



Research in Computing Science

ISSN: 1870-4069

Vol. 48
November 2010

Research in Computing Science

Series Editorial Board

Editors-in-Chief:

Grigori Sidorov, CIC-IPN, Mexico
Gerhard X. Ritter, University of Florida, USA
Jean Serra, Ecole des Mines de Paris, France
Ulises Cortés, UPC, Barcelona, Spain

Associate Editors:

Jesús Angulo, Ecole des Mines de Paris, France
Jihad El-Sana, Ben-Gurion Univ. of the Negev, Israel
Alexander Gelbukh, CIC-IPN, Mexico
Ioannis Kakadiaris, University of Houston, USA
Petros Maragos, Nat. Tech. Univ. of Athens, Greece
Julian Padget, University of Bath, UK
Mateo Valero, UPC, Barcelona, Spain

Editorial Coordination:

Blanca Miranda Valencia

RESEARCH IN COMPUTING SCIENCE, Año 9, Volumen 48, Septiembre de 2009, es una publicación mensual, editada por el Instituto Politécnico Nacional, a través del 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, Ciudad de México, Tel. 57 29 60 00, ext. 56571. <https://www.rcs.cic.ipn.mx>. Editor responsable: Dr. Grigori Sidorov. Reserva de Derechos al Uso Exclusivo del Título No. 04-2004-062613250000-102. ISSN: 1665-9899, otorgado por el Instituto Nacional del Derecho de Autor. Responsable de la última actualización de este número: el Centro de Investigación en Computación, Dr. Grigori Sidorov, Av. Juan de Dios Bátiz S/N, Esq. Av. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, C.P. 07738. Fecha de última modificación 7 de Septiembre de 2009.

RESEARCH IN COMPUTING SCIENCE, Year 9, Volume 48, September 2009, is a monthly publication edited by the National Polytechnic Institute through the Center for Computing Research. Av. Juan de Dios Bátiz S/N, Esq. Miguel Othón de Mendizábal, Nueva Industrial Vallejo, C.P. 07738, Mexico City, Tel. 57 29 60 00, ext. 56571. <https://www.rcs.cic.ipn.mx>. Editor in charge: Dr. Grigori Sidorov. Reservation of Exclusive Use Rights of Title No. 04-2004-062613250000-102. ISSN: 1665-9899, granted by the National Copyright Institute. Responsible for the latest update of this issue: the Computer Research Center, Dr. Grigori Sidorov, Av. Juan de Dios Bátiz S/N, Esq. Av. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, C.P. 07738. Last modified on September 7, 2009.

Advances in Computing Science



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



Instituto Politécnico Nacional, Centro de Investigación en Computación
México 2010

ISSN: 1870-4069

Copyright © Instituto Politécnico Nacional 2010
Formerly ISSN: 1665-9899

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

Electronic edition

Table of Contents

	Page
Search for Places on the Web for Wordnet Synonyms	5
<i>Rafael Guzmán Cabrera</i>	
Hidden Markov Framework for Lexical Tagging.....	15
<i>C. Razo Hernández, J.M. Benedí</i>	
Uso de Sistemas Multiagente en Órdenes de Traba de Manufactura.....	19
<i>José Antonio Gordillo Sosa</i>	
Espectroscopia eléctrica aplicada para determinar osteoporosis en muestras óseas	25
<i>Fernando Ireta M., René Martínez C., Eduardo Morales S., Bárbara González R.</i>	

Search for Places on the Web for Wordnet Synonyms

Rafael Guzmán Cabrera

Universidad de Guanajuato,
Facultad de Ingeniería Mecánica, Eléctrica y Electrónica,
Mexico

guzmanc@salamanca.ugto.mx

Abstract. In this paper, we present a method that allows us, for a polysemic word given in English, to find collocations linked in a significant way with each of its senses. The synonyms that compose every sense are used as pattern of search in the Web with the purpose of creating a corpus by sense. Finally, we apply techniques of data mining that allow us to select the most relevant collocations or lexical patterns for every sense. We especially in two types of collocations especially: associations and lexical sequences. The obtained results show us that we can find collocations between words using the Web as linguistic corpus, as well as the feasibility of incorporation of the lexical pattern obtained in systems of word sense disambiguation that can be used in turn for example in machines of translation and of information retrieval.

Keywords: Text mining, text recognizing, WordNet synonyms.

1 Introduction

With the so-called information society, the amount of stored data increases daily, which leads to an increase in the difficulty of processing this information using traditional methods. To overcome this problem, a series of techniques and tools have emerged in recent years that facilitate advanced data processing, as well as in-depth automatic data analysis. One of these tools is data mining, whose key idea is that data contain more hidden information than meets the eye. Data mining can be defined as the non-trivial extraction of implicit, previously unknown, and potentially useful information from data (Frawley, 1992). Web mining, on the other hand, focuses on the use of data mining techniques to automatically discover and extract information from Web documents and services (Etzioni, 1996). Web mining can be classified into three main areas:

- Web usage mining: this method attempts to extract information (habits, user preferences, or content and relevance of documents) from the sessions and behavior of users and browsers, that is, it allows for the discovery of website access patterns.
- Structured Web mining: Focused on discovering a model from the topology of network links. This model can be useful for classifying or grouping documents.

- Web content mining: Allows you to find common information from web documents. It can be further classified as:
- Text mining: If the documents are in plain text.
- Hypertext mining: Whether documents contain links to other documents or to themselves.
- Markup mining: if the documents are structured, that is, with marks.
- Multimedia mining: if the documents contain images, audio or video.

This paper focuses on the latter, mining web content, using the web as a linguistic corpus for extracting collocations between words. A corpus is a collection of text in electronic format that can be processed by a computer for various purposes, such as linguistic research and language engineering (Leech, 1997).

Corpus linguistics involves the study of languages based on examples of their use. We used Google as our search engine, although recent research conducted to study the possibility of using the web to disambiguate nouns preceded by an adjective (Rosso, 2005) shows that the results do not depend much on the search engine used. Several investigations have been carried out using Web mining as a tool.

For example, (Mihalcea, 2004) presents the main lines of research regarding the exploitation of the Web as a linguistic resource in Word Sense Disambiguation (WSD) systems.

Furthermore, (Celina, 2003) has used the Web to enrich tagged corpora, which then facilitate the task of WSD. There are several reasons for using the Web as a linguistic resource, among them that it provides a means of quickly and easily accessing a wide variety of information stored in electronic format in different parts of the world. Furthermore, it is free and available with a mouse click.

The size of the Web in July 1999 was estimated at 56 million addresses, 125 million in January 2001, and 172 million in January 2003. This represents a massive growth of over 300% in just under five years. In 1999, 800 million indexed Web pages were available; if we estimate the average Web page size to be between 7 and 8 kilobytes of plain text, we have nearly 6 terabytes of available text in 1999 and approximately 30 terabytes in 2003.

With these figures, the Web is clearly an immense corpus, given the amount of information available to us. Furthermore, the Web is multilingual, since approximately: 71% of the pages are written in English, 6.8% in Japanese, 5.1% in German, 1.8% in French, 1.5% in Chinese, 1.1% in Spanish, 0.9% in Italian, and 0.7% in Swedish, the remaining 11.1% is spread across other languages and dialects (Kilgarriff, 2003). In Natural Language Processing (NLP) research, the use of corpora is important to extract language models, such as combinations of meaningful words that allow knowing which words are related to each other or which of them are from a certain domain.

