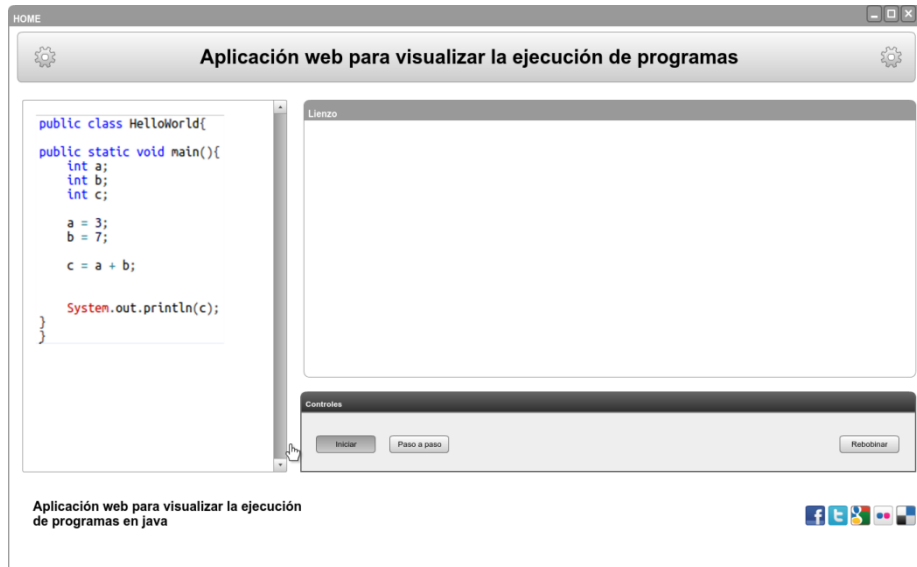


Fig. 4. Vista preliminar de la herramienta propuesta.

En la Figura 5 se muestra la captura de pantalla de la interfaz, una vez que ha empezado la ejecución del programa. Se aprecia que el programa ya ha creado tres variables (a, b y c) en el panel de visualización, las cuales aún no tienen ningún valor asignado. También puede apreciarse que el nombre del procedimiento actual es main(). La flecha muestra en el panel izquierdo la instrucción que se está ejecutando.

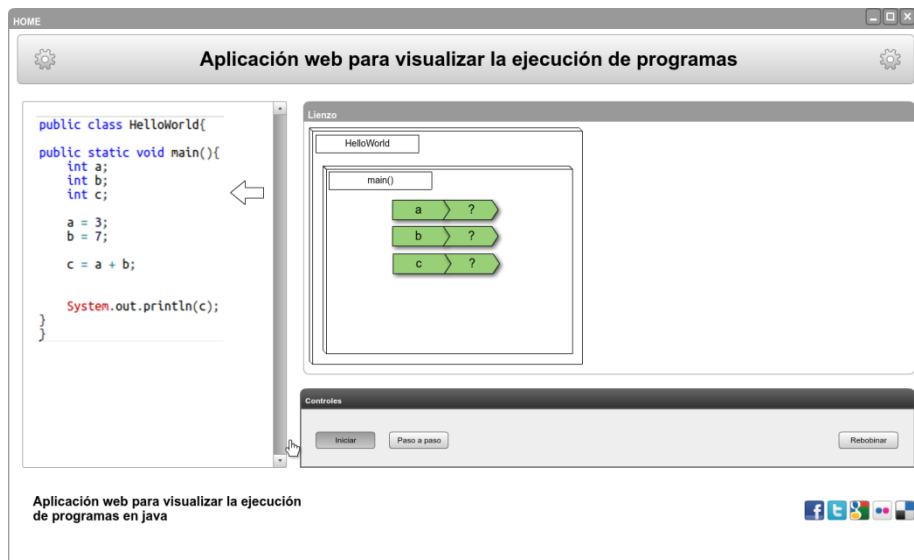


Fig. 5. Vista preliminar de la herramienta al ejecutarse el programa.

En la Figura 6 se muestra la interfaz una vez que se han asignado valores a las variables a y b. En el panel de visualización se puede apreciar la animación que dibuja $3 + 7 = 10$, para asignarlo a c. Estos pasos se van mostrando en la herramienta de manera animada. La Figura 7 muestra cómo el valor de 10, resultado de la operación de $a + b$, viaja del área de operaciones hacia la casilla de memoria donde está la variable c.

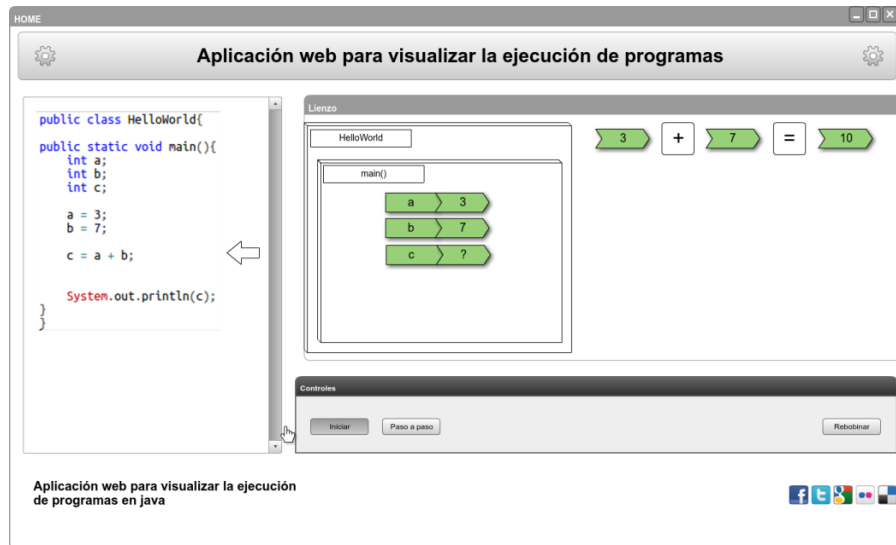


Fig. 6. Ejecución de una operación y una asignación.

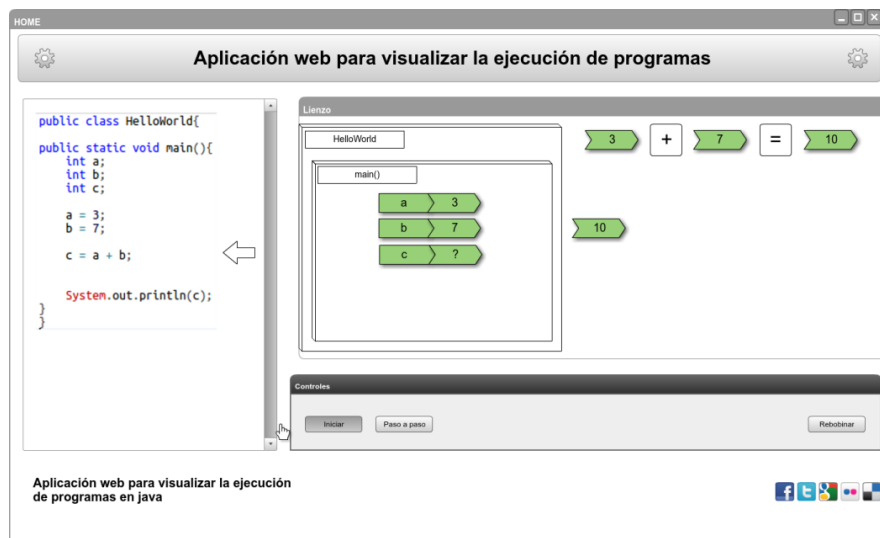


Fig. 7. Resultado de operación viaja a casilla de memoria.

En la Figura 8 se observa que el valor de 10 ha llegado a la casilla de memoria de la variable *c*, y el signo de interrogación ha desaparecido. Finalmente, la última instrucción del programa imprime el valor de la variable *c* a la consola, por lo que en el panel de visualización se presenta una consola donde el valor de *c* se imprime.

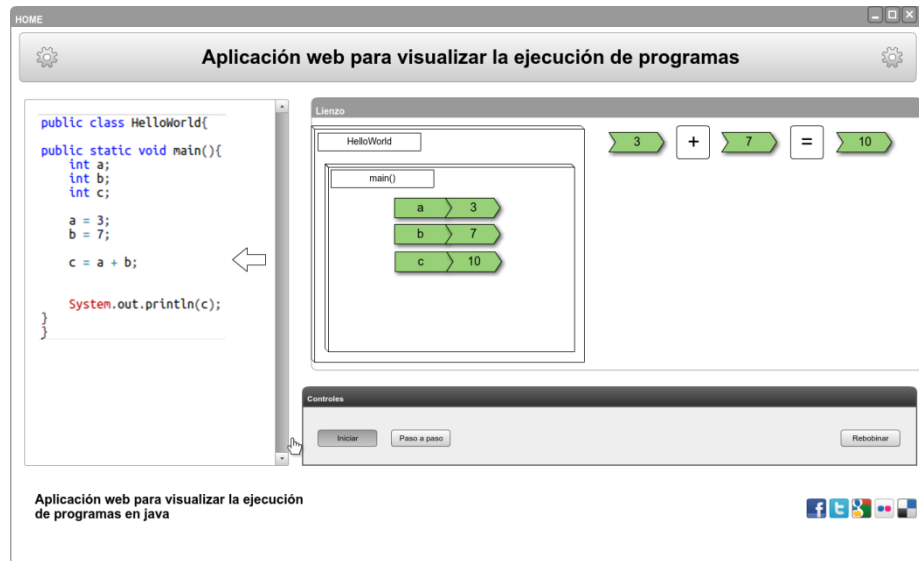


Fig. 8. Resultado de operación llega a casilla de memoria.

5. Conclusiones y trabajo futuro

El propósito de este artículo fue presentar una propuesta de herramienta web de apoyo para el aprendizaje de la programación a nivel universitario, la cual permitirá visualizar una animación que representará la ejecución de programas escritos en el lenguaje de programación Java. La herramienta propuesta apoyará en la comprensión de programas y entendimiento de los conceptos básicos de la programación, tales como creación de variables, asignación de valores a variables, uso de estructuras de control de flujo de programa y llamadas a funciones con paso de parámetros.

Como trabajo futuro se deben llevar a cabo las actividades propuestas en la metodología presentada, las cuales fueron concebidas después de realizar un análisis detallado de los requerimientos que la herramienta propuesta contendrá. Asimismo, se llevarán a cabo pruebas de funcionalidad y usabilidad con estudiantes y profesores.

Una vez que se tenga una versión estable de la herramienta web, se planea incorporar a una plataforma web de tutoriales interactivos que se ha diseñado y desarrollado por uno de los autores de este artículo. Esto último permitirá a los docentes crear tutoriales que utilicen esta herramienta web para sus ejemplos y ejercicios, principalmente para cursos relacionados con la programación.

Referencias

1. Berón, M., Uzal, R., Henriques, P. R., Varanda Pereira, M. J.: Comprensión de programas por inspección visual y animación. In: IX Workshop de Investigadores en Ciencias de la Computación (2007)
2. Bernardis, H., Berón, M., Riesco, D., Henriques, P., Pereira, M. J.: Extracción de información dinámica en programación orientada a objetos (Java). In: XIII Workshop de Investigadores en Ciencias de la Computación, Red de Universidades Nacionales con Carreras de Informática (RedUNCI) (2011)
3. Berón, M., Riesco, D., Montejano, G., Rangel, P., Pereira, M.: Estrategias para Facilitar la Comprensión de Programas. In: Workshop de Investigadores en Ciencias de la Computación, El Calafate, Santa Cruz, Argentina (2010)
4. Miranda, E., Berón, M., Montejano, G., Riesco, D., Henriques, P., Pereira, M. J.: Visualización de software: conceptos, métodos y técnicas para facilitar la comprensión de programas. In: XIII Workshop de Investigadores en Ciencias de la Computación, Red de Universidades Nacionales con Carreras de Informática (RedUNCI) (2011)
5. Berón, M., Uzal, R., Henriques, P. R., Varanda Pereira, M. J.: Inspección de código para relacionar los dominios del problema y programa para la comprensión de programas. In: X Workshop de Investigadores en Ciencias de la Computación (2008)
6. Berón, M.: Tesis Doctoral Inspección de Programas para Interconectar las Vistas Comportamental y Operacional para la Comprensión de Programas
7. Aquino, J. R., Hernández, J. H., Gutiérrez, D. A., Colorado, G. H.: Comparativa Entre Paquetes de Software de Apoyo a la Enseñanza de la Programación en Java
8. Berón, M., Henriques, P. R., Varanda Pereira, M. J., Uzal, R.: Herramientas para la comprensión de programas. In: VIII Workshop de Investigadores en Ciencias de la Computación (2006)
9. Moroni, N., Señas, P.: SVED: Sistema de visualización de algoritmos. In: VIII Congreso Argentino de Ciencias de la Computación (2002)
10. BlueJ. Disponible en: <http://bluej.org>
11. Jeliot. Disponible en: <http://cs.joensuu.fi/jeliot/>
12. jGRASP. Disponible en: <http://jgrasp.org>
13. Scratch. Disponible en: <http://scratch.mit.edu>
14. Scitools. Disponible en: <http://scitools.com>
15. CodeSurfer. Disponible en: <https://www.grammatech.com/products/codesurfer>
16. Imagix. Disponible en: <http://www.imagix.com/products/source-code-analysis.html>
17. Shrimp. Disponible en: <http://www.thechiselgroup.com/shrimp/>
18. Alma. Disponible en: <http://epl.di.uminho.pt/~gepl/ALMA/>

Análisis del logro académico de estudiantes en el nivel medio superior a través de minería de datos centrada en el usuario

Gabriel Maldonado¹, Guillermo Molero-Castillo², José Rojano-Cáceres³,
Alejandro Velázquez-Mena⁴

¹ Universidad Veracruzana, MSICU, Facultad de Estadística e Informática, México

² Universidad Veracruzana, CONACYT, México

³ Universidad Veracruzana, Facultad de Estadística e Informática, México

⁴ Universidad Nacional Autónoma de México, Facultad de Ingeniería, México

`gabriel.amhx@gmail.com, ggmoleroca@conacyt.mx,`
`rrojano@uv.mx, mena@fi-b.unam.mx`

Resumen. Existe la necesidad natural de buscar nuevas formas de analizar y procesar datos de diferentes fuentes. Una de estas formas es mediante una minería de datos centrada en el usuario. Se analizó el logro académico de estudiantes en el nivel medio superior del país a través de un algoritmo de agrupamiento particional. Se observó variados niveles de dominio, destacando Insuficiente y Elemental en más del 70% de la población evaluada, mientras que Bueno y Excelente fueron alcanzados por un número reducido de escuelas, 20%. Esto contrasta una notable diferencia entre los niveles alcanzados por los estudiantes, trayendo como consecuencia que éstos retrasen o detengan sus estudios universitarios debido a que obtienen un certificado sin tener los conocimientos para aprobar los exámenes de ingreso en las universidades.

Palabras Clave: DCU, logro académico, minería de datos, PLANEA.

Analysis of the Academic Achievement of Students at the High School Level through User-Centered Data Mining

Abstract. There is a natural need to find new ways to analyze and process data from different sources. One of these ways is through a user-centered data mining. The academic achievement of students at the high school level of the country was analyzed through a partitional clustering algorithm. Various

proficiency levels were observed, highlighting Inadequate and Elemental more than 70% of the study population, while Good and Excellent were hit by a small number of schools, 20%. This contrasts a notable difference between the levels achieved by students, bringing as a result they delay or stop their university studies because they receive a certificate without having the knowledge required to pass the entrance exams at universities.

Keywords. UCD, academic achievement, data mining, PLANEA.

1. Introducción

En la actualidad, debido al crecimiento de la recolección de datos y la evolución del poder de cómputo, se tiene información almacenada en diferentes fuentes. Esto permite contar con datos históricos útiles para explicar el pasado, entender el presente y predecir situaciones futuras [1]. Por tanto, es cada vez mayor la necesidad de buscar nuevas formas de analizar y procesar las fuentes de datos existentes para obtener información y conocimiento útil. Sin embargo, el volumen de información que alcanzan estas fuentes, es con frecuencia una limitante para el análisis de forma manual, por lo que se han desarrollado tecnologías especializadas que permiten procesar y obtener información de interés con el propósito de servir de apoyo en el proceso de la toma de decisiones [2, 3]. Precisamente, una de estas tecnologías es la minería de datos, la cual es un área de la ciencia de la computación que permite realizar el análisis inteligente de datos, que tiene como propósito resolver dos grandes retos [4]: a) trabajar con conjuntos de datos para extraer y descubrir información de interés, y b) usar técnicas adecuadas para analizar e identificar tendencias y comportamientos que faciliten una mejor comprensión de los fenómenos que ocurren en el entorno y sirvan de ayuda en el proceso de la toma de decisiones.

Para hacer minería de datos en la actualidad se tienen variados procesos que guían la planeación y desarrollo de los proyectos. Sin embargo, estos procesos, a pesar de su amplia variedad, tienen una reducida participación del usuario en cada una de sus etapas, trayendo como consecuencia desarrollos de minería de datos con limitaciones de usabilidad y accesibilidad [5]. Oportunamente, en los últimos años el diseño centrado en el usuario logra un importante posicionamiento en el desarrollo tecnológico a través de métodos que se centran en el diseño y la participación de los usuarios [6]. Algunas de las características del diseño centrado en el usuario son [7]: a) facilidad de aprendizaje, que mide qué tan fácil le resulta al usuario hacer una tarea; b) eficiencia, que mide qué tanto ha aprendido el usuario acerca del funcionamiento básico del software; c) cualidad de ser recordado, que mide con qué facilidad recuerdan los usuarios el uso del sistema; d) errores, que mide con qué facilidad los usuarios cometen errores y con qué facilidad los pueden resolver; y e) satisfacción, que mide qué tan sencillo le resultó al usuario utilizar el sistema. Además, el diseño centrado en el usuario es ampliamente utilizado en distintas aplicaciones, como: diseño de equipos médicos, creación de vehículos, productos de consumo, entre otros. Recientemente también se está proyectando su uso en la minería de datos, logrando con esto resultados con mayor satisfacción en los usuarios [8].

En este trabajo se presenta los resultados de la proyección de una minería de datos centrada en el usuario para analizar el logro académico de los estudiantes del nivel medio superior en el país. Para esto se utilizó como fuente de datos los registros del Plan Nacional para la Evaluación de los Aprendizajes (PLANEA), específicamente datos de escuelas del nivel medio superior, de carácter público, federal y estatal, y planteles particulares reconocidos por la Secretaría de Educación Pública (SEP). PLANEA tiene como propósito conocer en qué medida los estudiantes logran dominar aprendizajes esenciales en Lenguaje y Comunicación (comprensión lectora) y Matemáticas al término de la educación media superior. Los patrones de datos obtenidos son útiles como instrumento de información para padres de familia, estudiantes, maestros, directivos, autoridades educativas y sociedad en general.