However, the Web has several negative aspects, including the fact that the information found is very heterogeneous and disorganized, and there is a lot of junk information or information with tags that make it difficult to process. In addition, you can't be sure that everything you find is correct, since no one checks it. But thanks to the Web's redundancy, accurate information usually prevails.

The task of automatically finding semantic relationships between adjacent words has attracted the attention of many researchers in the field of natural NLP in recent decades. As a result of this research, important dictionaries of English collocations have been written, for example, the one created by (Benson, 1989).

There are also developed systems that allow analysis between collocations, such as the N-Gram Statistics Package. This system allows for the analysis of words in a corpus, such as the counting of frequent occurrences and various statistical measures that allow for an association between two or more words.

The possible applications of collocation detection and its relationship to the lexical meaning of the sentence are many and varied. Examples include translation from one language to another and the integration of lexical patterns into WSD systems.

In Section 2, we describe the method used to find meaningful collocations on the Web for a given word. In Section 3, we present the results of the experiments and some examples of the lexical patterns found by meaning. Finally, in Section 4, we present the conclusions of this work.

2 Methodology

In natural language, there are many combinations of words that frequently co-occur and correspond to a particular use of a word or a sense of a sentence. These combinations can be presented as an unbroken sequence of words, in this case simply called sequences, or the words in the combination may not occur contiguously in context, in this case called associations. Sequences and associations are two of the most important types of word collocations. The method used to obtain meaningful collocations consists of the following steps:

1. **WordNet Senses.** For a given polysemous word in English, we obtain its meanings from the WordNet lexical database, which is a lexical-conceptual database for English structured as a semantic network, so that access to lexical information is not restricted to merely alphabetical access. This is inspired by psycholinguistic theories about human lexical memory. WordNet stores information on words belonging to the syntactic categories of noun, verb, adjective, and adverb. The cost of having syntactic categories is a large amount of redundancy that conventional dictionaries lack.
2. **Synsets.** For each sense, we obtain a set of synonyms (synsets); Words in WordNet are organized into sets of synonyms, or synsets, each of which represents a different lexical concept. Each synset contains the list of synonymous words, as well as information on semantic relationships established with other words or synsets.
3. **Snippets.** Using each synonym (synset element) as a web search pattern and given a search engine (e.g., Google), we download snippets. When we perform a query on the web, the search engine provides us with a response, if the request is found, a set of web pages that match the request, as well as a brief summary of the content of each page so the user can choose the address that best matches the request. Each

set of addresses and summary provided by the search engine, stored in a plain text file, is a snippet. In our case, they represent commonly used expressions from WordNet synonyms for the given polysemous word.

4. Corpus by meaning. The snippets of all synonyms for a given sense are filtered and concatenated into a single file; that is, we create a WordNet corpus per sense. The corpus per sense is constructed by taking five context words from the right and five words from the left of the synonym.
5. Extraction of collocations. We found all significant collocations (sequences or associations) that occur in each corpus by meaning with the following criteria:
 - Force. A sequence or association is relevant if it is frequent, that is, if it occurs more than a predetermined threshold or cutoff value; and is defined by:
 - The threshold or cutoff frequency. It is defined as a value equal to the sum of the average frequency and the standard deviation and is the minimum frequency that words must have to pass this threshold. This ensures that only those occurrences that appear frequently in the contexts of WordNet's meaning are extracted, eliminating all words that may appear randomly.
 - Local dispersion. To find the set of sequences and associations that are representative of each of the WordNet senses of a given word, it is desirable that, in addition to having passed the previous frequency-based filter, they be in the context of all the synonyms that make up the sense. This is precisely what the local dispersion measure does; it allows us to discard those words that are not in the context of all the synonyms that make up the WordNet sense.
 - External dispersion. With this measure, what we seek is that the set of words that exceed the two previous measures is found only in the context of one of the senses, that is, all those words that appear in more than one sense are discarded.

The methodology described in these steps is applied to the extraction of lexical associations and uninterrupted sequences. We now present a description of each of these types of collocations.

Lexical associations.-Lexical associations are a set of words significantly linked to one of the WordNet senses of a polysemous word. To extract these, the synonyms that make up the WordNet sense in the corpus are identified. For each synonym found, the context words are entered into a table. Basically, for each context word found, the question "Does it exist in the table?" is asked. If the answer is no, it is included and its frequency value is initialized to one, while if the answer is yes, its frequency is increased by one. In this way, we go through the entire corpus. From the resulting table, we must select those context words that exceed the strength measure and the dispersion measures mentioned in point 5 of the methodology. It is worth mentioning that the words found in the table are not required to be adjacent to the synonym or between them; their position within the context varies within the defined window. In this way, we find those words that are significantly linked to the meaning.

Sequences.- To obtain uninterrupted word sequences, an automatic iterative process is performed that changes the window size, that is, the number of words taken to the left and right of the synonym, from one to five. For each window size, the process is as

Table 1. Number of snippets downloaded from the Web for the synonyms of Instance.

case	919
instance	924
example	983
illustration	987
representative	987

follows: the word or words are taken respecting their location relative to the synonym, that is, whether they are on the left or right. The result is a set of tables showing the context sequences to the left and right of the synonym for the different window sizes.

Statistics are obtained for each one. As with lexical associations, the resulting uninterrupted sequences are filtered, and only those that are significant in the corresponding WordNet sense are recovered. In this case, the words that make up the uninterrupted sequence are contiguous to each other, respecting their position within the context.

Sequences and associations are found using the redundancy of the Web as a corpus. This is intended to allow future incorporation of the results obtained into WSD systems, whose objective is to associate a word, given in a context, with a definition or meaning that distinguishes it from other meanings attributable to that word. Any NLP system requires a module with these characteristics. WSD is not an end in itself, but rather a necessary step for performing actions such as syntactic analysis or semantic interpretation in NLP tasks, as well as for the development of final applications such as information retrieval (Montes, 2000), text classification (Kosala, 2000), discourse analysis (Montes, 2002), and machine translation (Smrz, 2001), among others.

For example, a traditional information retrieval system will answer the question "What plants live in the desert?" with all documents containing the terms "plants" and "desert," regardless of their meaning. In some of these documents, the term "plant" would appear with the meaning of "living being," while in others, it would mean "industry." If the information retrieval system were able to distinguish the meanings of the query terms, it would return only the documents that use the meaning of "living being." To do this, the system must integrate a WSD module to disambiguate both the query terms and the terms in the indexed documents.

3 Results

We show the results obtained when applying the methodology to the word "instance." We chose this word because, in addition to being polysemous, it has two common synonyms among its meanings, which will allow us to observe the effect of the dispersion measures on the lexical patterns obtained. The WordNet meanings of "instance," as a noun, are:

1. case, instance, example -- (an occurrence of something)
2. example, illustration, instance, representative -- (an item of information that is representative of a type)

Table 2. Abstract of statistics for instance.

	Sense 1	Sense 2
Word: instance	Web	Web
Number of usage examples in the corpus	12684.0	15848.0
Number of distinct words	2831	3590
Average	4.5	4.4
Standard deviation	7.3	7.5
Cut-off frequency (mean + standard deviation)	11.8	11.96
Number of words that exceed measure 1	179	238
Number of words that exceed measure 2	87	67
Number of words that exceed measure 3	25	9

Table 3. Simple lexical associations for Instance

Sense 1		Sense 2
based	java	English
case	learning	free
code	multiple	government
data	net	library
date	number	Link
definition	org	members
documents	our	resources
example	process	section
examples	proposal	software
file	server	
index	use	
instance	will	
It		

Using synonyms for each sense as a Google search pattern, web snippets were downloaded containing examples of context of use, that is, commonly used expressions of the synonyms that make up the corresponding WordNet sense. The number of snippets downloaded per synonym is shown in Table 1.