2. Procesos actuales de minería de datos

En la actualidad, dada la amplia variedad de técnicas y herramientas existentes para llevar a cabo el análisis de datos, se requiere de procesos específicos para el desarrollo de proyectos de minería de datos [9]. Uno de los procesos utilizados es el KDD (Knowledge Discovery in Databases), que proporciona fases estructuradas de manera secuencial e iterativa. Sin embargo, este proceso no provee una descripción ampliada de las actividades que deben realizarse en cada una de sus fases, dejando la definición de cada una de éstas a criterio del equipo de trabajo [10]. Otro de los procesos es SEMMA (Sample, Explore, Modify, Model, and Assess) que se enfoca en aspectos técnicos del proyecto dejando de lado la fase de análisis y comprensión del negocio [11], teniendo como objetivos [12]: segmentar grupos de clientes, personalizar la gestión de relaciones, identificar clientes rentables e identificar clientes que se van con la competencia. SEMMA al igual que KDD no describen las actividades específicas que se deben realizar en cada una de éstas [10, 13]. En la Tabla 1 se muestra una síntesis de las principales características de éstos y otros procesos descritos.

Otro proceso es CRISP-DM (Cross Industry Standard Process for Data Mining) que es la guía actual más utilizada en el desarrollo de proyectos de minería de datos [14]. Este proceso sirve para guiar el desarrollo de un proyecto de a través de una secuencia dividida en fases [15]. Estas fases siguen un proceso iterativo, que es útil para realizar cambios en etapas anteriores. Sin embargo, no incluye tareas de control y monitoreo del plan de trabajo [11]. Por su parte, Microsoft propone soluciones para usuarios empresariales y no para usuarios finales, propiciando con esto que en ocasiones se obtenga aplicaciones no ergonómicas y con dificultad de uso [16]. Además, está vinculada a herramientas comerciales como SQL Server. Otros de los procesos utilizados es DWEP (Data warehouse engineering process), que basa su funcionamiento en el proceso unificado de desarrollo de software, abarcando desde la recolección de requerimientos hasta la fase de pruebas del producto final. Sin embargo, DWEP al ser un proceso orientado a ingeniería de software no comparte todas las fases que requiere la minería de datos, salvo en el análisis de datos y del problema [17].

Así, a pesar de la variedad de estos procesos, que han tenido que adaptarse a las necesidades de los proyectos y usuarios, éstos presentan algunas debilidades sobre todo en la fase de análisis del problema o del negocio, donde por lo general los requisitos del proyecto y necesidades de los usuarios no son recuperados adecuadamente, trayendo como consecuencia desarrollos de minería de datos con limitaciones de usabilidad y accesibilidad [18]. Estas necesidades se traducen en puntos de vista, preferencias, estrategias y decisiones de los usuarios [5].

Tabla 1. Principales características de los procesos de la minería de datos.

	KDD	CRISP-DM	SEMMA	DWEP	Microsoft
Autor	Fayyad <i>et al.</i> (1996)	SPSS, Daimler Chrysler y NCR (1996)	SAS Inc.	Lujan-Morán (2005)	Microsoft
Área de aplicación	Academia / Industria	Academia / Industria	Industria	Industria	Industria
Objetivos	Se centra en la comprensión del dominio de aplicación	Se centra en los objetivos empresariales del proyecto	Se centra en las características técnicas del desarrollo del proceso	Se centra en los flujos de trabajo para la integración de datos	Se centra en los objetivos empresariales del proyecto
Estructura	Fases	Fases y niveles	Fases	Fases y flujos de trabajo	Fases
Número de fases	5	6	5	4	6
Fases	Integración y recopilación de datos Selección, limpieza y transformación Minería de datos Evaluación e interpretación Uso del conocimiento	Comprensión del negocio Comprensión de los datos Preparación de los datos Modelado Evaluación Despliegue	Muestreo Exploración Modificación Modelado Evaluación	Inicio Elaboración Construcción Transición	Definir el problema Preparar los datos Explorar los datos Generar modelos Validar modelos Implementar y actualizar modelos
Herramientas	Libres y comerciales	Libres y comerciales	SAS	Libres y comerciales	Microsoft
Iteración	Si	Si	No	Si	Si

3. Minería de datos centrada en el usuario

La participación del usuario no sólo beneficia la obtención de mejores requerimientos, sino también la creación de proyectos de minería de datos personalizados, es decir, hechos a la medida [18, 19]. Además, centrarse en el usuario no significa enfocarse en uno solo, sino se debe tomar en cuenta a todos los involucrados en el proyecto, considerando las edades, capacidades, rasgos, diferencias y otras características de interés [20]. Precisamente, una de las disciplinas encargadas de recopilar las necesidades del usuario e incorporarlas en el producto final es el Diseño Centrado en el Usuario (DCU), que es ampliamente utilizado para el diseño y desarrollo de proyectos tecnológicos [6]. Un enfoque actual del DCU es el Modelo de Proceso de la Ingeniería y de la Accesibilidad (MPIu+a) definido por [21]. Este proceso no sólo incluye la usabilidad, sino también agrega el concepto de accesibilidad, permitiendo con esto el acceso a distintos tipos de usuario. Como pilares de este proceso destacan: a) la ingeniería de software, que sigue un ciclo de desarrollo en cascada (análisis, diseño, desarrollo y transición); b) el prototipado,

utilizado para explorar posibles mejoras acerca del diseño; c) la evaluación, que engloba y categoriza técnicas de evaluación existentes, como pruebas de caja blanca o caja negra; y d) usuario, que es el principal actor de todo el proceso y que interviene en cada una de sus etapas.

Una característica de MPIu+a es la usabilidad, que es la facilidad de uso de una aplicación interactiva [22]. Sin embargo, ésta sólo es una definición operativa, puesto que se logra a través de varias condiciones [7]: a) facilidad de aprendizaje, b) eficiencia, c) cualidad de ser recordado, d) errores, y e) satisfacción.

En consecuencia, existe la natural necesidad de definir un proceso de minería de datos que incluya el diseño centrado en el usuario. Mediante este tipo de procesos se diseñan e implementan aplicaciones personalizadas capaces de realizar el reconocimiento de patrones de datos. En este trabajo, se incluye como base procesos ampliamente utilizados por la comunidad científica: a) CRISP-DM y b) MPIu+a. Ambos aportan fundamentos significativos que de forma conjunta permiten definir etapas y acciones concretas para llevar a cabo proyectos de reconocimiento de patrones a través de la minería de datos centrada en el usuario. Como proyección general en la Fig. 1 se presenta la minería de datos centrada en el usuario, estructurada en cuatro etapas: a) análisis del problema, b) análisis y preparación de datos, c) adquisición de patrones, y d) presentación de patrones. Además, como parte de las etapas intervienen los usuarios, el prototipado y la evaluación.

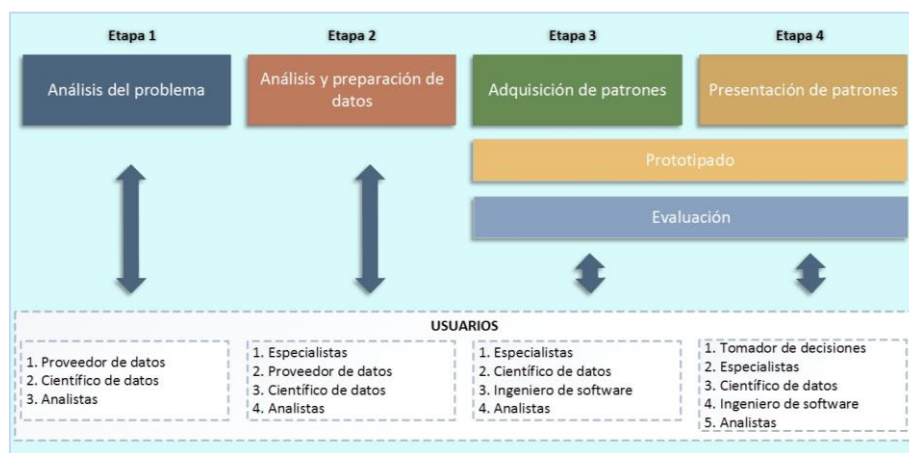


Fig. 1. Proyección general del proceso de minería de datos centrado en el usuario.

En la Fig. 2 se muestra las etapas y tareas generales del proceso. Tiene como característica principal la captura de requerimientos y necesidades de los usuarios para la construcción de interfaces personalizadas. Las primeras tres etapas corresponden al descubrimiento de patrones, mediante las cuales se hace un análisis del problema, análisis y preparación de los datos, así como la extracción de patrones significativos. Esta adquisición de patrones (Etapa 3) junto con la presentación (Etapa 4) representan la columna vertebral del diseño centrado en el usuario, en las cuales se definen

acciones específicas para orientar el desarrollo de aplicaciones personalizadas de minería de datos (software Ad hoc). Entre los tipos de usuarios identificados destacan: el proveedor de información, los especialistas del negocio, el tomador de decisiones, así como el equipo de desarrollo (analistas, científico de datos, ingeniero de software, entre otros). Además, como parte de proceso se incluyen también las fases de prototipado y evaluación, siendo la primera utilizada para crear representaciones del producto final, mientras que la segunda permite evaluar la calidad del producto obtenido.

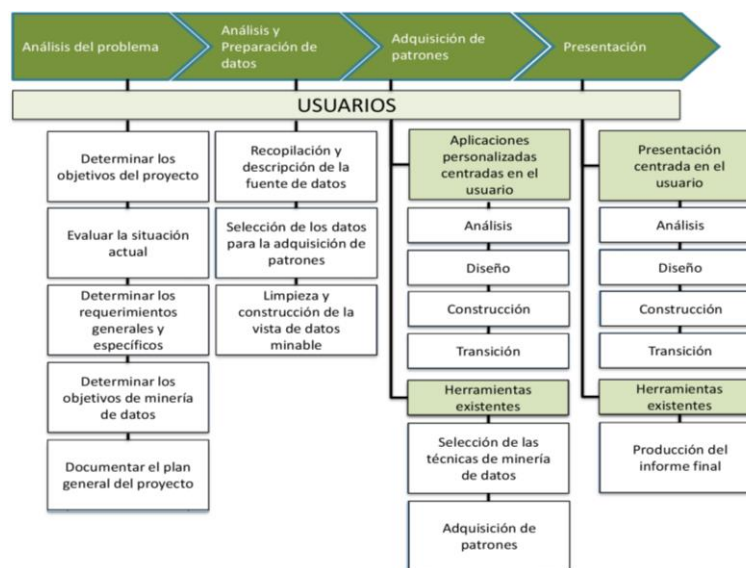


Fig. 2. Etapas y tareas del proceso de minería de datos centrado en el usuario.

Este proceso fue tomado como base para el análisis de grupos asociado al logro académico de estudiantes en el nivel medio superior. La fuente de datos utilizada corresponde al Plan Nacional para la Evaluación de los Aprendizajes (PLANEA).

4. Logro académico en el nivel medio superior

En la actualidad, la educación es uno de los pilares para el desarrollo social y económico de un país. Para obtener resultados satisfactorios se necesita de una educación de calidad; la cual se logra a través de los sistemas educativos que juegan un papel decisivo en el mejoramiento de la calidad educativa [23]. Así, al ser el logro académico un parámetro importante de medición sobre la calidad de la enseñanza, que proporcionan los sistemas educativos, surge el interés de conocer en qué medida los estudiantes alcanzan aprendizajes esenciales en diferentes dominios al término de cada nivel educativo, con el propósito de hacer un diagnóstico del desempeño y de los conocimientos alcanzados por los estudiantes.

Actualmente uno de estos diagnósticos se hace a través del Plan Nacional para la Evaluación de los Aprendizajes en el nivel medio superior (PLANEA-MS) de la Secretaría de Educación Pública [24, 25]. Por tanto, hacer el análisis del logro académico de estudiantes de nivel medio superior sirve de apoyo en la toma de decisiones del ámbito educativo para mejorar la calidad del desempeño académico en los estudiantes [26]. En el caso de la educación del nivel medio superior existen casos de alumnos que al término de sus estudios no obtienen los conocimientos necesarios para aprobar los exámenes de ingreso a las universidades del país, trayendo como consecuencia que retrasen o detengan sus estudios universitarios.

La evaluación que realiza PLANEA-MS está dirigido a los alumnos de toda la República Mexicana que cursan el último ciclo escolar, inscritos a algún plantel educativo, ya sea autónomo, estatal, federal o particular. Las áreas de competencia que evalúa son Lenguaje y Matemáticas. La primera evaluación de PLANEA-MS se realizó en marzo de 2015, participaron más de un millón de alumnos de educación media superior, de un total de 14548 instituciones, entre escuelas públicas, autónomas y privadas. Los logros académicos están comprendidos en cuatro niveles de dominio: a) I, insuficiente; b) II, elemental; c) III, bueno; y d) IV, excelente.

Derivado del análisis de datos de PLANEA, se obtuvo una vista de datos minable. La principal consideración fue determinar cuántas y cuáles son las variables apropiadas para el estudio. En la Tabla 2 se presenta las variables significativas de la vista de datos minable.

Tabla 2. Variables que conforman la vista de datos minable.

No.	Variable	Descripción	Tipo
1	Escuela	Nombre de la escuela	Nominal
2	Turno	Turno	Continuo
3	Entidad	Entidad	Continuo
4	Subsistema	Subsistema	Continuo
5	Sostenimiento	Sostenimiento	Continuo
6	Alumnos evaluados	Alumnos evaluados	Continuo
7	Contestaron_50+_lenguaje	50% o más de preguntas contestadas en lenguaje	Continuo
8	Contestaron_50+_matematicas	50% o más de preguntas contestadas en matemáticas	Continuo
9	P1_lenguaje_dominio_I	% de alumnos evaluados en lenguaje (Nivel I)	Continuo
10	P1_lenguaje_dominio_II	% de alumnos evaluados en lenguaje (Nivel II)	Continuo
11	P1_lenguaje_dominio_III	% de alumnos evaluados en lenguaje (Nivel III)	Continuo
12	P1_lenguaje_dominio_IV	% de alumnos evaluados en lenguaje (Nivel IV)	Continuo
13	P1_matematicas_dominio_I	% de alumnos evaluados en matemáticas (Nivel I)	Continuo
14	P1_matematicas_dominio_II	% de alumnos evaluados en matemáticas (Nivel II)	Continuo
15	P1_matematicas_dominio_III	% de alumnos evaluados en matemáticas (Nivel III)	Continuo
16	P1_matematicas_dominio_IV	% de alumnos evaluados en matemáticas (Nivel IV)	Continuo

Los requerimientos principales para la construcción de la aplicación fueron: a) tener una sección para generar la vista de datos minable, b) incluir un algoritmo particional tipo K-means, c) incluir el método del codo (elbow method) para analizar el número deseado de grupos, d) permitir al usuario crear un proyecto de minería de datos que sea dinámico y de fácil uso, y e) tener opciones de ayuda para retroalimentar al usuario sobre el funcionamiento de la herramienta. Como resultado en la Fig. 3 se muestra la pantalla principal de la aplicación para la adquisición de patrones centrada en el usuario. Los módulos implementados son: a) origen de datos, b) análisis de datos, c) selección de variables, d) minería de datos, y e) validación. El

algoritmo K-means tiene como particularidad establecer a priori el número de grupos de entrada; por lo que se implementó el método del codo (elbow method) para obtener el número deseado de grupos. Este método permite identificar los grupos a través de una representación visual [27].

Para el análisis del logro académico se identificó 6 como número deseado de grupos (Fig. 4). Esta consideración fue con base en la variada información de los subsistemas educativos, sostenibilidad y los niveles de dominio del logro académico.

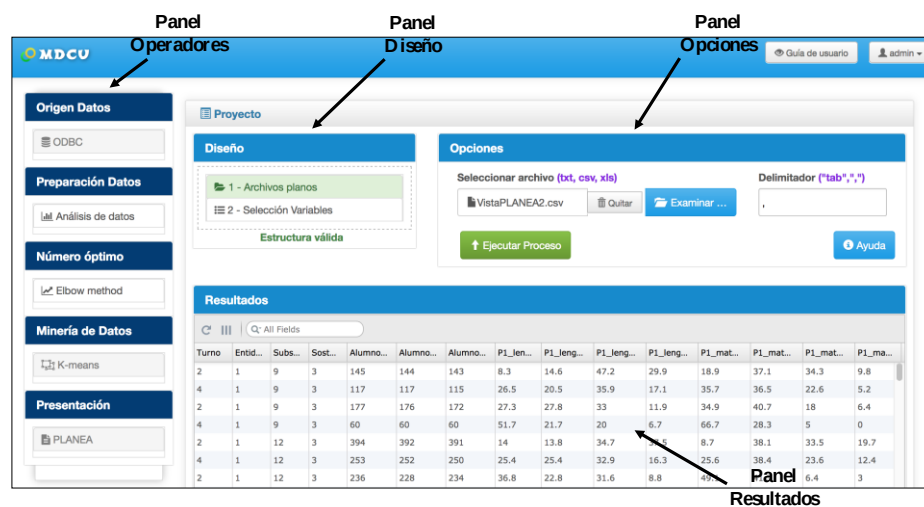


Fig. 3. Pantalla principal de la aplicación.

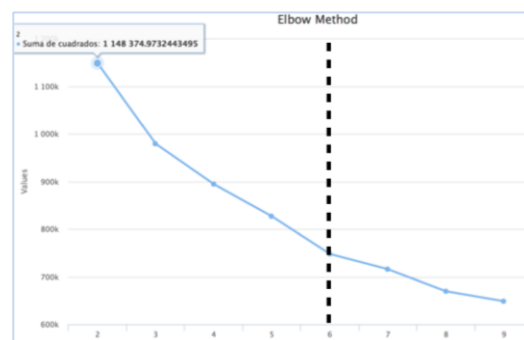


Fig. 4. Pantalla de la obtención del número deseado de grupos.

En la Fig. 5 se presenta un resumen de los grupos obtenidos por K-means. Como resultado se observó que el Grupo 1 alcanzó un mejor logro académico en Lenguaje en el nivel III (Bueno, 32%) y II en Matemáticas (Elemental, 39%). El Grupo 2 obtuvo un mayor logro académico para Lenguaje en el nivel I (insuficiente, 41%), mientras que en Matemáticas fue en el nivel I (insuficiente, 49%). El Grupo 3 destacó

por tener logros bajos, nivel I, insuficiente en Lenguaje (50%) e insuficiente en Matemáticas (59%). El Grupo 4 refleja el logro académico más bajo, con niveles Insuficientes en Lenguaje y Matemáticas con 64% y 78%, respectivamente. El Grupo 5 destacó por alcanzar altos porcentajes de insuficiencia en Lenguaje (39%) y Matemáticas (45%). Finalmente, el Grupo 6 obtuvo en Lenguaje niveles Bueno (30%) e Insuficiente (29%), mientras que en Matemáticas resaltan los niveles Insuficiente (35%) y Elemental (34%). Estos resultados contrastan que existen alumnos que al finalizar sus estudios obtienen el certificado sin tener los conocimientos necesarios para posteriormente aprobar los exámenes de ingreso en las universidades del país. Trayendo como consecuencia que retrasen o detengan sus estudios universitarios.

Grupo	Resumen	Grupo	Resumen
1	No. Instituciones: 3375 Promedio de alumnos evaluados: 26 Dominio Lenguaje y Comunicación: Nivel I (Insuficiente): 27% (911 de 3375) Nivel II (Elemental): 23% (776 de 3375) Nivel III (Bueno): 32% (1080 de 3375) Nivel IV (Excelente): 18% (608 de 3375) Dominio Matemáticas: Nivel I (Insuficiente): 32% (1080 de 3375) Nivel II (Elemental): 39% (1316 de 3375) Nivel III (Bueno): 19% (641 de 3375) Nivel IV (Excelente): 10% (338 de 3375)	2	No. Instituciones: 1727 Promedio de alumnos evaluados: 192 Dominio Lenguaje y Comunicación: Nivel I (Insuficiente): 41% (708 de 1727) Nivel II (Elemental): 21% (363 de 1727) Nivel III (Bueno): 26% (449 de 1727) Nivel IV (Excelente): 12% (207 de 1727) Dominio Matemáticas: Nivel I (Insuficiente): 49% (846 de 1727) Nivel II (Elemental): 32% (553 de 1727) Nivel III (Bueno): 13% (225 de 1727) Nivel IV (Excelente): 6% (103 de 1727)
3	No. Instituciones: 3132 Promedio de alumnos evaluados: 93 Dominio Lenguaje y Comunicación: Nivel I (Insuficiente): 50% (1566 de 3132) Nivel II (Elemental): 20% (626 de 3132) Nivel III (Bueno): 21% (658 de 3132) Nivel IV (Excelente): 9% (282 de 3132) Dominio Matemáticas: Nivel I (Insuficiente): 59% (1848 de 3132) Nivel II (Elemental): 27% (846 de 3132) Nivel III (Bueno): 10% (313 de 3132) Nivel IV (Excelente): 4% (125 de 3132)	4	No. Instituciones: 5816 Promedio de alumnos evaluados: 23 Dominio Lenguaje y Comunicación: Nivel I (Insuficiente): 64% (3722 de 5816) Nivel II (Elemental): 18% (1047 de 5816) Nivel III (Bueno): 14% (814 de 5816) Nivel IV (Excelente): 4% (233 de 5816) Dominio Matemáticas: Nivel I (Insuficiente): 78% (4537 de 5816) Nivel II (Elemental): 18% (1047 de 5816) Nivel III (Bueno): 3% (174 de 5816) Nivel IV (Excelente): 1% (58 de 5816)
5	No. Instituciones: 30 Promedio de alumnos evaluados: 969 Dominio Lenguaje y Comunicación: Nivel I (Insuficiente): 39% (12 de 30) Nivel II (Elemental): 21% (6 de 30) Nivel III (Bueno): 26% (8 de 30) Nivel IV (Excelente): 14% (4 de 30) Dominio Matemáticas: Nivel I (Insuficiente): 45% (14 de 30) Nivel II (Elemental): 34% (10 de 30) Nivel III (Bueno): 14% (4 de 30) Nivel IV (Excelente): 7% (2 de 30)	6	No. Instituciones: 459 Promedio de alumnos evaluados: 356 Dominio Lenguaje y Comunicación: Nivel I (Insuficiente): 29% (133 de 459) Nivel II (Elemental): 21% (96 de 459) Nivel III (Bueno): 30% (138 de 459) Nivel IV (Excelente): 20% (92 de 459) Dominio Matemáticas: Nivel I (Insuficiente): 35% (161 de 459) Nivel II (Elemental): 34% (156 de 459) Nivel III (Bueno): 19% (87 de 459) Nivel IV (Excelente): 12% (55 de 459)

Fig. 5. Resumen de los grupos obtenidos por el algoritmo K-means.

Para la presentación de los resultados obtenidos se realizaron sesiones de trabajo con especialistas en educación. Entre los requerimientos levantados destacan: a) los resultados deben mostrarse a nivel federal y estatal, b) el usuario debe pasar de un nivel a otro sin restricciones, y c) incluir gráficas para visualizar las variables representativas de PLANEA. Así, se incluyó una variedad de gráficas para tener un

mayor entendimiento sobre el logro académico registrado en PLANEA 2015. En la Fig. 6 se presenta un extracto de la presentación de patrones en el nivel estatal.

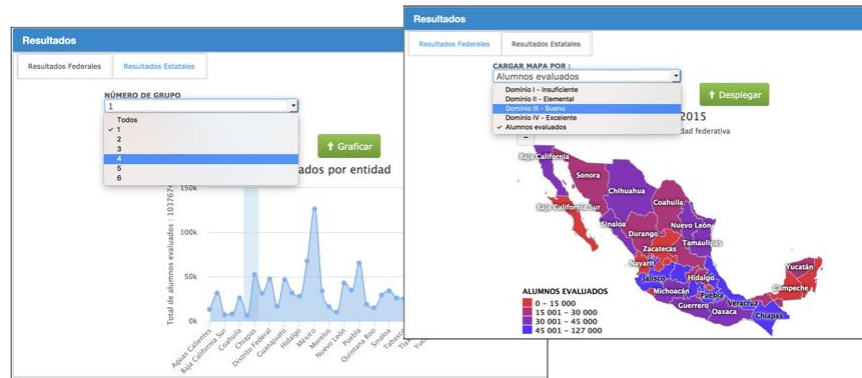


Fig. 6. Módulo de presentación de patrones sobre la fuente de datos PLANEA 2015.

Al dar clic en el mapa de la república mexicana se despliega una pantalla emergente con información de los grupos obtenidos mediante gráficas, dividido por variables, como: sostenimiento, subsistema y logro académico (lenguaje y comunicación, y matemáticas). En la Fig. 7 se visualiza también que el mayor número de escuelas en Veracruz pertenecen al subsistema de Tele Bachillerato, con 1022 instituciones. Con respecto al sostenimiento educativo, se observa que la mayor cantidad de instituciones son del tipo estatal, con 1212 escuelas, mientras que otro menor número están distribuidos entre particulares y federales.



Fig. 7. Distribución de las escuelas de acuerdo al tipo de subsistema educativo.

Las pruebas de usabilidad sobre la aplicación permitieron detectar mejoras en la interfaz para optimizar la experiencia del usuario. Estas pruebas se dividieron en dos etapas: a) con ocho usuarios con conocimientos de minería de datos para analizar la aplicación de adquisición de patrones, y b) con cuatro usuarios especialistas en educación para analizar el módulo de presentación de patrones. En la Fig. 8 se muestran los resultados obtenidos.

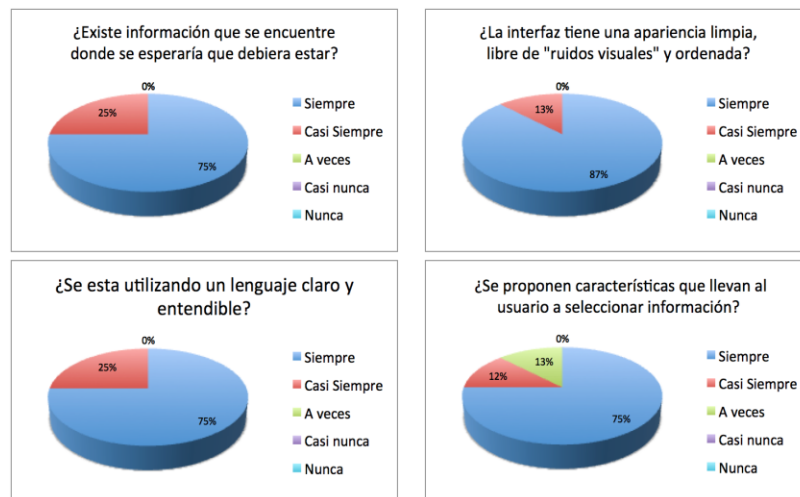


Fig. 8. Resumen resultados pruebas de usabilidad.

En general, las pruebas fueron favorecedoras, demostrando que la mayoría de los usuarios encontraron las cosas donde ellos esperaban, también mostraron agrado en el lenguaje en el que se les presentó la información, así mismo la mayoría coincidió en que la apariencia de la interfaz de usuario fue limpia y libre de ruidos visuales.

5. Conclusiones

Se realizó un análisis del logro académico de estudiantes en el nivel medio superior, tomando como base una minería de datos centrado en el usuario. Se utilizó como fuente de información la base de datos del Plan Nacional para la Evaluación de los Aprendizajes (PLANEA).

Resultado de los grupos obtenidos tras la aplicación del algoritmo K-means, se observó que existe una deficiencia en el logro académico alcanzado por los estudiantes y por tanto por las instituciones educativas, trayendo como consecuencia que estos no adquieran los conocimientos necesarios para continuar sus estudios superiores.

Se comprobó que el método del codo fue útil para detectar el número deseado de grupos en el algoritmo K-means. Seis fueron los grupos obtenidos, destacando los niveles insuficiente y elemental, tanto para matemáticas y lenguaje.

Las pruebas de usabilidad realizadas permitieron identificar mejoras sobre en la aplicación, destacando que la interfaz tiene una apariencia limpia y libre de ruidos visuales y que el lenguaje utilizado es claro y entendible.

Respecto a la satisfacción por parte de los usuarios que participaron en las pruebas de usabilidad, reflejaron resultados positivos, mostrando que la intervención del usuario en todo momento del proyecto ayuda a entender mejor el problema.

Agradecimiento. Este trabajo forma parte del proyecto “Infraestructura para agilizar el desarrollo de sistemas centrados en el usuario” financiado por el Consejo Nacional de Ciencia y Tecnología, en el marco de Cátedras CONACYT (Ref. 3053).

Referencias

1. López, C.: Minería de datos: técnicas y herramientas. Editorial Paraninfo (2007)
2. Hernández J., Ramírez M. J., Ferri C.: Introducción a la Minería de Datos. Pearson Educación, Editorial Pearson Prentice Hall, Madrid, España (2004)
3. Witten, I., Frank, E.: Data Mining: Practical machine learning tools and techniques. Morgan Kaufmann Publishers (2005)
4. Molero G.: Clasificador bayesiano para el pronóstico de la supervivencia y mortalidad de casos de cáncer de mama en mujeres de origen hispano (Tesis doctoral). Universidad de Guadalajara, México (2014)
5. Zhao, Y., Chen, Y., Yao, Y.: User-centered interactive data mining. In: Cognitive Informatics, 5th IEEE International Conference, 457–466 (2006)
6. Abras, C.; Maloney-Krichmar, D.; Preece, J.: User-centered design. Bainbridge, W. Encyclopedia of Human-Computer Interaction. Thousand Oaks: Sage Publications (2004)
7. Nielsen, J.: Usability engineering. Academic Press Limited, Massachuetts, Estados Unidos (1993)
8. Horberry, T., Burgess-Limerick, R., Steiner, L.: Human Centred Design for Mining Equipment and New Technology. In: Proceedings 19th Triennial Congress of the IEE, 9, 14 (2015)
9. Sumathi, S., Sivanandam, S.: Introduction to Data Mining and its Applications. Studies in Computational Intelligence, 29, editado por Springer-Verlag, Heidelberg, Alemania (2006)
10. Moine, J., Gordillo, S., y Haedo, A.: Análisis comparativo de metodologías para la gestión de proyectos de Minería de Datos. In: Presentado en el VIII Workshop Bases de Datos y Minería de Datos, 931–938 (2011)
11. Peralta, F.: Elementos para un mapa de actividades para proyectos de explotación de información. Facultad Regional Buenos Aires, Argentina (2013)
12. Vanrell, J., Bertone, R.: Modelo de Proceso de Operación para Proyectos de Explotación de Información. In: XVI Congreso Argentino de Ciencias de la Computación, 674-682 (2010)
13. Moine, J.: Metodologías para el descubrimiento de conocimiento en bases de datos: un estudio comparativo (Tesis doctoral). Universidad Nacional de la Plata, Buenos Aires, Argentina (2013)
14. KDnuggets.: Data Mining, Analytics, Big Data, and Data Science. www.kdnuggets.com/2014/10/crisp-dm-top-methodology-analytics-data-mining-data-science-projects.html
15. Chapman, P., Clinton, J., Kerber, R., Khabaza, T., Reinartz, T., Shearer, C., Wirth, R.: CRISP-DM 1.0 Step-by-step Data Mining Guide (2000)
16. Flores, H.: Detección de Patrones de Daños y Averías en la Industria Automotriz (Tesis de Maestría). Universidad Tecnológica Nacional (2009)
17. Luján-Mora, S.: Data warehouse design with UML (Tesis Doctoral). Universidad de Alicante, España (2005)
18. Ho, T., Nguyen, T., Nguyen, D.: Visualization support for a user-centered KDD process. In: Proceedings of the eighth ACM SIGKDD international conference on Knowledge discovery and data mining, Japan Advanced Institute of Science and Technology Tatsunokuchi, Ishikawa, 519–524 (2002)

19. Ho, T., Nguyen, T., Nguyen, D.: A User-Centered Visual Approach to Data Mining, The system D2MS, Estados Unidos, 213–224 (2002)
20. Martín, A.: MPIu+a Ágil: El modelo de proceso centrado en el usuario como metodología ágil, Universidad de Lleida (2010)
21. Granollers, T.: MPIu+ a. Una metodología que integra la Ingeniería del Software, la Interacción Persona-Ordenador y la Accesibilidad en el contexto de equipos de desarrollo multidisciplinares, Universidad de Lleida (2004)
22. Hassan, Y., Ortega, S.: Informe APEI sobre usabilidad, España (2009)
23. Arnaut, A., Giorguli, S.: Los grandes problemas de México, El colegio de México, Ciudad de México (2010)
24. SEP.: Plan nacional para la evaluación de los aprendizajes (PLANEA), www.inee.edu.mx/index.php/planea
25. PLANEA-MS: Manual para usuarios, Plan Nacional para las Evaluaciones de los Aprendizajes en el nivel medio superior (2015)
26. De Ibarrola, M.: Los grandes problemas del sistema educativo mexicano, Perfiles educativos, 34 (SPE), 16–28 (2012)
27. Kodinariya, T., Makwana, P.: Review on determining number of Cluster in K-Means Clustering, International Journal, 1(6), 90–95 (2013)

Análisis basados en n-gramas para la clasificación del nivel de Burnout en académicos universitarios

Mauricio Castro, David Pinto, Darnes Vilariño, Mireya Tovar

Benémerita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

mayka2000@yahoo.com.mx,
{davideduardopinto,dvilarinoayala,mireyatovar}@gmail.com
<http://www.lke.buap.mx/>

Resumen. En este artículo se presenta un estudio del nivel de estrés sobre la base de los comentarios abiertos emitidos por académicos universitarios usando reconocimiento del lenguaje natural. El enfoque presenta la aplicación de un algoritmo de clasificación automática y el análisis de características para dicho clasificador basadas en n-gramas de palabras. Los resultados de dos experimentos son mostrados: 1) el impacto en el incremento en el tamaño del n -grama en el proceso de clasificación, y 2) la verificación en la separación de dos clases, una con bajo y la otra con alto nivel de Burnout. Los resultados obtenidos muestran un porcentaje de clasificación correcta cercano al 60% cuando se consideran tres clases: alto, medio y bajo nivel de Burnout, mientras que dicho porcentaje se incrementa a 79% cuando únicamente se consideran dos clases: alto y bajo nivel de Burnout.

Palabras clave: Clasificación automática, nivel de Burnout, N-gramas de palabras.

N-gram based Analysis for Classifying the Burnout Level of University Professors

Abstract. In this paper we present a study of stress level using a set of open comments written by university professors employing natural language processing techniques. The approach uses a classification algorithm and feature analysis based on word n-grams. We show the results obtained in two experiments: 1) The impact of increasing the n-gram size in the classification process, and 2) An analysis of the degree of separation of two classes: one with low and another with high level of Burnout. The obtained results show a classification accuracy close to 60% when using three classes (high, medium and low level of Burnout), whereas, that percentage increases to 79% when we employed only two classes (high and low level of Burnout).

Keywords. Machine learning, Burnout level, word N-grams.

1. Introducción

El término estrés fue inicialmente acuñado por Hans Selye en 1936 al experimentar con ratas expuestas a diferentes condiciones externas nocivas para ellas. Selye le llamó “Stress” a la reacción de las ratas a esas condiciones externas, y describió la conducta de las ratas en su trabajo [6].

En su momento, el estudio del estrés se reducía a establecer los cambios fisiológicos que sufría el ser humano expuesto a amenazas del medio. Sin embargo, desde finales de la década de los 60, el estrés se asocia a amenazas de carácter organizacional. Freudenberg [1] estudia la conducta de personas con altos niveles de estrés y que fueron expuestas a duros retos laborales, llamándole “Síndrome de Burnout” a este padecimiento. Freudenberg describe este padecimiento como sentimientos de cansancio emocional y frustración asociados a una carga irracional de trabajo.

En el año de 1976, Cristina Maslash define al Síndrome de Burnout como cansancio Emocional, despersonalización y baja realización profesional, y de esta forma se abren las perspectivas de estudios cuantitativos que culminan para el año 1981 con la creación de un instrumento que está compuesto de 22 ítems organizados de la siguiente forma: 9 ítems relacionados con la dimensión de cansancio emocional, 5 con la dimensión de despersonalización y 8 con la dimensión de realización profesional. A este instrumento se le ha llamado “Maslash Burnout Inventory” [2,3] (en adelante MBI) y fue validado por medio del análisis factorial con una muestra de 1,025 personas de diferentes ámbitos laborales pero todos con profesiones relacionadas con el área de servicios como son: enfermeras, maestros, trabajadores sociales etc. subrayando las autoras, fuertes correlaciones de los resultados del análisis factorial con el contexto demográfico. Un análisis de las principales fortalezas y debilidades del instrumento se pueden ver en [4].

Paralelamente a los estudios cuantitativos de este síndrome con los instrumentos descritos anteriormente, en la literatura se habla de validaciones externas, es decir de observadores externos que por medio de estudios clínicos determinan si una persona tiene o no el síndrome. Así, en el mismo trabajo [2], se realiza una “validación externa” con ayuda de 40 médicos de la salud mental, los cuales realizaron una evaluación clínica de los encuestados con el instrumento MBI.

La identificación automática del síndrome de Burnout en las personas puede ser de gran beneficio pues permitiría atender este padecimiento y evitar complicaciones a largo plazo en su desarrollo personal y profesional. En este trabajo de investigación, existe el interés de detectar el nivel de Burnout en académicos universitarios utilizando los comentarios de texto escritos por los mismos docentes en un cuestionario sobre su percepción del desempeño y del ambiente laboral.

El resto de este trabajo se encuentra dividido como sigue. En la Sección 2. se presenta la metodología propuesta para el desarrollo de esta investigación. La Sección 3. muestra los datos utilizados, los experimentos, los resultados obtenidos, así como una discusión de dichos resultados. Finalmente, en la Sección 4. se dan las conclusiones de este trabajo y los experimentos planteados a futuro.