These snippets formed two corpora, one for each sense. The corpus for sense 1 was formed by concatenating 2,826 snippets, while the corpus for sense 2 was formed with 3,881 snippets; this number is higher because it includes one more synonym. Table 2 shows a summary of the results obtained. In the corpus for sense 1, 12,684 examples of common context use of the synonyms that comprise it were found, while in the corpus for sense 2, 15,848 examples of context use were found.

A total of 2,831 different context words were found in the 12,684 usage examples for sense 1. Of these, 179 exceed the threshold frequency (frequency greater than the average frequency plus the standard deviation). The words that, in addition to having exceeded the threshold frequency, are found in the context of all synonyms for sense 1

Table 5.- Sequences to the left of instance

Instance-1 Sequences		Instance-2 Sequences	
Customers	Bottle	Design	Page layout
Home	Us party	art	Visual arts
Yam	The bottle	Bouchard	Of the mouth
Party	Studies customers	Medical	Proactive core component
Resources	To the	The	Multimedia design at
This	The us party	Fanny Bouchard	Object web proactive core component
To	In the bottle	Graphic design	Illustration of the mouth
Tools			

are 87. In the end, we have 34 words that could help us disambiguate the word instance (25 for sense 1 and 9 for sense 2); these words are shown in Table 3.

In simple lexical associations, the significant words encountered do not necessarily have to appear contiguously. To illustrate the use of significant words associated with instance, we present some commonly used sentences. To reinforce the idea of using these words in lexical disambiguation systems, we separate the examples by meaning. The first meaning of instance relates to "the occurrence of something," while the second sense refers to "an item of information that is representative of a type."

Sense 1:

...another instance of the same process already running on the current machine....

...Enforcing a rule that only one instance of process is running is an interesting task....

Sense 2:

...ACTIVITY in this instance involves the use of government facilities and equipment for ...

.... Another difference between instance members and class members is that class

Since the examples presented are extracts from sentences, coupled with the fact that we are talking about ambiguous words, it may be difficult to clearly and concisely distinguish between one meaning and the other. In this case, it would be worthwhile to increase the number of immediate context words surrounding the ambiguous word. However, this task is not easy, even manually. It should be noted that in the Senseval-2 competition (a competition that compares the performance of different WSD systems in different languages), there was only 75% agreement between human annotators for English.

Table 6.- Sequences to the right of instance.

Sequences Instance-1		Sequences Instance-2	
Design	Code	And	A formal example with
Edu	Western reserve	Clients	Of the mouth
Index	Studies case	In	Livres d enfants children
Law	Studies catalog	For	Employees post a jobs
Studies	Studies in	Is	Of the mouth illustration
Study	Western reserve university	Of the	A formal example with a
Western	Studies catalogs resources	Of a	Government by john Stuart
		In the	Of the secretary general
		Livres d enfants	Livres d enfants children book's
		Employees post a	Employees post a jobs and
		Of judicial activism	Government by john Stuart mill
		And fine art	Of the mouth illustration of
		Of the secretary	
		Government by john	
		Of next at	

For uninterrupted word sequences around the instance senses, the results obtained for different window (V) values are shown in Table 4, as well as the number of different and significant sequences. Negative values in the window column represent the number of words taken to the left of the synonym.

As the window size increases, the number of sequences decreases. This is because the sequence (of one, two, or more words) must be part of the context of all synonyms and also have the same order of appearance. The correlation between the number of sequences and significant sequences is 0.94. This value tells us that as we increase the

number of sequences, the number of significant sequences will increase, which is to be expected since it presents a nearly normal distribution. For a sequence that has passed the strength and dispersion measures, the higher the frequency, the more significant it will be (Guzmán, 2005a).

Unbroken sequences typically begin or end with a polysemous word, instance in our case. For this reason, we separated the results obtained, differentiating the unbroken sequences not only by direction but also by their left or right location. Table 5 shows the unbroken sequences on the left, which are found in commonly used expressions such as customers instance or graphic design instance.

The significant sequences to the right of instance are shown in Table 6. We will find these sequences in commonly used expressions, such as instance design or instance studies case.

In our work, we have not ignored stop words, or empty words, such as prepositions, determiners, etc., which appear in almost every sentence. However, some of these words play an important role in assigning meaning to a sentence (Guzmán, 2005b). For example, in the case of "for," we found it significantly associated with sense 2 of "instance" and is used in commonly used expressions such as "for instance"; this sequence has a single meaning in WordNet (its meaning refers to an example).

4 Conclusions

The initial experiments conducted demonstrate the potential of the Web as a linguistic corpus. Furthermore, they demonstrate the feasibility of incorporating the extracted lexical patterns into disambiguation systems. The methodology can be applied to other morpho-syntactic categories, as well as to other languages, provided that a lexical database exists in these languages, such as WordNet for English, that allows us to determine the meanings attributable to a polysemous word. Furthermore, it can be applied to finite corpora. Although the number of significant collocations per meaning is generally encouraging, we must increase the size of the corpus to have more examples of contexts of use for each of the synonyms that make up the meaning in WordNet and thus obtain better lexical patterns

References

1. Benson, M.: The BBI combinatorial dictionary of English. John Benjamin Publ., Amsterdam, Philadelphia (1989)
2. Benson, M.: Collocations and General Purpose dictionaries. *International Journal of Lexicography*, 3, pp. 23–35 (1990)
3. Buscaldi, D., Rosso, P., Masulli, F.: The Upv-unige-CIAOSENSE WSD System. In: *Proceedings of Senseval-3*, pp. 77–82 (2004)
4. Celina, S., Gonzalo, J., Verdejo, F.: Automatic Association of Web Directories with Word Senses. *Computational Linguistics*, 29(3), pp 485–502 (2003)
5. Choueka, Y.; Klein, S.T., Neuwitz, E.: Automatic retrieval of frequent idiomatic and collocational expressions in large corpus. *Association for Literary and Linguistic Computing Journal*, 4(1):34–38 (1983)

6. Etzioni, O.: The World Wide Web Quagmire or Gold Mine? *Communications of the ACM*, 39(11), pp. 65–68 (1996)
7. Fellbaum, Ch.: *WordNet as Electronic Lexical Database*. MIT Press (1998)
8. Frawley, W., Piatesky-Shapiro, G.: *Knowledge Discovery in Databases: An Overview*, *AI Magazine*, pp. 213–228 (1992)
9. Guzmán-Cabrera, R., Rosso, P., Montes-y-Gomez, M., Gomez-Soriano, J. M.: Mining the Web for word sense discrimination, In: *Information and communication technologies international symposium*, Tetouan, Morocco (2005)
10. Guzmán-Cabrera, R., Montes-y-Gomez, M., Rosso, P.: Searching the Web for word sense collocations. In: *IADIS international conference*, Algarve, Portugal (2006)
11. Kilgariff, A., Greffentette, G.: Introduction to the Special Issue on Web as Corpus, *Computational Linguistics* 29(3), pp. 1–15 (2003)
12. Kosala, R., Blockeel, H.: Web mining research: a survey. *GIS KDD Explorations*, pp. 1–15 (2000)
13. Leech Garcide, G., McEnery, T.: Corpus annotation. *Linguistic Information from Computer Text Corpora, Grammatical Tagging*. chap 2, pp. 19–33 (1997)
14. Mihalcea, R.: Making Sense out of the Web. In: *Workshop on Lexical Resources and the Web for Word Sense Disambiguation*, IBERAMIA, Mexico (2004)
15. Montes-y-Gómez, M., López-López A., Gelbukh, A.: Information Retrieval with Conceptual Graph Matching. In: *11th International Conference on Database and Expert Systems Applications DEXA 2000*, Springer-Verlag (2000)
16. Montes y Gómez, M.: *Text Mining using Similarity between Semantic Structures*. Doctoral Thesis, Computing Research Center (CIC), National Polytechnic Institute (IPN), Mexico (2002)
17. Resnick, P. S.: *Selection and Information: A Class-Based Approach to Lexical Relationship*. Doctoral Thesis, University of Pennsylvania (1993)