2. Metodología

En diciembre del 2012 se aplicó el instrumento “Maslash Burnout Inventory” a 2,943 profesores de diferentes facultades de la Benemérita Universidad Autónoma de Puebla (BUAP), con la intención de buscar las causas del bajo rendimiento que se había observado años atrás en algunos docentes, y las quejas de estudiantes de mal trato de algunos profesores.

Al aplicársele a este instrumento el análisis factorial exploratorio y observar una nueva configuración de las dimensiones que resultaban validadas se generó un nuevo instrumento con 21 elementos, ordenando las dimensiones de la siguiente forma:

- 5 Elementos de la dimensión de cansancio emocional
- 8 Elementos de la dimensión despersonalización
- 8 Elementos de la dimensión realización profesional

Se calculó el Alfa de Gronbach a cada dimensión, variando ésta de entre 0.702 y 0.860. Con este nuevo instrumento se hizo el estudio de 2,943 docentes de la BUAP de un total de 5,395 trabajadores académicos. Los docentes pertenecen a 31 unidades académicas de entre las cuales, 8 unidades son del nivel medio superior (20.3% de la población), el resto pertenecen al nivel superior; de estos el 45.8% son del género femenino. Las edades oscilan entre los 23 y los 80 años con una antigüedad promedio en la BUAP de 12 años y 7.8 fuera de la BUAP. El 80% dice poseer nivel superior a licenciatura, y de estos el 13.7% tienen título de doctor con un 7.9% en proceso de titulación. El 67% tiene hijos y el 17.8% tiene más de 2 hijos. En el año 2012, el promedio de los grupos atendidos es de 32 y de alumnos atendidos 172 (el 52% de la planta docente atiende grupos de más de 30).

Tabla 1. Los resultados obtenidos.

Nivel de Burnout	Número de profesores	Porcentaje
Libre	789	26.8%
Bajo	865	29.4%
Medio	892	30.3%
Alto	350	11.9%
Grave	47	1.6%
Total	2,943	

Como puede verse en la Tabla 1, un 13.5% de docentes manifiestan altos niveles del síndrome de Burnout. Al final del test que se aplicó a los 2,943 docentes se les pidió que voluntariamente dieran un comentario abierto acerca de ellos o de lo que les pareciera negativo o positivo de su ambiente laboral, obteniendo en total 1,944 comentarios abiertos.

Así, el problema planteado en este trabajo de investigación sería:

- Puede detectarse el síndrome de burnout a partir de estos comentarios?, o

- Puede determinarse la no existencia del síndrome a partir de la forma en que se expresan en los comentarios abiertos?

Para este propósito, es necesario establecer una metodología para la ejecución de los experimentos que lleven a solucionar o acercarse a una solución para el problema planteado. Así, se plantea iniciar los experimentos estableciendo un conjunto de datos que emplee solamente dos clases: alto y bajo nivel de Burnout. A partir de dichos datos deseamos verificar si es posible detectar el nivel de Burnout de un académico universitario. Para tal objeto usaremos una técnica basada en aprendizaje automático, la cual considera utilizar parte de los datos como conjunto de entrenamiento y el resto como conjunto de prueba.

La Tabla 1 presenta las cinco clases que fueron usadas en el proceso de evaluación del nivel de Burnout en académicos universitarios mediante un análisis basado en un cuestionario. Son clases discretas que no son del todo excluyentes dado que, por ejemplo, un profesor con nivel grave de Burnout podría considerarse con un nivel alto y por tanto, podríamos unir a estas dos clases en un solo nivel. Así, para los experimentos planteados en este trabajo de investigación se decidió agrupar las cinco clases mostradas en la Tabla 1 en solamente tres clases: “Bajo”, “Medio” y “Alto”. Estas clases se muestran en la Tabla 2 y son el resultado de agrupar las clases “Libre” y “Bajo” en una sola clase llamada “Bajo”, y agrupar las clases “Alto” y “Grave” en una sola clase llamada “Alto”; la clase “Medio” se mantiene tal cual. La columna 2 de la Tabla 2 muestra la nueva distribución de las clases, así como la cantidad de profesores asociados con cada clase.

Es importante destacar que de los 2,943 profesores que fueron evaluados con el instrumento, solamente 2,035 de ellos escribieron un comentario abierto. Aun así, un análisis desarrollado sobre estos datos mostró que en varios casos, el mismo profesor había evaluado más de una vez (datos repetidos), por lo que se procedió a dejar solamente aquellos profesores que pusieron comentarios abiertos y evitando la repetición de personas. La descripción de los datos que quedaron se muestra en la tercera columna de la Tabla 2; el porcentaje final de profesores que cumple con estas características se puede observar en la última columna de la misma tabla.

Tabla 2. Conjunto de datos para los experimentos.

	Número de profesores sin filtrado	Número de profesores con filtrado	Porcentaje
Bajo	1,654	994	60%
Medio	892	587	66%
Alto	397	277	70%
Total	2,943	1,858	

Partimos del uso de los datos anteriormente descritos para determinar el comportamiento en el empleo de n -gramas como mecanismo de representación

de los comentarios abiertos realizados por los profesores, y con la finalidad de evaluar la exactitud de un modelo basado en aprendizaje automático para clasificar a los académicos universitarios de acuerdo a su nivel de Burnout. En la siguiente sección se explica cada uno de los experimentos llevados a cabo en este trabajo de investigación.

3. Resultados experimentales

En este trabajo se plantean dos experimentos: 1) la evaluación del rendimiento de un clasificador supervisado cuando se utilizan n -gramas como mecanismo de representación textual. En particular, se varía la longitud del n -grama, para $n = 1$ hasta $n = 5$, es decir, desde unigramas hasta quintigramas. Este experimento utiliza el conjunto de datos presentado en la sección anterior. La cantidad de profesores por clase, así como la distribución de cada clase se muestra en la Tabla 3. Como puede observarse, la clase “Bajo” tiene una mayor cantidad de datos (53% del conjunto de prueba), superando significativamente a la clase “Alto”, la cual representa únicamente el 15% de los datos del conjunto de prueba.

Tabla 3. Conjunto de datos usado en los experimentos.

Nivel de Burnout	Número de profesores	Porcentaje
Bajo	994	53%
Medio	587	32%
Alto	277	15%
Total	1,858	100%

3.1. Resultados obtenidos

Dado que los datos a clasificar son textos de escritura libre, es decir, sin ningún tipo de formato, es necesario realizar una transformación que permita modelar matemáticamente los conjuntos o clases. Hemos elegido un método de representación vectorial (bolsa de palabras) que utilice como característica a un n -grama. Dado que los datos son discretos, hemos optado por Naïve Bayes Multinomial como el algoritmo de clasificación automático. El componente multinomial puede parecer problemático en el contexto la clasificación de documentos, sin embargo, algunos autores [5] han discutido dichos problemas e incluso han propuesto diversas maneras para aliviarlos, incluyendo el uso de mecanismos de pesado para las características (por ejemplo, tf-idf) en lugar de utilizar solo las frecuencias y la normalización de la longitud del documento, lo cual ha producido un clasificador bayesian que es competitivo con, por ejemplo, máquinas de soporte vectorial.

Para el primer experimento realizado se utilizó un esquema de experimentación basado en 10-fold cross-validation, lo cual significa que parte del conjunto de

datos fue utilizado como entrenamiento y el resto como prueba. Se realizan 10 iteraciones y se reporta el promedio de los resultados. Usando el conjunto de datos que contiene 3 clases (ver Tabla 3) se obtienen los resultados mostrados en la Tabla 4. Cada columna reporta el valor de datos clasificados correcta e incorrectamente para n -gramas ($n = 1, 2, 3, 4, 5$), junto con su valor de exactitud en términos porcentuales.

Tabla 4. Comparación de resultados obtenidos con representación basada en n -gramas ($n = 1, 2, 3, 4, 5$).

	1-gramas	2-gramas	3-gramas	4-gramas	5-gramas
Correctos	1,154/54.05%	1,174/54.99%	1,267/59.34%	1,274/59.67%	994/53.50%
Incorrectos	981/45.95%	961/45.01%	868/40.66%	861/40.33%	864/46.50%

En la Figura 1 se muestra gráficamente el incremento en la exactitud del clasificador cuando la longitud del n -grama se incrementa desde $n = 1$ hasta $n = 4$, comenzando un decremento cuando el valor de n -grama alcanza $n = 5$. Desafortunadamente, el máximo valor de exactitud es 59.67%, lo cual se considera aún como un resultado pobre para tareas de clasificación en ambientes reales.

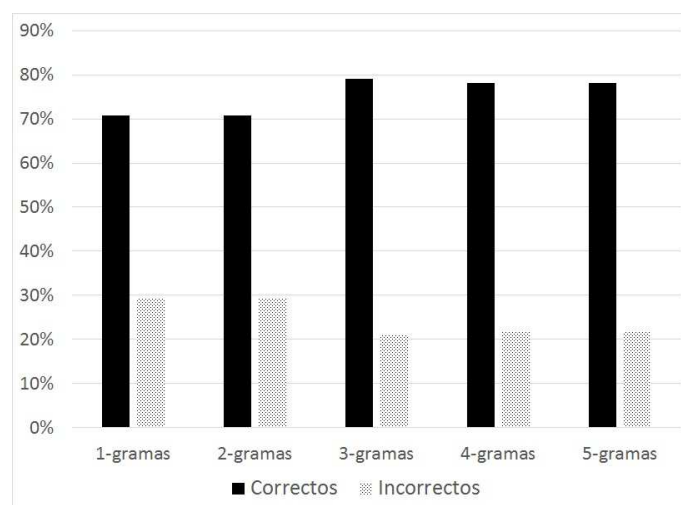


Fig. 1. Exactitud obtenida al clasificar usando n -gramas ($n = 1, 2, 3, 4, 5$).

Debido a la problemática anteriormente planteada, se ha decidido verificar el comportamiento del mismo clasificador cuando se utilizan únicamente las dos clases extremas (“Alto” y “Bajo”). El conjunto de datos usado en este segundo experimento se puede observar en la Tabla 5.

Tabla 5. Conjunto de datos usado en los experimentos.

Nivel de Burnout	Número de profesores	Porcentaje
Bajo	994	78.2%
Alto	277	21.8%
Total	1,271	100%

En la Tabla 6 se pueden observar los resultados cuando el algoritmo de clasificación se ejecuta considerando únicamente las dos clases mas alejadas. Se puede ver un incremento considerable en la exactitud de los resultados. El mejor valor obtenido es cuando el tamaño del n -grama es 3, con un ligero descenso cuando se incrementa el tamaño del n -grama. En la Figura 2 se muestra gráficamente el comportamiento descrito.

Tabla 6. Comparación de resultados obtenidos con representación basada en n -gramas ($n = 1, 2, 3, 4, 5$).

	1-gramas	2-gramas	3-gramas	4-gramas	5-gramas
Correctos	899/70.73%	890/70.03%	1,005/79.07%	994/78.21%	994/78.21%
Incorrectos	372/29.27%	381/29.97%	266/20.93%	277/21.79%	277/21.79%

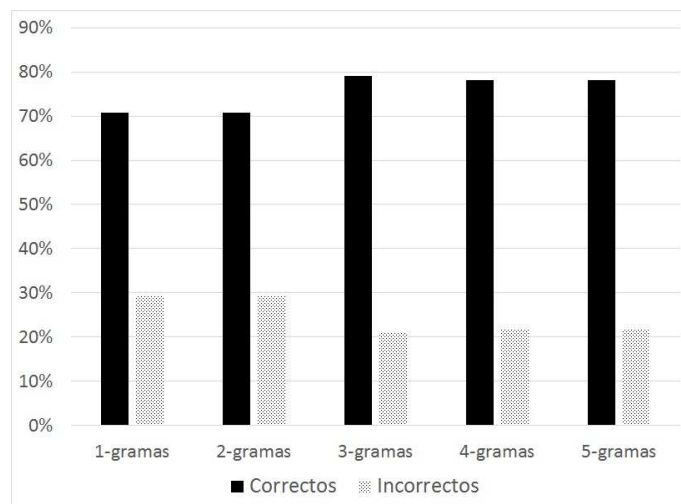


Fig. 2. Exactitud obtenida al clasificar usando n -gramas ($n = 1, 2, 3, 4, 5$).

Finalmente, en la Tabla 7 se enumeran algunos de los patrones de texto encontrados durante el proceso de extracción de características basado en n -gramas.

Algunos de estos patrones son evidentemente positivos y estarían asociado a profesores con bajo nivel de Burnout, mientras que otros de éstos son patrones asociados con un alto nivel de Burnout, por ejemplo: “mejorar las condiciones de trabajo”. Un estudio posterior debería determinar la polaridad de dichos patrones y a partir de esta, probablemente generar una relatoría de las posibles mejoras a fin de decrementar el nivel de Burnout de los académicos universitarios.

Tabla 7. Patrones encontrados en los comentarios.

Patrón de texto
comunicación directa con la responsable
comunicación entre directivos y docentes
con los compañeros de trabajo
docentes participen en las actividades
el ambiente de trabajo es
no tomar en cuenta a
en mi centro de trabajo
eso impide que haya conflictos
gracias por tomar en cuenta
numero de alumnos por grupo
para la toma de decisiones
por parte de las autoridades
por parte de los directivos
que haya conflictos o diferencias
que se tomen en cuenta
un buen ambiente de trabajo
una mejor relacion
y eso impide que haya
mejorar las condiciones de trabajo

4. Conclusiones

En este trabajo de investigación se ha explorado el proceso de la clasificación automática de comentarios abiertos hechos por profesores, con la finalidad de determinar el nivel de Burnout de dichos académicos universitarios. Se llevaron a cabo dos experimentos, variando la cantidad de clases o niveles de Burnout. Para cada experimento se utilizó una representación basada en n -gramas, la cual se evaluó desde longitudes $n = 1$ hasta $n = 5$.

Los mejores resultados que se reportan son de un 79.07% de exactitud, cuando se usan únicamente dos clases (nivel “Bajo” y “Alto”), lo cual significa que el modelo generado permite clasificar correctamente 8 de 10 comentarios abiertos y asociar esa clasificación para determinar si el profesor que expresó el comentario se encuentra con uno de estos dos posibles niveles de Burnout.

El análisis final permitió encontrar una serie de patrones textuales que refieren a aspectos académicos pero más laborales y que estarían asociados con

cierto nivel de Burnout. Es trabajo a futuro el desarrollar métodos computacionales que exploten estos patrones para apoyar mejor la tarea de disminuir el nivel de estrés en los académicos universitarios.

Referencias

1. Freudenberger, H.J.: Staff burn-out. *Journal of Social Issues* 30(1), 159–165 (1974), <http://dx.doi.org/10.1111/j.1540-4560.1974.tb00706.x>
2. Maslach, C., Jackson, S.E.: The measurement of experienced burnout. *Journal of Organizational Behavior* 2(2), 99–113 (1981), <http://dx.doi.org/10.1002/job.4030020205>
3. Maslach, C., Schaufeli, W.B., Leiter, M.P.: Job burnout. *Annual Review of Psychology* 52(1), 397–422 (2001)
4. Olivares, V.E. y Gil-Monte, P.: Análisis de las principales fortalezas y debilidades del maslach burnout inventory (mbi). *Ciencia & Trabajo* 33, 160–167 (2009)
5. Rennie, J.D.M., Shih, L., Teevan, J., Karger, D.R.: Tackling the poor assumptions of naive bayes text classifiers. In: *In Proceedings of the Twentieth International Conference on Machine Learning*. pp. 616–623 (2003)
6. Selye, H.: A syndrome produced by diverse nocuous agents. *Nature* 138(32) (1936)

Una propuesta de clasificación automática de polaridad para notas periodísticas

José A. Baez, Orlando Ramos, María J. Somodevilla, Ivo H. Pineda,
Darnes Vilariño, Concepción Pérez de Celis

Benemérita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

icc.bagatella@outlook.com, {orlandxrf, mariajsomodevilla,
ivopinedatorres, dvilariñoayala, mcpcelish}@gmail.com

Resumen. En este trabajo una técnica para determinar la subjetividad escrita en lenguaje natural es presentada, en particular la polaridad de un conjunto de noticias periodísticas. Un algoritmo cuantitativo para la clasificación automática de polaridad basado en un árbol de búsqueda binaria es propuesto. En dicho algoritmo se utiliza TF*IDF para obtener las palabras con mayor peso en las noticias, además de un lexicon de palabras positivas y negativas como soporte para crear el modelo de clasificación. Se realizaron experimentos utilizando y sin utilizar el algoritmo propuesto como preprocesamiento. Los resultados de la clasificación de polaridad con preprocesamiento fueron satisfactorios independientemente del número de noticias del corpus de prueba.

Palabras clave: Clasificación, aprendizaje automático, análisis de sentimientos, minería de opinión, polaridad, noticias.

A Proposal for Automatic Classification of News Reports Polarity

Abstract. This paper presents a technique for determining the subjectivity written in natural language, in particular the polarity of news reports. A quantitative algorithm for automatic polarity classification, based on a binary search tree, is proposed. TF*IDF is applied in order to determine the words more weight in the news; besides a lexicon of positive and negative words, as a support for building the classification model, is also consulted. Experiments were performed using and without using the proposed algorithm as preprocessing. The classification results of preprocessing polarity were satisfactory regardless of the number of news in the test corpus.

Keywords. Classification, machine learning, sentiment analysis, opinion mining, polarity, news.

1. Introducción

Algunos años atrás, cuando se deseaba conocer la opinión escrita en lenguaje natural de las personas sobre algún tema o servicio específico, las empresas especializadas en esta tarea, se dedicaban a la ardua labor de realizar encuestas, claramente definidas y orientadas a los temas de los cuales se deseaba conocer la opinión de las personas (usuarios), esto claro de forma manual. El paso siguiente está a cargo de un equipo especializado de personas, encargado de analizar de forma manual el contenido de los datos recabados, para así determinar la opinión de los usuarios.

Actualmente el Análisis de Sentimientos (AS) o Minería de Opinión (MO) es uno de los tópicos recientes y más estudiados del Procesamiento de Lenguaje Natural (PLN), esto para determinar de manera automática o semi-automática, la polaridad de un texto en lenguaje natural de un determinado autor (usuario) sobre alguna temática específica. Las aplicaciones sobre esta tarea son diversas, por ejemplo conocer la opinión de los usuarios de Twitter sobre alguna persona (político, artista), empresa, gobierno, campañas sobre productos, noticias, servicios, etc.

Dada la problemática planteada en los párrafos anteriores, en este trabajo se presenta un método para clasificar noticias de manera automática, con ayuda de técnicas, algoritmos de PLN y recursos externos, como lexicones de palabras positivas y negativas, para asignar la polaridad correspondiente a cada noticia.

La estructura del artículo es como sigue, en la segunda sección se describe el estado del arte relacionado con el AS y MO, la tercera sección presenta los conocimientos fundamentales referentes al aprendizaje automático y clasificación. El planteamiento del problema y el conjunto de datos usado es presentado en la cuarta sección, en la quinta sección presentamos nuestra propuesta, los resultados obtenidos se describen en la sexta sección y en la séptima sección las conclusiones.