Hidden Markov Framework for Lexical Tagging

C. Razo Hernández¹, J. M. Benedí²

¹ Universidad de Guanajuato,
Facultad de Ingeniería Mecánica, Eléctrica y Electrónica,
Mexico

² Universidad Politécnica de Valencia,
Departamento de Sistemas Informáticos y Computación,
España

jbenedi@dsic.upv.es

Abstract. In this work we present a lexical tagger for English using hidden Markov models with a probabilistic model of distribution of words in categories. The corpus used and set of labels corresponding to the categories is described. After, the model use is described and the form which it is estimated. Finally, we show the realized experiments and the obtained results.

Keywords: Tagger, POS tags, HMM, NLP.

1 Lexical Tagged

Tagged is the assignation of category to which the words in a corpus belong (POST = Part-Of-Speech Tagging). Its purpose is help to improve applications of the natural language processing in which its required to know the sense of the words, automatic translation, information retrieval, text classification and extraction of information among others.

The corpus used in the experiments realized in this work is the part of the Wall Street Journal that has been processed in the project Penn Treebank. It contains approximately a million words distributed in 25 directories. It was automatically labeled, analyzed and manually reviewed. The size of the vocabulary is greater to 49,000 words and the set of POS tags is 45 tags. For the experiments, the corpus is divided in training corpus (directory 00-20), tuning corpus (directory 21-22) and test corpus (directory 23-24).

2 Proposed Model

The used model, combines a probabilistic model of distribution of words in categories with hidden Markov models. For the probabilistic model, a list of words is used in which appears each word and the frequency of assignation of category (C_w). The hidden Markov model used is a model of the left to right.

In order to find the set of lexical labels corresponding to a sequence of words, the sequence of observable symbols is calculated (labels) that will be emitted by the most

probable states sequence. This calculation is carried out using HMM and the distribution word-category as in the equation (1):

$$\tilde{C} = \operatorname{argmax}_{s,e} \prod b_1(e_1) \Pr(w_1|e_1) a_{12} b_2(e_2) \Pr(w_2|e_2) a_{23} \dots a_{n-1,n} b_n(e_n) \Pr(w_n|e_n). \quad (1)$$

$\Pr(w_i|e_i)$ calculates using distribution word-category C_w and the rest with the hidden Markov model. A modification to the Viterbi algorithm that allowed us to calculate the expression (1) and to obtain the corresponding set of labels [1].

The estimation of the Markov model and the model based on categories is made separately by simplicity. The hidden Markov model is training with a corpus tagged. A detailed description of the methods of training for HMM can be found in [3].

The parameters of the distribution word-category, $C_w = \Pr(w|e)$, calculates agreement with the equation [2]:

$$\Pr(w|e) = \frac{N(w, e)}{\sum_w N(w', e)}, \quad (2)$$

where $N(w, e)$ is the number of times that the word w has been tagged with the POStag e and $\sum_w N(w', e)$ is the sum of all the words that have been tagged with that POStag.

The word w can belong to different categories. It can also occur, that in the training set does not appear a word and therefore its probability $\Pr(w|e)$ is not defined. This is solved adding the term $\Pr(UNK|e)$ to all the categories, which represents the probability for words unknown in the test set. In order to estimate this probability, three approaches are used.

The first approach consists of assigning a small probability equal to $\Pr(UNK|e)$ for all the categories.

The second approach consists of the supposition of that the distribution $\Pr(UNK|e)$ in the test set, is very similar to the corresponding one in a tuning set. For that reason, the distribution of the frequency of appearance of the words unknown tagged with $\Pr(UNK|e)$ in this set of tuning is considered. For the possible null values of the considered distribution, a very small probability is added. These first two approaches are described in [1].

The third approach, is based on a study of Demartas and Kokkinakis [4], they concludes that the probability distribution of the unknown words is very similar to the one of which they frequently is equal to 1 and very different from the distribution of the well known words. Based on this, the estimation becomes by means of the equation (3):

$$\Pr(w|e) = \frac{\Pr(e|w^{unkonwn})P(w^{unkonwn})}{P(e_i)}. \quad (3)$$

$\Pr(e|w^{unkonwn})$ and $P(e_i)$ are calculated from the training set and $P(w^{unkonwn})$ it is calculated from the tuning set. For the null values of the considered distribution, a very small probability is added.

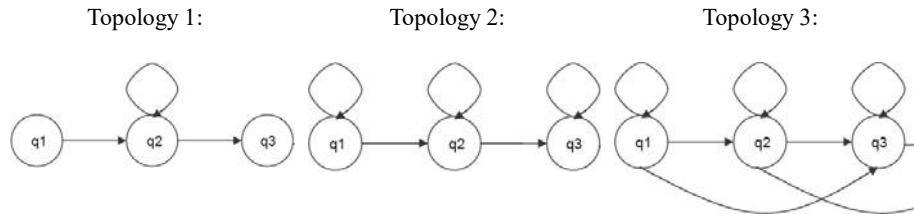


Fig. 1. Topologies used in the HMM.

Table 1. Baseline for both used vocabularies.

Size Vocab.	Precision	Exactitude
1000	80.46%	80.46%
37075	85.63%	85.63%

Table 2. Result of the best model.

Model	Precision	Exactitude
C_w	80.46%	80.46%
$\square, C_w$	87.50%	87.21%

Table 3. Results for test corpus.

Model	Precision	Exactitude
C_w	81.12%	81.12%
$\square, C_w$	87.25%	86.79%

Table 4. Better result for vocabulary complete.

Model	Precision	Exactitude
C_w	85.63%	85.63%
$\square, C_w$	91.69%	91.69%

Table 5. Results for test corpus with complete vocabulary.

Model	Precision	Exactitude
C_w	84.09%	84.09%
$\square, C_w$	90.02%	89.13%

3 Results

In order to find the best model, experiments were made using different topologies and number of states in the hidden Markov model. In addition two sizes of different vocabularies were used considering the three approaches for unknown words.

Like baseline, the so large distribution was made tagged of corpus of test using only word-category for of vocabularies both used in table 1.

1. Experiments on topology and states:

Experiments were done changing the number of states in the HMM and using 3 different topologies.

2. Experiments on approaches of unknown words:

Considering the 3 topologies, different number of states and using the three approaches for unknown words, the best result is obtained using topology 1, 8 states in the HMM and using approach 3 for not known words as it is in table 2. Using this model, test corpus was tagged; the obtained results we show in table 3.

3. Experiments with complete vocabulary:

Using a vocabulary of 37,075 words, the results improve really. The best model is obtained with topology 3 and 3 states for the HMM and using approach 3 for unknown words, the results are in table 4.

4 Conclusions

The obtained results are good comparing with baseline. The precision was 81.12% like baseline, value that our model improves up to 87.25% for the vocabulary of 1000 words. For the vocabulary of 37,075 words, baseline is 84,09% and our model it improvement up to 90.02%.

References

1. Sánchez, J.A., Nevado, F., Benedí, J.M.: Lexical decoding based on the combination of category-based stochastic models and word-category distribution models. In: Proceeding of COLING (2006)
2. Sánchez, J.A., Benedí, J.M.: Combination of n-grams and stochastic context-free grammars for language modelling. In: Proceeding of COLING (2000)
3. Juang, B.H., Rabiner, L.: Fundamentals of Speech Recognition. Prentice-Hall (1993)
4. Molina, A.: Desambiguación en procesamiento del lenguaje natural mediante técnicas de aprendizaje automático. Technical report, Universidad Politécnica de Valencia (2004)

Uso de Sistemas Multiagente en Órdenes de Traba de Manufactura

José Antonio Gordillo Sosa

Universidad Tecnológica del Suroeste de Guanajuato,
México

antgor@antoniogordillo.com

Resumen. En este trabajo de investigación se presenta el proceso de desarrollo e implementación de un algoritmo de distribución de recursos aplicado a la programación de un sistema multi-agentes, en un entorno de fabricación de piezas cerámicas. La idea central de la misma gira en torno a la búsqueda del balance entre la distribución de órdenes de trabajo aplicando modelos de temperatura entre los agentes. Se desarrolló una simulación completa de la tarea anterior aplicando la herramienta de programación de agentes JADE.

Palabras clave: Multi-agentes, distribución dinámica, algoritmo de optimización, programación de órdenes de trabajo.

Application of Multi-Agent Systems in Manufacturing Work Orders

Abstract. This research paper presents the developing and implementation process of an algorithm applied to the scheduling of a multi-agent system in a ceramic manufacturing environment. The main focus revolves around the search for a balance in the distribution of work orders by applying temperature models among the agents. A complete simulation of the aforementioned task was developed using the JADE agent programming tool.

Keywords: Multi-agent, dynamic distribution, optimization algorithm, work order scheduling.

1. Introducción

Una de las actividades más importantes y al mismo tiempo, la más compleja a ser desarrollada dentro de los sistemas de manufactura actuales, es la asignación de tareas. Se requiere que sea lo suficientemente dinámica y adaptable como para enfrentar fallas mecánicas, órdenes emergentes y un sinnúmero de eventualidades más. Actualmente,

la industria está experimentando cambios muy importantes en lo que se refiere a la aplicación de elementos de software al interior de los procesos productivo.

Los agentes inteligentes integran una de las áreas con más rápido crecimiento en este rubro. En la red Internet, por ejemplo, un agente de software (también conocido como agente inteligente) se identifica como una aplicación de cómputo que puede recolectar información o realizar servicios diversos sin necesidad de contar con una presencia real del usuario.

Las características principales de estos agentes de acuerdo a Jennings & Wooldridge [1] son:

- Autonomía: los agentes deben ser capaces de desarrollar la mayoría de sus tareas sin requerir la intervención directa de humanos y tener un cierto grado de control sobre sus propias acciones y su propio estado interno.
- Habilidad Social: los agentes deben ser capaces de interactuar con otros agentes.
- Responsividad: los agentes deben percibir su entorno, y responder con suficiente oportunidad a los cambios que puedan ocurrir en él.
- Proactividad: en el lapso de emisión de respuesta hacia el entorno, los agentes deben mostrar un comportamiento enfocado a cubrir sus metas y tomar la iniciativa en aquellos casos en los que sea apropiado.

Los agentes se utilizan regularmente en aplicaciones tales como administración personalizada de la información, de procesos industriales, de transacciones de comercio electrónico entre otros, haciendo énfasis especialmente en la administración del flujo de trabajo en un entorno de manufactura.

Una vez implementado un SMA (Sistema Multiagentes) en este entorno de trabajo, las características del mismo, tales como interacción local y dinámica no lineal – entre otras – permitirán que la distribución de órdenes a ser procesadas pueda realizarse en tiempos y con un número total de fallas cada vez menores.

El algoritmo implementado, presentado en el presente trabajo, se centra en tratar de reducir el tiempo de búsqueda del espacio que pueda recibir las piezas fabricadas, tratando de generar un equilibrio en la distribución de las órdenes.

Para lograr lo anterior, se utilizan conceptos tales como Temperatura, Calor Relativo y Calor Latente, los cuales se utilizan para programar los agentes en el sistema y, de este modo, revisar el impacto sobre el rendimiento global de la operación.

2. Estado del arte

El objetivo de este trabajo es el diseño, simulación y ejecución de la programación de órdenes de trabajo en las líneas de producción de una empresa del sector cerámico. Se propone una arquitectura de agentes JADE [4] integrados al área de prensado como medio de enlace para distribuir dinámicamente los diferentes lotes producidos hacia el área de esmaltado, programados en base a un algoritmo de equilibrio de cargas.

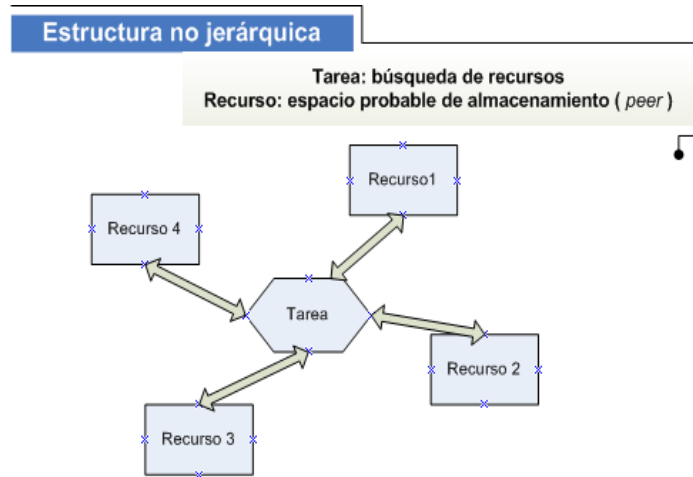


Fig. 1. Entorno más con ardillas.

3. Metodología

Debido al tipo de aplicación dentro del sistema, el estudio supone una estructura descentralizada, no jerárquica, en la que se programará el comportamiento individual de cada agente en el sistema con base en su autonomía y cooperación interna (Ver figura 1).

La comunicación entre agentes se basa en la programación de comportamientos (Behaviours en el lenguaje de programación) mediante los cuales se simula la distribución de órdenes de trabajo del área de producción. Con esto, cada agente interactúa con el resto de manera dinámica, distribuyendo las piezas producidas en las líneas con capacidad de recepción.

El entorno debe adaptarse a la distribución dinámica de tareas. Al aplicar la programación de agentes, se tendrán dos comportamientos a ser observados: el del agente individual y el del sistema en general.

Para el agente individual, el objetivo será minimizar el tiempo de respuesta y en el caso de este algoritmo específico, la indicación correcta del estado de su temperatura.

El sistema deberá tender a una distribución uniforme de la carga de órdenes generadas sobre todos los agentes disponibles y con esto, estará mejor preparado para aceptar trabajo adicional y adaptarse más adecuadamente a fallas surgidas durante la operación.

La herramienta de desarrollo para implementar esta estrategia es la plataforma de agentes JADE, utilizando como base de distribución de tareas un algoritmo de optimización que asigna valor de temperatura a cada agente en el sistema, dependiendo de su carga de trabajo.