2. Trabajos relacionados

En esta sección abordaremos los trabajos relacionados con el análisis de sentimientos y la minería de opinión. El trabajo [10] del SemEval 2015 en la tarea 12, utilizó un clasificador de aprendizaje supervisado automático para predecir cada polaridad de opinión (positiva, negativa y neutra). El clasificador se usa en combinación con un proceso de selección basado en la probabilidad para entidades y la detección de atributos de la categoría, teniendo una bolsa de palabras, lemas, bigramas después de verbos y un lexicon basado en características, alcanzando buenos resultados en el dominio de laptops y restaurantes.

Técnicas sobre análisis de sentimientos y minería de opinión se presentan en [5], los enfoques más populares son utilizar lexicon subjetivo, modelos de ngramas y aprendizaje automático. En cambio las técnicas que proponen en el proceso del análisis de sentimientos para textos son: generar un lexicon

(para extraer el conocimiento de los sentimientos), detectar la subjetividad (clasificar el texto en nivel de su naturaleza subjetiva y objetiva), detección de polaridad de sentimientos (la clasificación de los sentimientos en clases semánticas), estructuración de sentimientos (basada en las 5W: why, where, when, what, who) y resumen, visualización y seguimiento de sentimientos (la visualización es generada de manera prudente en un grafo de acuerdo a una dimensión o combinación de dimensiones).

En [1] identificaron tres subtareas: definición del objetivo, separación de las buenas y malas noticias sobre el contenido de sentimientos buenos y malos expresados en el objetivo, y el análisis de opinión claramente marcado que se expresa de forma explícita, sin necesidad de interpretación. El conjunto de datos que utilizan proviene de las aplicaciones NewsBrief¹ y MedISys² de EMM³. En sus experimentos utilizaron WordNet [11], SentiWordNet [3], MicroWNOp [2].

Los experimentos realizados en los trabajos anteriores utilizaron diferentes ventanas alrededor del objetivo, mediante el cálculo de una puntuación de las palabras de opinión identificadas y la eliminación de las palabras que estaban en las mismas palabras de opinión, de tiempo y palabras categoría. Además, se utilizó un recurso incorporado de las palabras de opinión con la polaridad asociada, que denotaron como Tonalidad JRC. Cada uno de los recursos empleados fue mapeado en cuatro categorías, que fueron dadas con puntuaciones diferentes: positivo (1), negativo (-1), positivo alto (4) y alta negativo (-4). Los mejores resultados fueron obtenidos con la combinación de la Tonalidad JRC y MicroWN, en una ventana de 6 palabras.

La tarea para determinar la polaridad de un texto, es una tarea complicada, más aún cuando el texto contiene palabras positivas y negativas en la misma oración. El problema ya no resulta trivial, es decir; se tiene que analizar el contexto que quiere expresar el autor. Por ejemplo, una palabra negativa seguida de una positiva: *el alcalde rechazó darles audiencia*, si solamente se analiza un conteo de palabras positivas y negativas, en el ejemplo sería: una positiva y una negativa. Sin embargo se debe categorizar su polaridad como negativa teniendo en cuenta el contexto. El método propuesto en este trabajo incorpora una mejora en este sentido.

3. Aprendizaje automático y clasificación

En esta sección se describe el aprendizaje automático y clasificación en el contexto de la predicción de polaridad en textos.

¹ Noticias de última hora y noticias en vivo de los últimos minutos/horas. Noticias clasificadas de acuerdo a los sujetos.

² Análisis en Tiempo Real de Noticias de Medicina y temas relacionados con la salud, alertas tempranas por categoría y país.

³ Europe Media Monitor

3.1. Aprendizaje automático

El aprendizaje automático es una rama o subdisciplina de las ciencias de la computación fundamentada en el cómputo suave y el cómputo granular. Dicha disciplina estudia la construcción de algoritmos que puedan aprender de datos de análisis y posteriormente hacer predicciones con otros conjuntos de datos. Para la labor del aprendizaje, el algoritmo debe ser capaz de construir un modelo basado en un conjunto de entrenamiento con el fin de realizar predicciones en el conjunto de datos de prueba.

Los conjuntos de datos utilizados en el aprendizaje automático deben poseer características específicas, y deben corresponder a uno de los siguientes tipos:

- Continuos. Es decir, pueden ser cadenas de texto plano.
- Categorizados. Lo que se puede entender como una discretización de valores ya sea numéricos o alfanuméricos.
- Binarios. Todos los de tipo “verdadero” o “falso”.

De manera similar a la minería de datos, la función de los algoritmos es encontrar patrones y tomar decisiones ajustadas correctamente. El aprendizaje se divide en 2 tipos: supervisado y no supervisado. Si en el conjunto de datos las instancias son marcadas con la respuesta correcta entonces el aprendizaje es de tipo supervisado. En contraste con el aprendizaje no supervisado las instancias no están marcadas por lo cual los investigadores lo utilizan para descubrir datos útiles [6]. Algunos autores describen el tipo semisupervisado como un conjunto parcialmente marcado para un aprendizaje corto y predicción casi desde cero.

El aprendizaje automático tiene varias aplicaciones más allá de desarrollar inteligencia artificial, por ejemplo, según Forbes.com [7], el sitio Amazon planea crear un algoritmo para automatizar el control de acceso a empleados sin la necesidad de la intervención humana que garantice y revoque permisos. Otra aplicación es la identificación de fallas cardíacas que desarrolla IBM lo que supone que una computadora pueda aprender de la información de un paciente y determinar si tiene, tendrá o tuvo una falla cardíaca, sin la necesidad de un médico que lo determine. Y también se ha utilizado en Singapur para desarrollar una aplicación de teléfono móvil que predice ataques y convulsiones, el sistema aprende la diferencia de movimiento común del usuario y cuenta con el patrón de movimientos que ocurren durante una convulsión o un ataque. Entre otras aplicaciones.

3.2. Clasificación

La clasificación automática de textos es un gran reto hoy en día debido a la gran cantidad de información que se puede encontrar en la Web, y una manera de enfrentar este reto es el aprendizaje automático. Sin embargo para que se pueda entrenar a una computadora para clasificar estos textos es necesaria la intervención humana en el preparativo de los datos de entrenamiento, es decir,

una o muchas personas expertas en el área deberían ser las que clasifiquen la información, lo cual consume una enorme cantidad de tiempo y esfuerzo.

La clasificación automática puede ser dividida en 2 tareas bien definidas: la clasificación supervisada que provee de un conjunto de datos de entrenamiento que han sido etiquetados manualmente por un ser humano, y la clasificación no supervisada en la cual no posee ningún esfuerzo humano, la computadora es encargada de inferir alguna función para poder aprender de los datos y etiquetarlos. Hay una tercera forma llamada clasificación semi-supervisada o híbrida, en la cual solo algunos datos son etiquetados con intervención humana.

Entre los métodos más utilizados para la clasificación automática se encuentran: Expectativa Máxima (EM), Clasificador de Naïve Bayes, Máquinas de soporte vectorial, Algoritmo de los vecinos más cercanos, árboles de decisión y Redes neuronales artificiales. Se necesitan grandes cantidades de información para obtener una alta precisión, y la dificultad de obtener datos etiquetados lleva a una importante pregunta: ¿Qué otras fuentes de información pueden reducir la necesidad de datos etiquetados? [9].

Esta investigación se enfoca en proponer un modelo algorítmico con el fin de lograr una clasificación automática híbrida de los datos de entrenamiento, intentando eliminar la necesidad de la intervención humana en el etiquetado de los datos con los que la computadora aprenderá.

4. Planteamiento del problema

Se cuenta con una gran colección de noticias en Español recuperadas y digitalizadas de distintos diarios digitales del estado de Puebla, los cuales abarcan una gran variedad de temas. Es necesario obtener la polaridad de estas noticias dependiendo del reportaje debe clasificarse positiva o negativa.

Debido a la gran cantidad de texto que estas noticias contienen e igualmente el alto número de noticias con las que se cuentan causaría un gran esfuerzo humano el clasificar cada una de ellas como buena o mala noticia, ¿Será posible etiquetar estas noticias con algún método algorítmico?

4.1. Conjunto de datos de aprendizaje

Para esta investigación se usarán diversos recursos, principalmente se cuenta con un corpus de noticias que se clasificarán mediante el algoritmo 1. Para el entrenamiento se utilizó un corpus conteniendo 2904 noticias. En la etapa de pruebas se utilizaron cuatro corpus los cuales contenían 100, 500, 1000 y 1762 noticias cada uno.

Los corpus de entrenamiento serán reclasificados con el algoritmo de clasificación automática propuesto. Se cuenta con dos lexicones obtenidos de Minqing Hu y Bing Liu [4,8], estos lexicones contienen 4783 palabras con polarización negativa y 2005 palabras con polarización positiva, los cuales originalmente están escritos en el idioma inglés y fueron traducidos de manera automática mediante el traductor de Google al idioma Español. Por otra parte se utiliza un diccionario que contiene las conjunciones del idioma Español.

5. Propuesta de solución

Se propone un algoritmo cuantitativo para la clasificación automática de textos basados en un árbol de búsqueda binaria el cual generará además un diccionario pesado mediante la medida TF*IDF para obtener la importancia de las palabras en las noticias y su polaridad. El resultado de este algoritmo es un conjunto de datos de prueba. Se utilizará además el algoritmo de Naïve Bayes para el entrenamiento y prueba, esta última será apoyada por el diccionario generado mediante el algoritmo 1.

El modelo que se propone es un modelo cuantitativo que utiliza una ventana máxima de 3 palabras para analizar frases y determinar la polaridad de la frase analizada. Dicho modelo se enfoca en revisar el texto redactado por una persona de una noticia que ha aparecido en un diario y se ha digitalizado. Para la evaluación de las frases se utilizan los dos lexicones mencionados en la sección 4.1, los cuales fueron traducidos automáticamente al idioma español y enlistan 2005 palabras que tienen una polaridad positiva y 4783 palabras con una polaridad negativa, ninguna palabra esta repetida y si una palabra aparece en un lexicon, significa que no aparecerá en el otro, es decir, son lexicones disjuntos. La revisión del texto de la noticia no se hace palabra por palabra pues esto causa una pérdida de contexto, sino que se analiza un conjunto de tres palabras continuas reconocidas por los lexicones. El análisis del texto obedece al algoritmo 1.

Algoritmo 1. Clasificación cuantitativa de polaridad.

```
N: Noticia
cp: Contador Positivo
cn: Contador Negativo
D: Documento
w: palabra
LP: Lexic\'on Positivo
LN: Lexic\'on Negativo
For each N do:
  cp := 0
  cn := 0
  For i:=0 to len(D) do:
    If w[i] pertenece LP then:
      If w[i+1] pertenece LP then:
        If w[i+2] pertenece LP then:
          cp:= 5
        Else If w[i+2] pertenece LN then:
          cn:= 3
        Else:
          cp:= 2
      Else If w[i+1] pertenece LN then:
        If w[i+2] pertenece LP then:
          cn:= 3
        Else If w[i+2] pertenece LN then:
          cp:= 3
```

```
Else:
    cn:+= 2
Else If w[i+2] pertenece LP then:
    cp:+= 2
Else:
    cp:+= 1
Else If w[i] pertenece LN then:
    If w[i+1] pertenece LP then:
        If w[i+2] pertenece LP then:
            cp:+= 1
        Else If w[i+2] pertenece LN then:
            cn:+= 3
        Else:
            cn:+= 1
    Else If w[i+1] pertenece LN then:
        If w[i+2] pertenece LP then:
            cn:+= 3
        Else If w[i+2] pertenece LN then:
            cn:+= 5
    Else:
        cn:+= 2
Else If w[i+2] pertenece LN then:
    cn:+= 1
Else:
    cn:+= 1
If cp > cn then: "positiva"
Else If cp < cn then: "negativa"
Else: "positiva"
```

Un ejemplo del cálculo de la polaridad se presenta en la Figura 1, la cual representa un árbol de búsqueda binaria. El primer nivel del árbol representa la polaridad de la primera palabra a analizar, el segundo nivel es similar al primer nivel pero con la palabra que le sigue y el tercer nivel es la palabra que está a 2 palabras de distancia de la primera.

Una palabra positiva es aquella que se encuentra en el lexicón positivo, de igual manera para las palabras negativas, mientras que una palabra neutra es aquella que no aparece en ningún lexicón o que sea una conjunción. Cuando la frase de 3 palabras seguidas es terminada de evaluar, un contador toma el valor expresado en la última hoja del árbol en la que se haya movido. Si solo quedan 2 palabras finalmente por analizar entonces se analizará una frase de 2 palabras y el contador se incrementará dependiendo del nodo al cual se haya movido, de igual manera se realiza si solo queda una última palabra por analizar.

6. Experimentos

Un primer experimento utilizó como entrenamiento el corpus con la polaridad original. Posteriormente el corpus de entrenamiento es modificado en su

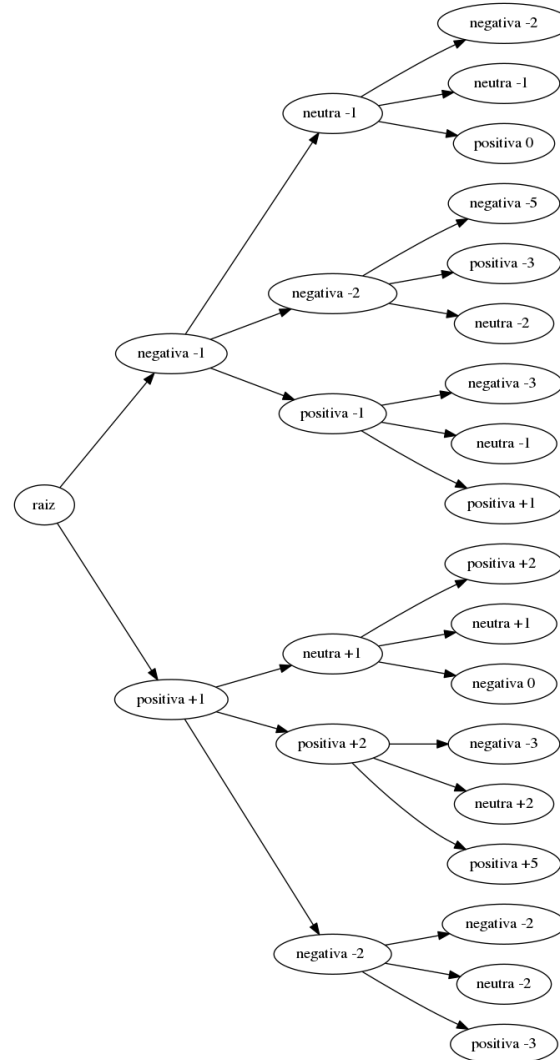


Fig. 1. Árbol de búsqueda binaria que expresa como es evaluada una frase de 3 palabras dentro de un texto.

atributo de polaridad, el cual es calculado utilizando el árbol de decisión binario propuesto. Además se genera un diccionario de apoyo para la prueba del algoritmo de Naïve Bayes. Este diccionario contiene palabras clasificadas como positivas y negativas y la relevancia de la palabra en las noticias del entrenamiento utilizando la medida TF*IDF. Para los experimentos se presentan dos ejemplos de clasificación, los cuales analizan noticias que no son un caso trivial de clasificación.

6.1. Clasificación negativa de noticia

La siguiente noticia es un ejemplo de la complejidad al determinar una polaridad utilizando clasificación automática. Texto original:

El gobernador rechazó darles audiencia para abordar la represión y los presos políticos en Puebla. El gobernador de Puebla Rafael Moreno Valle no recibió a diputados y senadores izquierdistas quienes buscaron una audiencia con él Los legisladores de izquierda Aida Valencia Ricardo Monreal José Arturo López Cándido Manuel Huerta Ladrón de Guevara Alfonso Durazo Juan Luis Martínez Rodrigo Chávez Gerardo Villanueva Loreta Ortiz Luisa María Alcalde María Fernanda Romero Jaime Bonilla y los senadores Mario Delgado David Monreal Martha Palaxof Adán Augusto López Manuel Bartlett y Rabindranath Salazar acudieron a Casa Puebla para reunirse con el gobernador de Puebla Sin embargo personal de las instalaciones les informó que no serían recibidos por Moreno Valle. Luego de esperar algunos minutos afuera de Casa Puebla los legisladores ingresaron a las instalaciones de Gobierno donde aguardaron la reunión con Moreno Valle la cual fue solicitada con antelación por Fernando Jara de la dirigencia estatal de Morena en Puebla. La diputada Luisa María Alcalde escribió "Nos dan aviso en las oficinas del gobernador Moreno Valle de Puebla Rafa Gobernador que rechaza audiencia" La administración estatal había condicionado la reunión de legisladores de Morena para que se llevara a cabo en Casa Aguayo con el secretario de Gobierno Luis Maldonado. La renuencia de Moreno Valle para reunirse con los diputados izquierdistas contrasta con las reuniones que el propio gobernador ha sostenido con diputados locales del PRI PAN Panal PRD y con los líderes magisteriales del SNTE a quienes ha recibido personalmente

El analizador es capaz de juzgar las frases:

- 1 | Frase: rechazo darles audiencia
- 1 | Frase: abordar la represion
- 1 | Frase: no recibió a
- 1 | Frase: no serian recibidos
- 1 | Frase: rechaza audiencia la administración
- 1 | Frase: renuencia de moreno

Esta noticia alcanza un puntaje negativo de 6 unidades en el algoritmo, por lo tanto será clasificada como negativa.

6.2. Clasificación positiva de noticia

Un segundo ejemplo para tratar de determinar una polaridad utilizando clasificación automática. Texto original:

Ayuntamiento de Quecholac presenta su bando de policía y gobierno. El ayuntamiento de Quecholac encabezado por el alcalde Néstor Camarillo medina presentó mediante una reunión general de la sindicatura la Regiduría de Gobernación la dirección de seguridad pública la Regiduría de Hacienda y el juzgado calificador el bando de policía y gobierno para el municipio este ordenamiento municipal fue reformado y adecuado a la actualidad ya que el

pasado bando era obsoleto por lo que ahora estará a la par de las reformas constitucionales y se pretende con el mismo regular las acciones conductas y movimientos mercantiles y de servicios que los ciudadanos reciben a través del ayuntamiento lo anterior fue dado a conocer por la síndico municipal Araceli Campos Jiménez quien explicó que el bando de policía y gobierno fue aprobado en sesión de cabildo el pasado de Abril y publicado en el diario oficial del gobierno del estado en Noviembre del 2015 en la reunión se dijo que no solamente fue una decisión del ayuntamiento el actualizar el reglamento municipal de la administración sino que también es en cumplimiento a la constitución federal local y la Ley orgánica municipal el bando de policía y gobierno se encuentra publicado en el portal de transparencia del Ayuntamiento de Quecholac.