4. Entorno de trabajo

En torno a la sección de prensado, se asigna un agente en cada línea de trabajo. Cada uno de ellos, recibirá la orden de trabajo enviada por un agente de distribución y la aceptará o rechazará de acuerdo a si “temperatura”, es decir, a la carga de trabajo que tenga en ese momento.

El agente de distribución maneja un “pizarrón de estado” en el cual se escriben las temperaturas de cada agente en el sistema. Ésta será la base mediante la cual la distribución de tareas tratará de mantenerse en equilibrio. Cuando un agente esté “caliente” el nivel calorífico correspondiente se asentará en el pizarrón y se buscará asignar el trabajo a un agente con menor temperatura, lo cual ayudará al proceso de “enfriamiento” del sistema.

La estructura anterior se implementa mediante tecnología de Agentes JADE con las siguientes características:

- Comunicación programada. Los agentes distribuyen las órdenes a lo largo de la red interna mediante protocolo TCP/IP
- Dinámica no-lineal. El agente de distribución lee la temperatura tanto individual como colectiva en el sistema y se decide por la opción de menor energía interna, que puede ser diferente en cada revisión.
- Reconfiguración sobre demanda. El lenguaje de programación permite ajustar, incluso en tiempo real la estructura multiagente para responder a las diferentes situaciones que puedan presentarse.

5. Experimentación

El algoritmo de solución propuesto, busca implementar una solución adecuada al problema de la distribución de órdenes de trabajo, mediante la comunicación y posterior comunicación de cada agente entre sí al interior del SMA. Las características generales implementadas en este sistema son:

- Generación aleatoria del total de órdenes de trabajo a ser procesadas, así como valores de referencia de las mismas. Estos últimos representan el nivel de prioridad de las tareas a desarrollarse.
- Generación aleatoria de valores de temperatura asociados a cada uno de los agentes en el sistema, referenciando con ellos su nivel de carga de trabajo. Comparativamente, para un mayor valor generado se tendrá un agente con más actividad y por tanto, con una temperatura más alta. Un valor bajo indicará un agente con menor temperatura, es decir, tendrá un menor nivel de carga en el sistema al momento de buscar opciones para la asignación de trabajos a ser procesados.

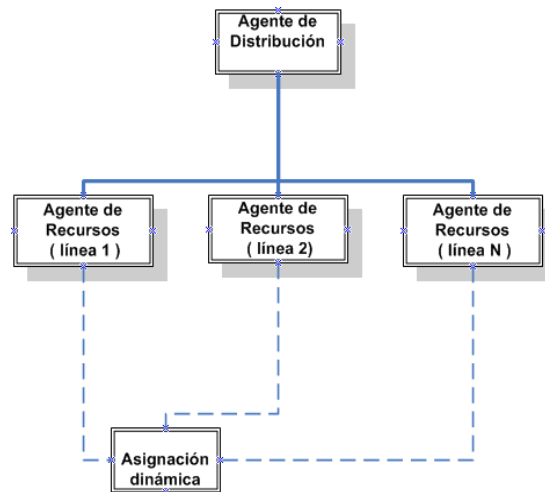


Fig. 2. Distribución de agentes en el sistema.

- Al terminar el proceso de generación y asignación de valores de carga y temperaturas, el sistema realiza una autoevaluación en busca de un equilibrio general: al presentarse un agente con mayor carga de trabajo (y por consecuencia, un incremento de temperatura) distribuye una parte del mismo hacia otro con menor carga. Esto permite una distribución más adecuada de las cargas de trabajo en el área de procesamiento.

En este contexto, el simulador de almacenamiento de recursos permite manejar los parámetros siguientes:

- Cantidad de piezas para distribuir. Desde 1 hasta n piezas.
- Cantidad de espacios de almacenamiento. Desde 1 hasta n espacios.
- Capacidad de almacenamiento en cada espacio. El simulador puede adaptarse a la capacidad específica requerida del sistema y, en tiempo de ejecución, comunicar a cada agente la capacidad admitida en ese momento.
- Cantidad de piezas almacenadas en cada espacio. Varía durante la ejecución. Al igual que la capacidad de almacenamiento, el simulador comunicará en el momento requerido, la cantidad almacenada.

7. Conclusión

En esta etapa del estudio, no se pretende hacer un desarrollo complejo o crear mecanismos de alto nivel para cooperación y negociación, solamente centrarse en un mecanismo simple, confiable y práctico para programación dinámica de tareas, debido a que se intenta utilizar como parte de un diseño integral basado en agentes y sistemas de manufactura para una fábrica real. Con esto se simplifica la complejidad del sistema,

se reduce el retardo de comunicación y por consecuencia, el tiempo de programación/reprogramación.

Los resultados obtenidos demuestran que integrando un esquema SMA acoplado al área de proceso bajo el algoritmo de Equilibrio de Temperaturas, se logra desarrollar un esquema de trabajo real con un desempeño eficiente de distribución de tareas.

Asimismo, una vez observado lo anterior, puede predecirse que es posible responder a posibles situaciones de emergencia, tales como:

- Reconfiguración para nuevos pedidos de emergencia Reconfiguración tras cancelación de pedidos.
- Reconfiguración para optimización de la programación actual Reconfiguración por ruptura de la máquina o producto destruido.

8. Trabajos futuros

Como líneas futuras de investigación dentro de este trabajo pueden destacarse las siguientes:

- Implementación y comparación de eficiencia de otros algoritmos de programación como medida de reducción de tiempos de proceso.
- Desarrollo de interfaz gráfica del simulador para una mayor facilidad de uso.
- Reestructuración del código para permitir la introducción manual de datos en tiempo real, para enfrentar contingencias.
- Estudio detallado de requerimientos mínimos de *hardware* y *software* que permita un funcionamiento óptimo del simulador.
- Implantación en la fábrica de un sistema multiagente para la programación de producción en tiempo real.

Referencias

1. Jennings, N.R., Wooldridge, M.: Software Agents. IEEE Review, pp. 17–20 (1996)
2. Poli, R., Di Chio, C., Langdon, W.: Exploring Extended Particle Swarms: a Genetic Programming Approach. In: Proceedings of the conference on Genetic and Evolutionary Computation, pp. 169–176 (2005)
3. Parsopoulos, K.E., Vrahatis, M.N.: Evolutionary Computing and Optimization: Particle Swarm Optimization Method in Multiobjective Problems. In: Proceedings of the ACM Symposium on Applied Computing (2002)
4. Chien-Chung, S., Chaiporn, J.: Ad Hoc Multicast Routing Algorithm with Swarm Intelligence. Mobile Networks and Applications, 10(1-2) (2005)
5. JADE: Java Agent Development Framework (2005) <http://jade.tilab.com>.
6. Camorlinga, S., Barker, K.: Multiagent Systems Storage Resource Allocation in a Peer-to-Peer Distributed File System (2003)

Espectroscopia eléctrica aplicada para determinar osteoporosis en muestras óseas

Fernando Ireta M., René Martínez C., Eduardo Morales S.,
Bárbara González R.

Universidad de Guanajuato,
División de Ingenierías campus Irapuato-Salamanca,
México

`fireta@salamanca.ugto.mx`

Resumen. La osteoporosis es considerada como la pérdida de calcio en los hueso, en el presente trabajo se muestran los resultados al desarrollar una técnica que emplea espectroscopia de impedancia eléctrica para determinar la pérdida de calcio en muestras óseas, las cuales se sometieron a un proceso de descalcificación químico controlado, a fin de simular el proceso de osteoporosis o pérdida de calcio y probar que es posible medir una variación de la pérdida de calcio en dichas muestras, empleando la técnica de Espectroscopia de Impedancia Eléctrica, la cual relaciona los parámetros del valor real y reactivo de la impedancia con respecto a la frecuencia de la señal de entrada, en un diagrama Cole-Cole.