El analizador es capaz de distinguir las siguientes frases:

- +1 | Frase: dirección de seguridad
- +1 | Frase: reformado y adecuado
- 1. | Frase: pasado bando era
- 1. | Frase: obsoleto por lo
- +1 | Frase: reformas constitucionales y
- 1. | Frase: dado a conocer
- 1. | Frase: no solamente fue
- +1 | Frase: decisión del ayuntamiento
- +1 | Frase: cumplimiento a la

Esta noticia contiene frases que denotan tanto positividad como negatividad, y siguiendo el algoritmo propuesto da como resultado que el contador positivo alcanza 5 unidades y el contador negativo alcanza 4 unidades. Por lo que la noticia será clasificada como positiva.

7. Resultados

El algoritmo propuesto tiene una complejidad $O(n)$, donde “n” es el número de palabras contenidas en el corpus a preprocesar. Después de la reclasificación de la polaridad, se ejecuta Naïve Bayes. Los resultados se muestran en la Tabla 2. Se aplicó el clasificador Naïve Bayes usando como entrenamiento noticias que se desconoce su criterio de clasificación, el cual contiene 2904 noticias. Las pruebas se harán 4 diferentes corpus y los resultados se especifican en la Tabla 1.

Tabla 1. Clasificador Naïve Bayes sin preprocesamiento.

Corpus	Numero de noticias	Presicion	Recall
1	100	34%	66%
2	500	37.6%	62.4%
3	1000	35.2%	64.8%
4	1762	53.97%	46.03%

Los resultados muestran precisiones con un promedio de 40.19%. Se observa que el numero de noticias afecta significativamente en la precisión. El corpus 4 duplica al corpus 3 y supera a este en aproximadamente 18% en precisión.

Tabla 2. Clasificador Naïve Bayes con preprocesamiento.

Corpus	Numero de noticias	Presicion	Recall
1	100	61%	39%
2	500	66%	34%
3	1000	64.9%	35.1%
4	1762	72.99%	27.01%

Estos resultados tienen precisiones con un promedio de 66.22%. Se puede observar una mejora de 26.03% como promedio en la precisión. Con respecto a la Tabla 1 a diferencia de los resultados de Naïve Bayes sin procesamiento, la precisión no incrementó significativamente con respecto al tamaño del corpus. Por lo tanto se puede concluir que el algoritmo propuesto exhibe una buena precisión en clasificación en corpus con pocas noticias.

8. Conclusiones

La propuesta reporta resultados satisfactorios en esta primera etapa de desarrollo. Aunque se enfrenta un reto de mejora, el algoritmo es capaz de reducir el esfuerzo humano en la clasificación de noticias. Sin embargo, se cree que se puede incrementar su eficacia utilizando diversas métricas tales como, técnicas de similitud semántica para determinar la carga emocional del texto en el diccionario de apoyo al algoritmo de clasificación Naïve Bayes.

No se presentan comparativas con los trabajos de SemEval, porque los corpus utilizados están en idioma Inglés, y además son para dominios específicos como laptops y restaurantes. Por otra parte solo se han encontrado reportados en la bibliografía predicción de polaridad de noticias en Twitter.

Se encuentra en progreso la utilización de otros clasificadores como los árboles de decisión J48 y las máquinas de soporte vectorial con el objetivo de incrementar la precisión en los resultados.

Referencias

1. Balahur, A., Steinberger, R., Kabadjov, M., Zavarella, V., Van Der Goot, E., Halkia, M., Pouliquen, B., Belyaeva, J.: Sentiment analysis in the news. arXiv preprint arXiv:1309.6202 (2013)
2. Cerini, S., Compagnoni, V., Demontis, A., Formentelli, M., Gandini, G.: Language resources and linguistic theory: Typology, second language acquisition,

- english linguistics. chapter micro-wnop: A gold standard for the evaluation of automatically compiled lexical resources for opinion mining. Milano, IT 23 (2007)
3. Esuli, A., Sebastiani, F.: Sentiwordnet: A publicly available lexical resource for opinion mining. In: Proceedings of LREC. vol. 6, pp. 417–422. Citeseer (2006)
4. Hu, M., Liu, B.: Mining and summarizing customer reviews. In: Proceedings of the tenth ACM SIGKDD international conference on Knowledge discovery and data mining. pp. 168–177. ACM (2004)
5. Kaur, A., Gupta, V.: A survey on sentiment analysis and opinion mining techniques. *Journal of Emerging Technologies in Web Intelligence* 5(4), 367–371 (2013)
6. Kotsiantis, S.B., Zaharakis, I., Pintelas, P.: Supervised machine learning: A review of classification techniques (2007)
7. Laura, H.: Six novel machine learning applications. In: *Forbes* (January 2014), <http://www.forbes.com/sites/85broads/2014/01/06/six-novel-machine-learning-applications/#5d6091f967bf>
8. Liu, B., Hu, M., Cheng, J.: Opinion observer: analyzing and comparing opinions on the web. In: Proceedings of the 14th international conference on World Wide Web. pp. 342–351. ACM (2005)
9. Nigam, K., Andrew, K.M., Thrun, S.: Text classification from labeled and unlabeled documents using em. Kluwer Academic Publishers 39, 103–134 (2000)
10. Saías, J.: Sentiue: Target and aspect based sentiment analysis in semeval-2015 task 12. In: Proceedings of the 9th International Workshop on Semantic Evaluation (SemEval 2015). pp. 767–771. Association for Computational Linguistics, Denver, Colorado (June 2015), <http://www.aclweb.org/anthology/S15-2130>
11. Strapparava, C., Valitutti, A., et al.: Wordnet affect: an affective extension of wordnet. In: LREC. vol. 4, pp. 1083–1086 (2004)

Búsqueda paralela exhaustiva aplicada a cadenas de ADN y ARNi

Jesús García-Ramírez, Ivo H. Pineda T., María J. Somodevilla,
Mario Rossainz, Concepción Pérez de Celis

Benemérita Universidad Autónoma de Puebla,
Facultad de Ciencias de la Computación, México

gr_jesus@outlook.com,
{ipineda,mariasg,rossainz,perezdecelis}@cs.buap.mx

Resumen. Dentro de la bioinformática la búsqueda de cadenas de ADN y ARNi es un problema que requiere procesamiento masivo, por esta razón el enfoque paralelo es una buena opción para reducir los tiempos de ejecución. En este artículo se propone un algoritmo paralelo y se muestran los resultados de su implementación mediante la interfaz de paso de mensajes, para encontrar si es que una sub-cadena esta contenida dentro otra cadena de ADN o ARNi, así mismo se propone una forma de guardar los datos para que estos sean enviados dentro de la red de computadoras.

Palabras clave: Bioinformática, búsqueda paralela, cadenas de ADN, ARNi.

Parallel Exhaustive Search for DNA and ARNi

Abstract. Search DNA strands and RNAi is a problem that needs massive parallel processing, in order to reduce processing time. This paper presents a parallel algorithm that allows you to search for assemblies DNA and RNAi is proposed. the MPI standard, which allows identifying substrings found which are similar.

Keywords. Bio-informatics, parallel search, DNA, RNAi.

1. Introducción

Los métodos de búsqueda en anchura y profundidad son técnicas para encontrar solución a problemas dentro de un gran campo de soluciones. Los métodos anteriores presentan desventajas asociadas a las complejidades de tiempo, las cuales son de orden exponencial[8]. Muchos enfoques se han desarrollado en los últimos años como las heurísticas no garantizan que se obtenga una solución óptima. Otro enfoque es la Búsqueda Exhaustiva Paralela (BEP), la cual se realiza desde diferentes puntos o en orden distinto dentro del espacio de soluciones, ya

que ésta se realiza como procesos independientes, se puede implementar mediante un enfoque paralelo.

Implementar un enfoque paralelo en las búsquedas puede ser de gran ayuda, ya que se reducen tiempos de ejecución, siempre y cuando se cuenten con los recursos necesarios, para que se pueda entregar los resultados correctos y completos. Existen diferentes herramientas para implementar algoritmos paralelos, como lo es: Message Passing Interface (MPI) que utiliza los recursos de diferentes computadoras conectadas en red, asignando tareas a cada una, de esta forma se pueden asignar distintas búsquedas a cada computadora (nodo), los programas son escritos de manera secuencial y las librerías son llamadas para enviar y recibir mensajes, de esta manera se realizan implementaciones paralelas [1]; programación multicore, en este enfoque se utilizan todos los núcleos que tienen los procesadores de las computadoras, asignando a cada uno búsquedas diferentes; programación con CUDA, las tarjetas gráficas que se utilizaron inicialmente para procesamiento de gráficos, ahora se utilizan para realizar programación paralela, aprovechando la totalidad de los núcleos con los que cuentan las tarjetas.

Para este trabajo se pretende implementar una búsqueda exhaustiva paralela con datos de cadenas de ARNi mediante un enfoque paralelo (MPI), de esta manera se pueden reducir tiempos de ejecución, optimizando la manera en que se realizan las búsquedas cadenas de ARNi de gran longitud para encontrar algún patrón de coincidencia en dichas cadenas. El presente trabajo se divide en las siguientes secciones, la sección 2. se aborda aspectos relacionados con el estado del arte sobre los métodos de búsqueda en tradicionales y se enfatiza sobre los métodos establecidos y la importancia de poder contar con una herramienta de bajo costo para lograr resultados similares. Posteriormente en la sección 3. se presenta la metodología seguida en la misma se una subsección 3.1. relacionada con los tipos de datos utilizados, los cuales por las características de los mismos merecen una descripción haciendo énfasis en una propiedad de las levaduras que es el cambio con el tiempo de las características de las cadenas de ARNi. Aspecto importante de la metodología es la posibilidad de poder realizar un ensamblaje de cadenas, elemento importante en cualquier análisis para este tipo de información. Finalmente se detallan los resultados del método y el impacto con respecto a la calidad de la búsqueda y el tiempo logrado para diferentes cadenas.

2. Estado del arte

Dentro de la bioinformática un tema que se ha estudiado desde los años ochenta a la fecha es la búsqueda en cadenas de ADN desde su representación computacional (archivos de texto con una representación como cadenas texto). Dada la amplia variedad de algoritmos de búsqueda que se han propuesto, mención especial merece el algoritmo BLAST inspirado en el algoritmo conocido como "Máxima subsecuencia común" (Longest Common Subsequence LCS) el cual se concibió inspirado en la técnica conocida como Programación Dinámica [2], dicha técnica resuelve subproblemas y se guarda su resultado para encontrar la solución global cuando todos son resueltos [3].

Un nuevo algoritmo para la búsqueda en cadenas de ADN es reportado por William Pearson y David Lipman se discute en [4]. En él se presenta una modificación al algoritmo de búsqueda de cadenas de ADN FASTP, a dicha modificación se le da el nombre de FASTA, el cuál es más rápido que el algoritmo FASTP. FASTP es un algoritmo que sirve para comparar cadenas de ADN, buscar proteínas o cadenas de ADN y comparar dichas cadenas con las que contenga alguna base de datos.

Gulio Pavesi et al. reportan un algoritmo que realiza una búsqueda exhaustiva en patrones de cadenas de ADN. En él se buscan subcadenas de forma que se encuentre un número determinado de veces en la cadena [5].

Por lo regular, las cadenas de ADN se guardan en archivos de texto, los cuales contienen miles de datos. Por esta razón el tiempo de búsqueda en este tipo de archivos es considerable, Ning Zemin et al. proponen el algoritmo SSAHA (Sequence Search and Alignment by Hashing Algorithm) que implementa búsquedas de subcadenas de ADN y crea una tabla Hash para realizarlas con una mayor rapidez [6], otro enfoque que también utiliza tablas Hash es reportado en [15].

La forma de representar las cadenas de ADN puede ayudar a mejorar la búsqueda, así como muestran Sonhammer y Durbin, quienes proponen que la forma de representación sea una forma matricial, lo cual permite ver las cadenas de ADN de manera gráfica [7], al igual que en [14], donde se reporta un sistema en el cual se pueden visualizar de manera gráfica las cadenas de ADN, ya que muchas veces la observación puede ser de gran ayuda para el desarrollo de proyectos de éste tipo.

En los últimos años se han desarrollado diferentes enfoques para éste tipo de búsquedas, como se mencionó anteriormente el enfoque paralelo puede ser una buena alternativa para implementarlas, por ejemplo, en [10] se muestra un enfoque paralelo que utiliza tarjetas gráficas CUDA, el cuál puede ser más rápido a medida que el número de cores de las tarjetas va aumentando, hay que tomar en cuenta que se debe tener el hardware adecuado para que éste enfoque resulte realmente eficiente, ya que, el tiempo de carga de datos en la memoria de la tarjeta es costoso por el bus que implementa. El trabajo presentado en [13] se utiliza un enfoque multi-hilo el cual tiene la ventaja de que los datos son manipulados directamente en la computadora y no se pierde tiempo en enviar los datos como en el enfoque anterior.

Otras formas de realizar la búsqueda dentro de las cadenas son mencionadas a continuación, [11] propone que se aplique el metodo de búsqueda k-mer, donde la principal desventaja que se ve es la memoria RAM, ya que este método requiere la utilización de espacio en memoria. Otros enfoques implementan algoritmos iterativos como es reportado en [12].

En este trabajo se propone una la BEP, la cual se realiza de una manera paralela implementando MPI, siendo enviado a cada nodo un diccionario con los datos extraídos de las bases de datos encontradas, de esta manera se realiza la búsqueda de manera exhaustiva, balancenado la carga de trabajo para así reducir los tiempos de ejecución tomando en cuenta el identificador de cada nodo. Para

las pruebas se aumenta la longitud de la cadena a ser buscada para determinar si impacta en el tiempo de ejecución.

3. Metodología propuesta

En esta sección se propone una forma de realizar una BEP dentro de una red de computadoras. Las búsquedas se realizaron en archivos de texto que contienen una representación de cadenas de ARNi, con su respectivo nombre. La arquitectura que se adopta en esta propuesta es tener un nodo principal, el cual realiza el pre-procesamiento del archivo, dejando los datos de los archivos en un diccionario, los cuales contienen las características y las cadenas, el cual se envía a todos los nodos para realizar la búsqueda de manera exhaustiva. La arquitectura, así como el funcionamiento general del sistema se puede observar en la Fig. 1, en la que existe un nodo principal y n nodos para el procesamiento paralelo.

3.1. Datos a usar: levaduras

Las levaduras son los agentes de la fermentación y se encuentran naturalmente en la superficie de las plantas, el suelo es su principal hábitat encontrándose en invierno en la capa superficial de la tierra. En verano, por medio de los insectos, polvo y animales, son transportados hasta el fruto, por lo que su distribución se produce al azar. Existe un gran número de especies que se diferencian por su aspecto, sus propiedades, sus formas de reproducción y por la forma en la que transforman el azúcar. Las levaduras del vino pertenecen a varios géneros, cada uno dividido en especies. Las especies más extendidas son *Saccharomyces ellipsoideus*, *Kloeckera apiculata* y *Hanseniaspora uvarum*, las cuales representan por sí solas el 90% de las levaduras utilizadas para la fermentación del vino. Como todos los seres vivos, tienen necesidades precisas en lo que se refiere a nutrición y al medio en que viven. Son muy sensibles a la temperatura, necesitan una alimentación apropiada rica en azúcares, elementos minerales y sustancias nitrogenadas, tienen ciclos reproductivos cortos, lo que hace que el inicio de la fermentación sea tan rápido, pero así como se multiplican, pueden morir por la falta o el exceso de las variables mencionadas.

Tomando como referencia estos elementos es que se decide utilizar el conjunto de datos asociados a la especie de levadura conocida como *Schizosaccharomyces pombe* o *S. pombe*, también llamada "fission yeast" en inglés. *S. pombe*, la cual se divide por fisión binaria y produce dos células de igual tamaño. Es muy usada como organismo modelo en el estudio del ciclo celular, ARNi, entre otros. Estas características sobre el comportamiento del hongo unicelular eucariota nos motivó al diseño de la herramienta para poder encontrar de manera rápida y eficiente información asociada a las cadenas de la levadura.

El pre-procesamiento de los datos se realiza mediante un archivo de texto de entrada que contiene el nombre de una cadena de texto, así como sus características. El balance de datos a procesar se realiza en cada nodo con base en

su número de identificador. Posteriormente se realiza la búsqueda en cada nodo de procesamiento y si se encuentra una subcadena preestablecida dentro de la cadena de ARNi se avisa al nodo principal y se imprime en que cadena de ARNi se encontró la subcadena.

Para realizar la búsqueda en cada nodo se realiza un pre-procesamiento, el cual es implementado mediante el análisis del archivo de texto en el que se guardan las características y las siguientes líneas se realiza una unión de todas estas para dejarlas como una sola cadena, posteriormente se distribuye la carga a cada nodo de procesamiento que se tenga dentro de la red de computadoras estableciendo los límites de donde se empieza y donde se termina a realizar la búsqueda en cada nodo. Finalmente se realiza la búsqueda en cada nodo dependiendo de la carga de trabajo que se le asigne. En el Algoritmo 1 se muestra el algoritmo para el nodo principal donde se realiza el pre-procesamiento de los archivos, donde, si la línea contiene el carácter punto y coma (;), significa que la línea contiene las características, de lo contrario se concatena la cadena para formar la cadena general, finalmente los datos son enviados a todos los nodos donde se realizará la búsqueda. En el Algoritmo 2 se muestra el algoritmo para los nodos de procesamiento donde se establece los rangos de búsqueda en el nodo, si se encuentra la cadena retorna un mensaje al nodo principal.

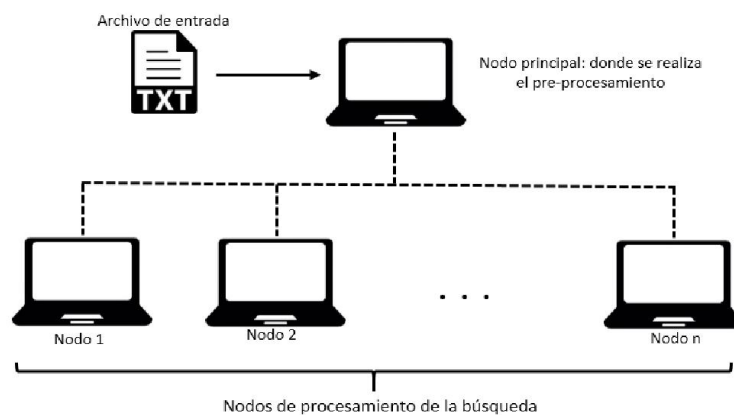


Fig. 1. Arquitectura general del sistema.

3.2. Pre-ensamble de cadenas

En este enfoque se propone un algoritmo para encontrar cadenas candidatas para realizar ensambles en cadenas de ADN de manera paralela con la interfaz de

Algoritmo 1 Pseudocódigo para el nodo principal.

Entrada: Archivo de entrada con cadenas de ARNi.

Salida: Diccionario con el nombre y las características de la cadena (*diccionarioCadenas*).

```
1: for linea in ArchivoCadenas do
2:   if linea contiene ";" then
3:     diccionarioCadenas[linea]=Cadena
4:     Nombre=linea
5:   else
6:     Cadena=linea
7:   end if
8: end for
```

Algoritmo 2 Pseudocódigo para los nodos de procesamiento.

Entrada: Diccionario con las cadenas de texto (*diccionarioCadenas*), Cadena a buscar (*cadBusqueda*).

Salida: Mensaje de en qué cadena encontró *cadBusqueda*.

```
1: NoProcesadores=número de nodos de procesamiento
2: Id=identificador del nodo
3: Longitud=diccionarioCadenas.length()
4: LongitudProcesamiento=int(Longitud/(NoProcesadores)+1)
5: Inicio=id*LongitudProcesamiento+1
6: Fin=Inicio+LongitudProcesamiento-1
7: if Fin es mayor que Longitud then
8:   Fin=Longitud
9: end if
10: for Inicio to Fin do
11:   if diccionarioCadenas in diccionarioCadenas[Inicio] then
12:     Print Se encontró la cadena en: Inicio
13:   end if
14: end for
```

paso de mensajes (Messages Passing Interface, MPI), el cual da la opción para realizar procesamiento paralelo en una red de computadoras o programación multicore.

En esta implementación se utilizó el lenguaje de programación Python y las librerías mpi4py y biopython, así como la base de datos Pombe database que contiene cadenas de una levadura presente en la cerveza, la cual contiene siete archivos con diferentes lecturas de ARN-i. Este enfoque se realiza en una computadora con seis núcleos.

Los procesos realizan las siguientes tareas: el primer nodo será denominado nodo de distribución el cual realiza una lista con los distintos archivos que componen la base de datos, la cual envía a cada proceso un archivo para realizar el procesamiento de las cadenas en dicho archivo (Algoritmo 1); los demás nodos son denominados nodos de procesamiento en los cuales se realizan la búsqueda de la cadena en el archivo correspondiente (Algoritmo 2).

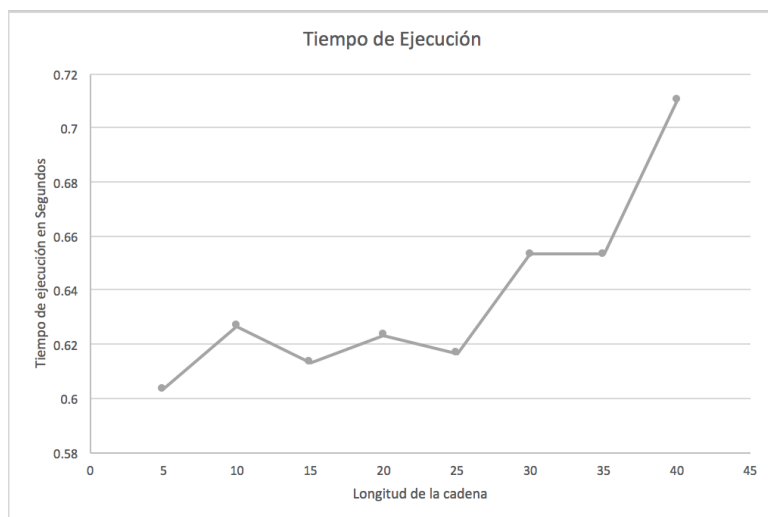


Fig. 2. Gráfica de tiempos de ejecución dependiendo del tamaño de la cadena de búsqueda.

El Algoritmo 3 realiza un balanceo en la carga, ya que envía un archivo a cada nodo de procesamiento, realizando esto de manera paralela y reduciendo el tiempo de procesamiento.

El Algoritmo 4, es el procesamiento que se realiza, primero se crea una lista, en la cual se guardarán las posibles cadenas para realizar el ensamble, de esta manera se hace un primer filtro para descartar cadenas que no contienen la subcadena para hacer el ensamble. Después se realiza el algoritmo de búsqueda de la cadena más larga, de la cual se puede encontrar la cadena común más larga y así determinar cual es la mejor cadena para realizar el ensamble.

Algoritmo 3 Balanceo de carga.

Entrada: nombre de los archivos de la base de datos.

Cadena: cadena a ser buscada para cada archivo en la base de datos: enviar el archivo que corresponde a cada nodo y la cadena a buscar

4. Experimentos y resultados

Para los experimentos se utilizó el lenguaje de programación Python con las librerías para implementar MPI, en una máquina virtual, con un sistema operativo Linux Ubuntu con un solo núcleo y una frecuencia de reloj de 2.6 GHz, 2GB de Memoria RAM y un disco duro virtual de 10 GB.

Algoritmo 4 Búsqueda de la cadena más larga.

Entrada: nombre del archivo a procesar.
 cadena:cadena a ser buscada
 1: Se abre el archivo de cadenas del archivo lista
 2: **for** cada cadena Cad en el archivo: **do**
 3: **if** cadena se encuentra en el Cad: **then**
 4: añadir Cad a la lista
 5: **end if**
 6: **end for**
 7: **for** cada valor en lista: **do**
 8: Aplicar Longest Common Subsequence(LCS) para encontrar la cadena más larga
 y es donde se podría hacer el ensamble
 9: **end for**

Los archivos que se utilizaron para las pruebas fue la base de datos de cadenas de ARNi que se encuentra en el sitio oficial de Pombase [9]. Dichos archivos están organizados de la siguiente manera. Primero se tiene una descripción de la cadena de ARNi que se va a procesar, seguida de la representación con letras, el cual contiene 5280 cadenas de ARNi.

Se realizaron experimentos con diferente número de nodos para determinar si la carga de trabajo se realiza correctamente. En la Tabla 1 se reporta cuantas coincidencias encuentra por nodo, por lo que se puede observar se distribuye la carga de trabajo a medida que se aumenta el número de nodos, ya que se van encontrado menor número de coincidencias en cada nodo, además la suma de coincidencias debe ser la misma para todas las pruebas ya que lo único que cambia es el número de nodos en los que se ejecuta el programa.

Tabla 1. Número de coincidencias encontradas en cada nodo del sistema.

Número de nodos	Número de coincidencias en cada nodo
3	562,508,575
4	417,405,396,427
5	329,341,294,329,352
6	270,292,260,249,272,302
7	233,249,230,206,236,241,250
8	205,212,214,191,176,220,207,220

Para determinar si la longitud de las cadenas que se tiene que buscar tiene un impacto directo sobre el tiempo de ejecución, se realizaron experimentos con diferentes longitudes de cadena para su búsqueda, como se puede observar en la Figura 2 los tiempos de ejecución van aumentando conforme la longitud de la cadena de búsqueda crece, esto es lógico, ya que, el número de comparaciones va aumentando a medida que la longitud de la cadena va creciendo.

5. Conclusiones

Haber adoptado un algoritmo con enfoque paralelo en problemas que requieran búsquedas exhaustivas o comparaciones dentro de un espacio de soluciones muy grande, puede ayudar a encontrar la solución del problema en menos tiempo, ya que el trabajo que correspondería a un solo procesador en una computadora convencional, se distribuye entre una red de computadoras o en diferentes procesadores dentro una sola computadora.

Como se pudo observar en la Figura 2 y en la Tabla 1 la distribución del trabajo en diferentes computadoras ayudó a que se redujeran los tiempos de ejecución, también dependiendo de la longitud de la cadena a buscar se aumenta el tiempo de ejecución.

Después de varios experimentos hemos llegado a la conclusión que un posible enfoque puede consistir en tener a disposición mayores cantidades de datos y poder contar con la opinión de expertos del área para brindar opciones de búsqueda pensando en subcadenas de una determinada longitud y con significado. Como elemento importante es el hecho que con las configuraciones actuales de cómputo se pueden obtener resultados favorables sin necesidad de tener una computadora de alto costo y tener resultados en tiempos de ejecución aceptables.

Referencias

1. Foster I.: Designing and Building Parallel Programs. Addison Wesley. primera edición (1995)
2. Altschul S, Gish W., Miller W., Myers E., Lipman D.: Basic Local Alignment Search Tool. Journal of molecular biology 215(3), pp. 403–410 (1990)
3. Levitin, A.: Introduction to Design and Analysis of Algorithm. Pearson Education, tercera edición (2012)
4. Pearson W., Lipman D.: Improved tools for biological sequence comparison. In: Proceedings of National Academy of Science (Biochemistry), Vol. 85, pp. 2444–2448 (1988)
5. Pavesi G., Mauri G., Pesole, G.: An algorithm for finding signals of unknown length in DNA sequences. Bioinformatics, 17(1), S207–S214 (2001)
6. Ning Z., Cox A., Mullikin J.: SSAHA: A Fast Search Method for Large DNA database. Genome Research, 1725–1729 (2001)
7. Sonnhammer E., Durbin R.: A dot-matrix program with dynamic threshold control suited for genomic DNA and protein sequence analysis. Gene, Vol. 104 (1996)
8. Russel S., Norving P.: Artificial Intelligence, A Modern Approach. Pearson Education, tercera edición (2010)
9. Wood V, Gwilliam R, Rajandream MA, Lyne M, Lyne R, Stewart A, Sgouros J, Peat N, Hayles J, Baker S, et al.: The genome sequence of Schizosaccharomyces pombe. Nature (2003)
10. Encarnaco G., Sebastiao N.: Advantages and GPU implementation of high-performance indexed DNA search based on suffix arrays. In: Proceedings de Conferencia Internacional en Cómputo de Alto rendimiento y simulación, pp. 49–55 (2011)

11. AIslam T., Pramanik S.; Ji X., James R. Cole J., Zhu Q.: Back translated peptide K-mer search and local alignment in large DNA sequence databases using BoND-SD-tree indexing. In: Proceedings de IEEE 15th International Conference of Bioinformatics and Bioengineering (BIBE), 1–6 (2015)
12. Cheng P., Chen H., Kao J.: Protein surface search in DNA-binding protein prediction by Delaunay triangulation modeling. In: Proceedings de International Computer Symposium, pp. 783–788 (2010)
13. Erodula K., Bach C., Bajwa H.: Use of Multi Threaded Asynchronous DNA Sequence Pattern Searching Tool to Identifying Zinc-Finger-Nuclease Binding Sites on the Human Genome. In: Proceedings de Eighth International Conference on Information Technology, pp. 985–991 (2011)
14. Glisovic N.: System for DNA visualization and clustering in searching through information. In: Proceedings de International Symposium on Computational Intelligence and Informatics, pp.169–170 (2010)
15. Faro S., Lecroq T.: Fast searching in biological sequences using multiple hash functions. In: Proceedings de International Conference on Bioinformatics & Bio-engineering, pp. 175–180 (2012)

Un modelo ontológico espacial para el tratamiento de desórdenes de conducta infantil

María J. Somodevilla, Andrea M. P. Tamborrell,
Ivo H. Pineda, Concepción Pérez de Celis

Facultad de Ciencias de la Computación,
Benemérita Universidad Autónoma de Puebla, México

mariajsomodevilla@gmail.com, andie-3-@live.com,
{ivopinedatorres,mcpcelish}@gmail.com

Resumen. Los desórdenes de conducta son un comportamiento de patrones persistentes en niños y adolescentes en presencia de los cuales los derechos de otros o reglas básicas son violados. Considerando este problema, se propone la creación de un modelo ontológico para el tratamiento espacial de desórdenes de conducta en el idioma español. Este modelo ontológico considera la interacción de ontologías individuales, el cual permitirá responder a consultas espaciales basadas en semántica, tales como lugares de incidencia de cierto trastorno de conducta o interacciones medicamentosas en el tratamiento de cierto problema de conducta. En el diseño del modelo se aplican los conceptos de reusabilidad e integración de ontologías las cuales han sido individualmente axiomatizadas para lograr consistencia del sistema.

Palabras clave: Desórdenes de conducta, modelo ontológico espacial.

An Ontological Model for Child Conduct Disorders

Abstract. Conduct disorders are repetitive behavior patterns in child and adolescents where rules or basic rights of third parties are broken. Considering this problem, the creation of an ontological model for spatial treatment of conduct disorders in the Spanish language is proposed. This ontological model considers the interaction of individual axiomatized ontologies, which will be able to answer spatial queries based on semantic, such as incidence places of conduct disorder or drugs in the treatment of a behavioral problem. Concepts of ontologies' reusability and integration to achieve the system consistency along the model design have been incorporated.

Keywords. Conduct disorder, spatial ontological model.

1. Introducción

Los desórdenes de conducta son una condición severa caracterizada por un comportamiento hostil y, en ocasiones violencia física. Infantes con problemas de conducta exhiben crueldad, que incluye desde maltrato físico a personas o animales hasta el acoso (*bullying*). Desde la infancia y adolescencia los problemas de conducta normalmente desarrollan en el adulto una personalidad antisocial, que necesita ser corregida con tratamiento tan pronto como sea detectada, si este problema es detectado a tiempo, mejor será el pronóstico [3]. Una ontología representa un modelo conceptual describiendo cierto dominio, tipos de objetos y conceptos que existen, así como sus propiedades y relaciones.

Considerando la gravedad de los desórdenes de conducta infantil, las tecnologías de la información podrían utilizarse como un soporte para su detección, así como una guía de orientación para padres. La Web se ha convertido en el principal medio de comunicación en nuestros días y son las ontologías las que proveen su base de conocimiento semántica. En este trabajo se propone desarrollar un sistema de ontologías que relaciona tres ontologías: desórdenes de conducta infantil, espacial de México y de psicofármacos.

2. Ontologías biomédicas

Existe trabajo previo en la FCC en el desarrollo de ontologías espaciales biomédicas a través de los proyectos en [10, 11, 12]. Dentro de estos proyectos se desarrollaron las ontologías *GeoOntoMex* y *HealthOntoMex*.

Diferentes ontologías se han creado en el ámbito médico como la *Disease Ontology*, la cual tiene como finalidad agrupar un vocabulario médico controlado desarrollado en el campo bioinformático, en colaboración con el Centro de Medicina Genética de la Universidad Northwestern, Chicago. *Disease Ontology* fue diseñada para facilitar una relación de las enfermedades y las condiciones asociadas a ellas para códigos médicos particulares como ICD9CM [1], SNOMED [2] y otros [9].

La ontología SNOMED [3] incluye una terminología de 364,000 conceptos de cuidado médico, con significados únicos y definiciones basadas en lógica formal organizadas en jerarquías [9]. DrOn [13] es una ontología para consultar fármacos por ingredientes, por disposición molecular y terapéutica, y por efectos secundarios.

La mayor parte de los trabajos ontológicos desarrollados en el ámbito de medicina están desarrollados en el idioma inglés y los pocos trabajos hechos de desórdenes de conducta no abordan específicamente este problema en infantes, por lo cual este trabajo representaría una gran diferencia en dicha área y en los países de habla hispana.

La publicación de ontologías médicas se puede realizar en sitios como *biontology.org* para que estas puedan ser de libre acceso en todo el mundo. Dicho sitio contiene una gran base de datos con ontologías del área de la medicina y permite a cualquier usuario subir su trabajo ontológico para que este pueda ser consultado por personas alrededor del mundo. Además, este sitio permite al usuario que lo visita ver a detalle

los datos de la ontología tales como el nombre, el identificador único e incluso tiene un apartado para visualizar la ontología de manera gráfica en línea.

3. Desórdenes de conducta

Los desórdenes de conducta son una serie de problemas comportamentales y emocionales que se presentan en niños y adolescentes. Los problemas pueden involucrar comportamiento impulsivo o desafiante, consumo de drogas o actividad delictiva [7]. De acuerdo a [4], infantes con desorden de conducta pueden presentar cuatro tipos de comportamiento negativo como se detalla en 3.1.

3.1. Tipos de comportamiento

- Agresión hacia personas o animales: (a) acoso y amenaza, (b) iniciar peleas físicas, (c) daño físico serio con el uso de armas, (d) asalto sexual.
- Destrucción de propiedad ajena: (a) daño intencional a la propiedad de otros, (b) iniciar incendios de manera deliberada.
- Acciones como engañar, mentir y robar: (a) allanamiento, (b) mentir para obtener cosas o favores o para evitar obligaciones, (c) robar sin confrontación.
- No respetar las reglas: (a) faltar a clases, (b) escapar de casa y/o pasar las noches sin permiso de los padres.

3.2. Causas de los desórdenes de conducta

Diferentes factores contribuyen a que una persona desarrolle desórdenes de conducta, incluyendo daño cerebral, abuso o negligencia, problemas genéticos, problemas escolares y experiencias traumáticas [4].

Existen diferentes tipos de factores de riesgo como:

- Biológicos: historial familiar de comportamiento negativo, desafiante o desorden de comportamiento.
- Modo de vida: violencia, divorcio, rechazo familiar.

3.3. Tratamientos e intervenciones

Está reportado en la bibliografía que los desórdenes de conducta se tratan con fisioterapia individual o grupal y/o con psicofármacos.

3.3.1. Intervenciones psicosociales

Sin intervención es probable que los desórdenes de comportamiento negativo (*Disruptive Behavior Disorders*) progresen [6]. Hay un número de tratamientos prometedores que están disponibles y que cuando son completados pueden llevar a grandes beneficios. Una investigación en [1] demostró que las personas que completan la

“Terapia de Intervención Padre-Hijo” muestran un cambio significativo al terminar dicha terapia.

En [5] se han identificado 16 tratamientos basados en evidencias para los desórdenes de comportamiento. Dos ejemplos son:

- PMT [4] “Entrenamiento para padres”. Está dirigido a padres y les enseña a identificar antecedentes, comportamientos resultantes y consecuencias asociadas a sus hijos así como también a ellos mismos. Por último, el entrenamiento se enfoca en reforzar conductas deseadas.
- PCIT [5] “Terapia de interacción padre-hijo”. Enfatiza mejoras en la relación entre los padres e hijos y ofrece herramientas para ayudar a manejar comportamientos que son negativos.

3.3.2. Psicofármacos

Los psicofármacos se clasifican en ansiolíticos, antidepresivos, antimaníacos y antipsicóticos. Medicamentos antipsicóticos o neurolepticos son ampliamente utilizados en el tratamiento para el tratamiento de agresión aguda y crónica en varias poblaciones [14].

Estudios indican que antipsicóticos atípicos generalmente son más eficaces que un placebo en el tratamiento de agresiones, pero tienen diferentes efectos en cada persona. En menores de edad con problemas de desórdenes de comportamiento, la *risperidona* es el medicamento más estudiado.

4. Propuesta de solución

En esta sección se define el problema a resolver, se establece la metodología a utilizar en el desarrollo, así como los criterios de evaluación. Finalmente se llevan a cabo todas las etapas de desarrollo del sistema de ontologías propuesto.

4.1. Definición del problema

Los desórdenes de conducta infantil representan un problema de salud a nivel global, entonces las tecnologías de la información podrían utilizarse como un soporte para su detección, así como una guía de orientación para padres. En este sentido se realizará la creación de un sistema de ontologías reutilizable, usable y actualizable para tratar semánticamente los desórdenes de conducta en infantes así como sus tratamientos y las áreas geográficas de influencia.

4.2. Metodología

La metodología de diseño utilizada es una modificación de la propuesta por [2]. Los pasos de la metodología general para construcción de ontologías son: (1) Elicitación de términos, (2) Identificación de módulos, (3) Integración de módulos y (4)

Evaluación global del sistema de ontologías, donde los módulos son las ontologías individuales.