Palabras clave: Espectroscopia eléctrica, osteoporosis.

Electrical Spectroscopy Applied to Determine Osteoporosis in Bone Samples

Abstract. Osteoporosis is considered as the loss of calcium in the bones, in the present work the results are shown by developing a technique that uses electrical impedance spectroscopy to determine the loss of calcium in bone samples, which were subjected to a process of controlled chemical decalcification, in order to simulate the process of osteoporosis or calcium loss and prove that it is possible to measure a variation of the loss of calcium in said samples, using the technique of Electrical Impedance Spectroscopy, which relates the parameters of the real and reactive value of the impedance with respect to the frequency of the input signal, in a Cole-Cole diagram.

Keywords: Electrical spectroscopy, osteoporosis.

1. Introducción

La osteoporosis ya no es considerada una enfermedad que afecta solo a las mujeres ni una consecuencia natural del envejecimiento; porejemplo en los hombres se sabe que son más los que presentan fracturas de la cadera debidas a osteoporosis que las mujeres. La pérdida mineral ósea se produce tanto en el hueso trabecular como en el cortical y la fortaleza de un hueso depende de la cantidad de masa mineral que posee. Una disminución en la densidad mineral ósea (DMO) causa la subsecuente reducción de la fuerza del hueso y un incremento del riesgo de fractura.

Las técnicas de medición más comunes del DMO, puede ser Ultrasonido o por Rayos X, por lo que se propone una nueva técnica para la medición de la pérdida de calcio que es por medio de un método no invasivo a través de Espectroscopia de Impedancia Eléctrica, y en el estado del arte no hay un estudio similar al propuesto en este proyecto.

La Espectroscopia de Impedancia Eléctrica es una técnica no invasiva que se basa en el principio de que los tejidos se comportan como conductores o dieléctricos de la corriente eléctrica dependiendo de su composición [1], por ejemplo un tejido blando es buen conductor; y un tejido óseo se considera [5] como un buen aislante

El principio de medida usado en la espectroscopia de impedancia eléctrica (EIE) es que la impedancia eléctrica de un medio con estas características es función de la frecuencia y viene usualmente descrita por (1) como sigue:

$$Z = Re(Z) + j Im (Z), \quad (1)$$

donde $Re(Z)=Z'$ es la parte real de la impedancia eléctrica Z y $Im(Z)=Z''$ la parte imaginaria de Z . El valor de $Im(Z)$ es normalmente negativo lo que indica efectos capacitivos, y dado que el hueso es un tejido biológico este no contiene la parte inductiva, y si representamos la parte imaginaria de la impedancia eléctrica de un sistema biológico respecto a la parte real de dicha impedancia a varias frecuencias (diagrama Cole-Cole), la figura resultante se aproxima a un semicírculo figura 1 según un modelo bien descrito en la literatura [2], donde el intervalo de frecuencia es por lo general de 40Hz a 100 MHz.

Este diagrama correspondería generalmente a un circuito eléctrico RC en paralelo como el de la figura 2, donde R_s es la resistencia de contacto de los electrodos y R_1, C_1 nos representa el compuesto de los cual esta formado el material, al que se le aplico la prueba de espectroscopia de impedancia eléctrica.

La Espectroscopia de Impedancia Eléctrica[3] puede separar las diferentes contribuciones a la respuesta de impedancia de un material[4], mediante la medición de dicha respuesta en un intervalo amplio de frecuencias, ya que los materiales tienen una constante de tiempo la cual indica que llega un momento en que el material deja de ser capaz de seguir al campo eléctrico, y ya no contribuye al valor de la impedancia; y es cuando se dice ocurre una dispersión o relajación, por lo que el diagrama de la figura 1 corresponde a un componente que tiene una sola relajación.

Una relajación se obtiene por medio de un circuito equivalente RC paralelo. La frecuencia de relajación corresponde a la condición de igualdad de R y C del circuito (2):

$$R = 1/\omega_x, \quad (2)$$

en donde R y C son los componentes del circuito equivalente y ω_x es la frecuencia angular donde $\omega_x = 2\pi f_x$, y en la impedancia como función de la frecuencia, la relajación aparece como un máximo, centrado en ω_x .

Así, se obtiene un semicírculo en el plano complejo o gráficas Cole-Cole la figura 3 corresponde al diagrama Cole-Cole de un material con dos relajaciones.

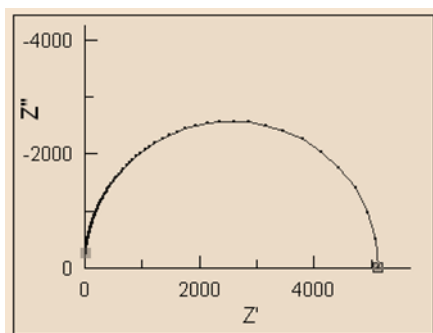


Fig. 1. Diagrama Cole-Cole.

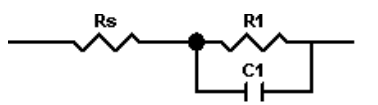


Fig. 2. Circuito eléctrico equivalente del semicírculo de la figura 1.

2. Desarrollo

Para el estudio de espectroscopia de muestras de hueso descalcificado, primeramente se cortaron 10 muestras de hueso de cerdo de aproximadamente 3mm de espesor con un diámetro de 2 cm, las cuales se lavaron y se quitó el máximo de cartílago y sangre, después se colocaron en una solución descalcificante la cual es como la empleada por los patólogos para descalcificar. Que consiste de una solución compuesta de HCl (ácido clorhídrico) concentrado al 0.001 % usándose la cantidad de 0.002 ml y 10 ml de $H_2C=O$ (formol), y se controla el tiempo que actúa la solución descalcificante.

Una vez descalcificadas las piezas, estas son lavadas con agua destilada para neutralizar el ataque y son observadas en el microscopio para ver como han sido descalcificadas, en la figura 4 se aprecia una muestra ósea sin descalcificar y otra con 5 días de descalcificación.

El siguiente paso es que las muestras son secadas en un horno a 25°C por 5 días hasta lograr un peso constante, el cual es medido para compararlo con el peso de la pieza sin descalcificar y así conocer la cantidad de pérdida de Ca^{+2} . Después de esto, las piezas son pulverizadas usando un mortero de ágata, en la tabla I se indican los

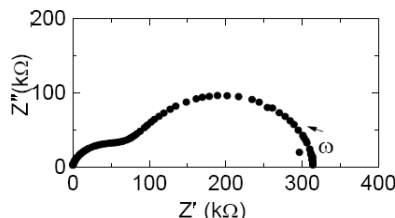


Fig. 3. Diagrama Cole-Cole de un material con dos relajaciones.