Las características principales del sistema de ontologías son las siguientes:

- Diseño Orientado a Dominios. Dominio: Salud, subDominio: salud mental infantil.
- Ontología Orientada a la Reutilización: GeOntoMex [10], dominio: Geográfico. Esta ontología representa la división política administrativa de México.
- Ontología Modular: 3 ontologías individuales conectadas por axiomas semánticos.
- Ontología Incremental: Permite la definición de nuevos axiomas estructurales y de comportamiento.

La técnica que propone [2] para la extracción de términos que integrarán la ontología es la generación de preguntas de competencia relacionadas con el problema a resolver. Dichas preguntas se espera que el sistema de ontologías responda al final de su realización. Las siguientes son algunas de las preguntas de competencia:

1. ¿Cuáles son los síntomas del desorden de conducta “D”?
2. ¿Es “T” un tratamiento adecuado para la actitud “A”?
3. ¿Qué causas originan un desorden de conducta “D”?
4. ¿Cuáles son los tratamientos indicados para un desorden de conducta “D”?
5. ¿En qué área geográfica han sido reportados casos del desorden de conducta “D”?

4.3. Evaluación de la ontología

La consistencia de la ontología es evaluada por el razonador a través de los axiomas que se definen para establecer: las relaciones jerárquicas (taxonómicas), (2) las relaciones de tipo: *data properties*, es decir, preguntarse ¿cuáles características son suficientes y necesarias para definir a un concepto? y (3) las relaciones de tipo *object properties*, es decir, ¿se han cubierto todas las relaciones explícitas de las preguntas de competencia?

Por otra parte, se utilizan tres métodos de evaluación para la ontología, los cuales consisten en: (1) evaluación del conocimiento y conceptualización, es decir, el grado de cumplimiento del listado de las preguntas de competencia previamente establecidas, (2) evaluación de la calidad, grado de cumplimiento con los principios de diseño y (3) evaluación de rendimiento, velocidad de las respuestas que el razonador generará.

4.4. Diseño conceptual de ontologías individuales

El sistema se compone de tres ontologías individuales. Las ontologías de desórdenes de conducta y la de psicofármacos fueron expresamente creadas para este proyecto y la ontología espacial fue adaptada y reutilizada.

4.4.1. Ontología de desórdenes de conducta

La ontología de desórdenes de conducta consta de diferentes clases, individuos, propiedades de individuos y propiedades de objeto (relaciones entre clases) como se muestra en la Figura 1. Las clases son: Causa, Desorden de conducta, Paciente, Síntoma y Tratamiento.

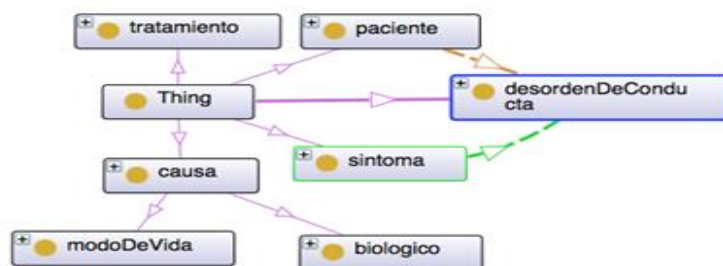


Fig. 1. Ontología OntoDesorden.

Las relaciones entre clases se presentan en la Tabla 1.

Tabla 1. Relaciones entre clases.

Relación	Dominio	Rango
tieneCausas	Desorden de conducta	Causa
tieneSíntomas	Desorden de conducta	Síntoma
tieneTratamiento	Desorden de conducta	Tratamiento
recomiendaTratamiento	Paciente	Tratamiento
tieneDesorden	Paciente	Desorden de conducta

Los individuos de las clases tienen las propiedades que se presentan en la Tabla 2.

Tabla 2. Propiedades de las clases.

Clase	Dominio
Causa	causaNombre causaTipo causaDescripcion
Desorden de conducta	DesordenNombre desordenTipo
Paciente	pacienteEdad pacienteGenero pacienteID
Síntoma	sintomaNombre sintomaIntensidad sintomaDescripcion
Tratamiento	tratamientoNombre tratamientoTipo tratamientoDuración

4.4.2. Ontología de psicofármacos

La ontología de psicofármacos se muestra en la Figura 2. Las clases son: Ansiolítico, Antidepresivo y Antipsicótico.

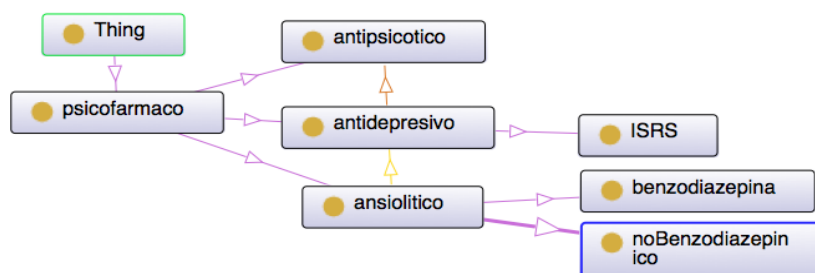


Fig. 2. Ontología Psicofármaco.

Las propiedades entre clases de la ontología psicofármacos se describen en la Tabla 3.

Tabla 3. Propiedades de las clases.

Clase	Dominio
Psicofarmaco	farmacoNombre farmacoOtroNombre farmacoDescripcion
Antipsicotico	antipsicoticoGenerico antipsicoticoGeneracion
ISRS	ISRSGenerico
noBenzodiazepinico	noBenzoGenerico
benzodiazepina	benzoAccion benzoGenerico

Las relaciones entre clases se muestran en la Tabla 4.

Tabla 4. Relaciones entre clases.

Relación	Dominio	Rango
interaccionMedAns	Ansiolitico	Antidepresivo
interaccionMedDep	Antidepresivo	Antipsicotico

4.4.3. GeOntoMex

Esta ontología es reutilizada porque es importante dar respuesta a consultas semánticas tales como la pregunta de competencia 5.

En la figura 3 se muestra el diseño conceptual de la ontología. Para el sistema se utiliza la clase de ExtensionGeopolitica que contiene la definición taxonómica política administrativa de México.

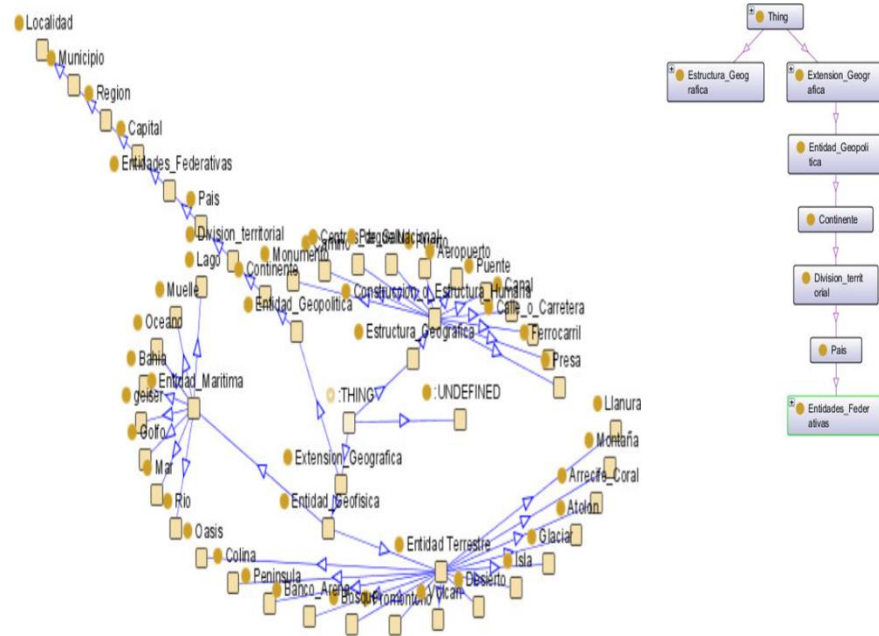


Fig. 3. (a) GeOntoMex.

(b) Reutilización de GeOntoMex.

4.5. Creación del sistema de ontologías

El sistema de ontologías representa la interacción de las ontologías individuales a través de las relaciones entre éstas (*object properties*). Esta forma de relación garantiza un modelo de diseño modular, interactivo e iterativo es decir se mantienen aparte los individuos de cada entidad lógica y al mismo tiempo pueden cooperar con los individuos de otras ontologías para generar nuevo conocimiento. La figura 4 muestra el esquema contextual del sistema de ontologías propuesto. Se muestran las tres ontologías y las relaciones *Psicofarmaco sePrescribe para Desorden* y *Desorden seLocaliza en Mexico*.

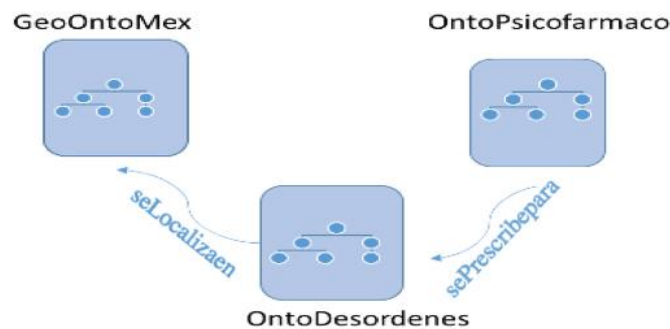


Fig. 4. Sistema de ontologías propuesto.

4.6. Consultas en DL Query

Las reglas DL (Description Logic) se definieron para responder una selección de las preguntas de competencia enlistadas en la sección 4.2. En la Tabla 5 se muestran dichas preguntas de competencia, así como la regla que le da respuesta.

Tabla 5. Definición de DL queries.

Pregunta de Competencia	Sintaxis de Consulta
¿Cuáles son los síntomas generales de los desórdenes de conducta?	9sicofá and sintomaPerteneceA value "Desorden de Conducta"^^ xsd:string
¿Qué causas originan un desorden de oposición desafiante?	causa and causaOriginaUn value "Desorden de Oposicion Desafiante"^^ xsd:string
¿Cuáles son los tratamientos indicados para un desorden de conducta "D"?	9sicofármaco and psicofarmaco-SeRecomiendaParaDesorden value desorden2

En la sección 5 se presentan los resultados de la clasificación reportados por Prótegé 5.0.0.

5. Resultados

Al ejecutar un razonador sobre la ontología, esta no registra ningún error por lo cual se concluye que la ontología es consistente.

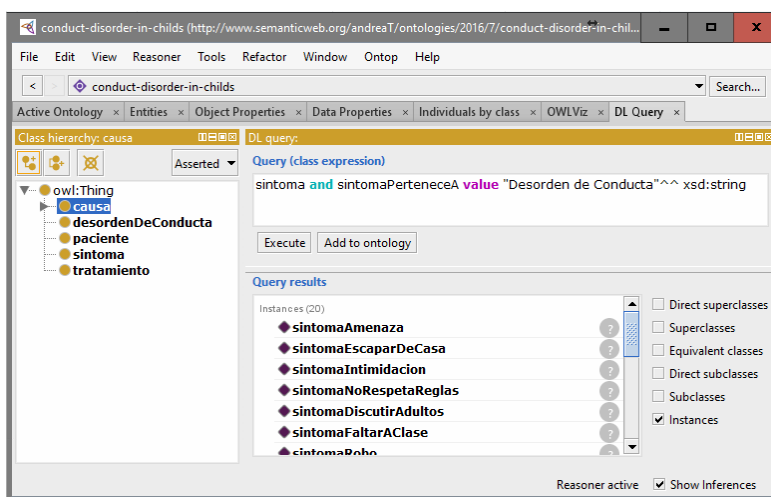


Fig. 5. Ejecución de primera pregunta de competencia.

5.1. Ejecución de DL query

En las figuras 5, 6 y 7 se muestra la ejecución las preguntas de competencia descritas en la tabla 5.

En la Figura 5 se muestran todos los síntomas que han sido capturados como individuos de la clase síntoma y que se asocian a los desórdenes de conducta infantil.

En la Figura 6 se muestran los resultados de la ejecución de la expresión de clase que determina las causas que originan un desorden de oposición desafiante.

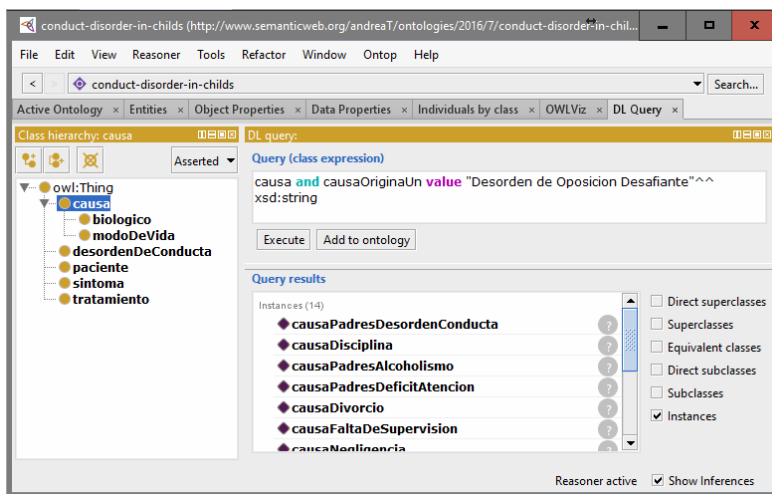


Fig. 6. Ejecución de segunda pregunta de competencia.

La Figura 7 presenta los resultados de la consulta asociada con el tratamiento con psicofármacos indicados para el desorden2. Los individuos por el momento son genéricos pues se está comprobando está información con los expertos.

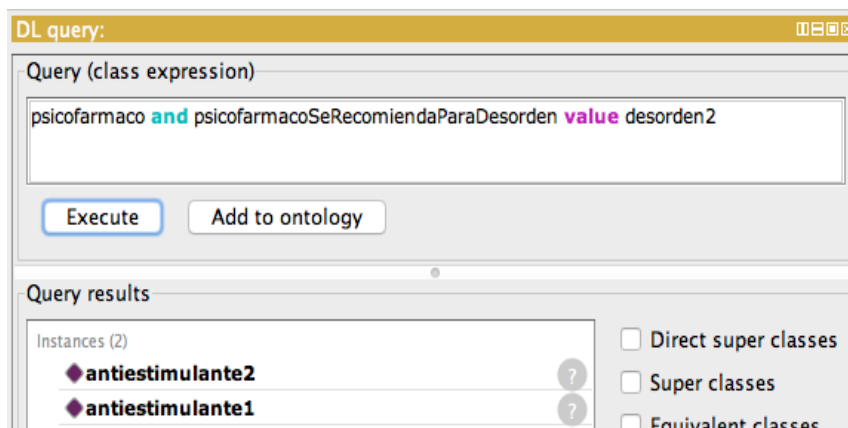


Fig. 7. Ejecución de tercera pregunta de competencia.

6. Conclusiones

Las alteraciones de la conducta en la infancia, han pasado a ser no solo el motivo de consulta más frecuente en psiquiatría infanto-juvenil, sino un motivo de alarma para la sociedad en general. Se propuso un sistema de ontologías que integra a través propiedades de los objetos tres ontologías individuales: desórdenes de conducta, geográfica de México y de psicofármacos. Se implementan además las relaciones entre clases que posteriormente se usaron en la definición de reglas DL para dar respuesta a las preguntas de competencia. El sistema de ontologías propuesto es consistente de acuerdo con la validación realizada por Protégé 5.0.0, donde fueron validados los axiomas de clase, de relaciones y de individuos.

Como trabajo a futuro se está trabajando en la definición una metodología, para el poblado semiautomático del sistema de ontologías propuesto, soportado por conceptos y herramientas de PLN e IR. Esta metodología comprende desde la búsqueda de casos clínicos, síntomas y tratamientos en diferentes sitios Web de la especialidad hasta la inserción de individuos. La población semiautomática permitirá el rápido crecimiento de individuos en la base de conocimientos para lograr una mejor precisión y orientación en los casos de desórdenes de conducta.

La presente propuesta tendrá un impacto social puesto que se plantea que pueda ser utilizado como base de conocimiento, en un sistema para la toma de decisiones relacionado con desórdenes de conducta infantil en México.

Referencias

1. Boggs, S.R., Eyberg, S.M., Edwards, D.L., Rayfield, A., Jacobs, J., Bagner, D., Hood, K.K.: Outcomes of parent-child interaction therapy: A comparison of treatment completers and study dropouts one to three years later. *Child & Family Behavior Therapy* (2004)
2. Bravo, M., Perez, J., Velazquez, J., Sosa, V., Montes, A., Lopez, M.: Design of a shared ontology used for translating negotiation primitives. *International Journal of Web and Grid Services*, 237–259 (2006)
3. Child Mind Institute. Guide to conduct Disorder, What is it?
4. Desórdenes de la Conducta. American Academy of Child Adolescent Psychiatry (2004)
5. Eyberg, S.M., Nelson, M.M., Boggs, S.R.: Evidence-based psychosocial treatments for children and adolescents with disruptive behavior (2008)
6. Gathright, M., Tyler, L.: Disruptive Behaviors in Children and Adolescents (2012)
7. Medline Plus.: Trastorno de Conducta Biblioteca Nacional de Medicina de los EEUU (2015)
8. National Institutes of Health. SNOMED CT. U.S. National Library of Medicine (2016)
9. Sarrión, A., Corrales, Y.: Un acercamiento a las ontologías médicas y su importancia. IV Jornada Científica de la SOCECS
10. Somodevilla, M.: Desarrollo de Sistemas de Ontologías para descubrir patrones de estilo de vida asociados con enfermedades crónicas no transmisibles. SOGM-ING15-I VIEP (2015)
11. Somodevilla, M.: Desarrollo de técnicas para el procesamiento de consultas con datos espaciales biomédicos. SOGM-ING14-I VIEP (2014)
12. Somodevilla, M.: Desarrollo de Técnicas para Poblado Automático de Ontologías Espaciales Biomédicas. SOGM-ING15-I VIEP (2016)

María J. Somodevilla, Andrea M. P. Tamborrell, Ivo H. Pineda, Concepción Pérez de Celis

13. Werry, J.S., Aman, M.G.: Methylphenidate and haloperidol in children: Effects on attention, memory and activity (1975)
14. Duncan, J., Eilbeck, K., Narus, S.P., Clyde, S., Thornton, S., Staes, C.: Building an Ontology for Identity Resolution in Healthcare and Public Health. *Online journal of public health informatics* 7 (2) (2015)
15. He, Y., Sarntivijai, S., Lin, Y., Xiang, Z., Guo, A., Zhang, S., Smith, B.: OAE: The ontology of adverse events. *Journal of biomedical semantics*, 5(1) (2014)

Impreso en los Talleres Gráficos
de la Dirección de Publicaciones
del Instituto Politécnico Nacional
Tresguerras 27, Centro Histórico, México, D.F.
noviembre de 2016
Printing 500 / Edición 500 ejemplares