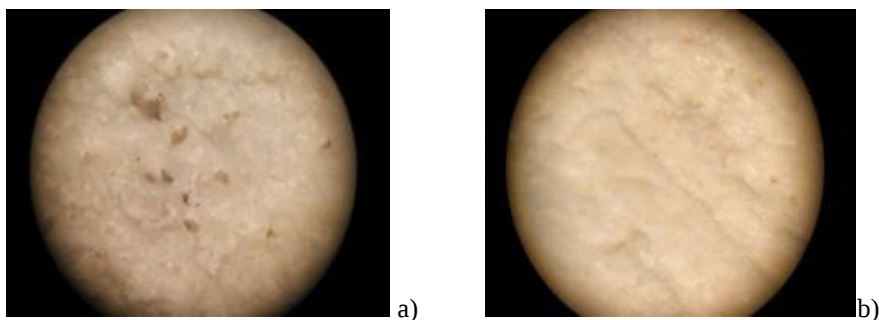


Fig. 4. Muestras óseas sin descalcificar (a) y con 5 días de descalcificación (b).

valores de calcio determinado por pérdida de peso después del proceso descalcificante.

Finalmente todas las muestras son molidas hasta formar un polvo fino; se preparan las 10 muestras con un peso de 0.280 gramos cada una; y esto es lo que se colocara en la celda de prueba para gráficos Impedancia RX y los de Cole-Cole de las pruebas de Espectroscopia.

Se anexo el medidor de fuerza porque la presión aplicada a la celda de prueba altera los resultados de la impedancia, para lo cual se realizó una prueba de Impedancia contra ángulo a distintas presiones en KgF, el resultado mostrado en la figura 6 nos indica si es posible una variación de la impedancia por el cambio de presión en la muestra ósea, realizar la prueba de espectroscopia de Impedancia Eléctrica.

a) Medición de Espectroscopia de Impedancia Eléctrica en muestras descalcificadas

Para la realización de las pruebas de espectroscopia de impedancia eléctrica se muestra en la figura 5 el arreglo experimental empleado para la realización de dicha prueba.

En el diagrama se tiene una celda de prueba, que consiste en dos electrodos de bronce dentro de un cilindro de material aislante, el cual tiene dos terminales de contacto para adaptarse al puente medidor de impedancia Agilent HP4294A.

La celda tiene acoplado un medidor de fuerza marca Mecmesin con un intervalo de 0 a 25 KgF, y los datos capturados por el HP4294A son procesados en una PC, a fin de obtener los resultados.

Tabla 1. Cantidad de calcio perdida en cada unade las 10 muestras descalcificadas.

Muestra	Calcio(gr)
1	0.103
2	0.164
3	0.244
4	0.324
5	0.384
6	0.473
7	0.603
8	0.694
9	0.771
10	0.864

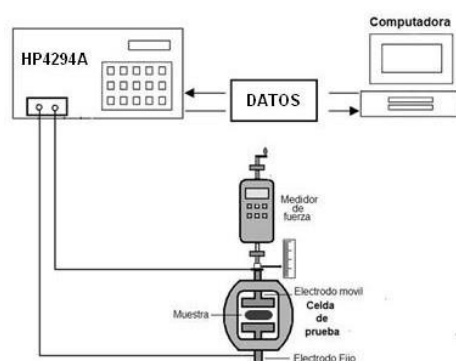


Fig. 5. Arreglo experimental para realizar las pruebas de espectroscopia de Impedancia Eléctrica.

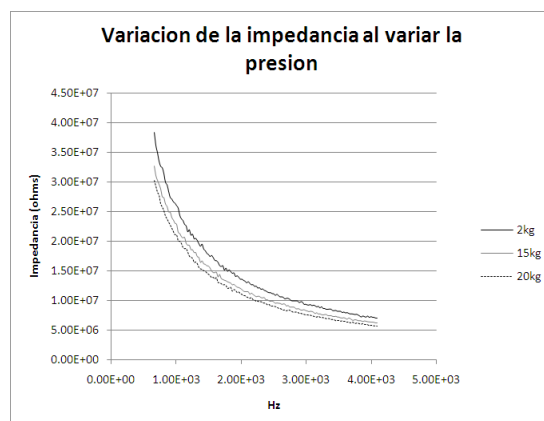


Fig. 6. Variación de la Impedancia con diferentes valores de presión en la celda de prueba.



Fig. 7. Arreglo experimental empleado para la medición de Espectroscopia.

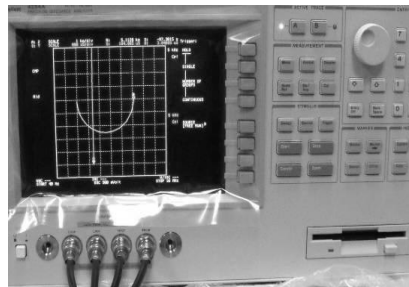


Fig. 8. Medición de circuito de prueba RC y sudiagrama Cole-Cole.

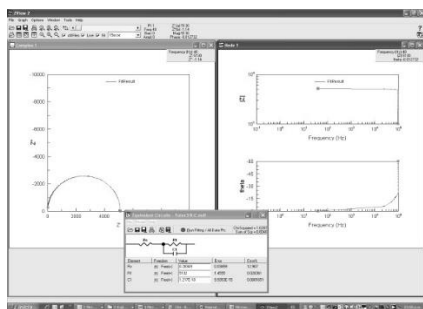


Fig. 9. Grafica Cole-Cole y de Impedancia-ángulo en el programa Z-view.

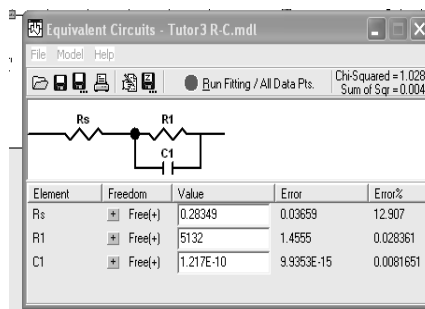


Fig. 10. Circuito equivalente simulado con el programa Z-view del circuito RC de prueba.

Y como apreciamos en la figura 6 si es posible registrar un cambio en la impedancia cuando aplicamos 2, 15 y 20 KgF, por lo que las pruebas a realizar serán a una presión de 15 KgF, para evitar alterar los datos de la espectroscopia, en la figura 7 se muestra la celda de prueba con el medidor de fuerza y el puente medidor de impedancia HP4294A.

3. Pruebas y Resultados

Para la realización de las pruebas de espectroscopia de Impedancia Eléctrica primeramente se requiere calibrar el medidor de Impedancia HP4294A, para lo cual el puente de impedancias cuenta con un protocolo de calibración utilizando una resistencia de precisión de 100%. Después de calibrar se midió un circuito RC en paralelo de $R = 5.1 \text{ k}$ y $C = 120 \text{ pF}$ con el puente de Impedancias, en un intervalo de 40hz a 5Mhz cuyo resultado se aprecia en la figura 8 y observamos como en el equipo se aprecia el diagrama Cole-Cole.

Para comprobar que la medición realizada con el puente sea correcta, los datos capturados en la prueba son procesados en el programa Z-view a fin de graficar el

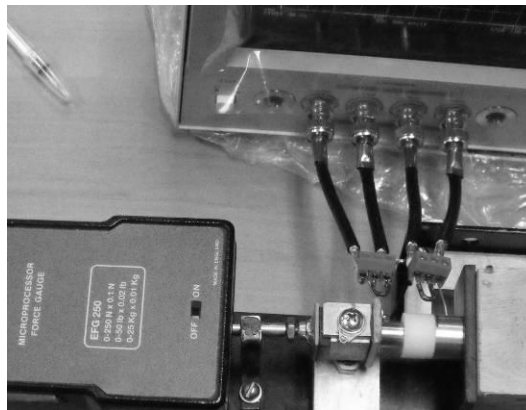


Fig. 11. Prueba de Espectroscopia de Impedancia Eléctrica a las 10 muestras.

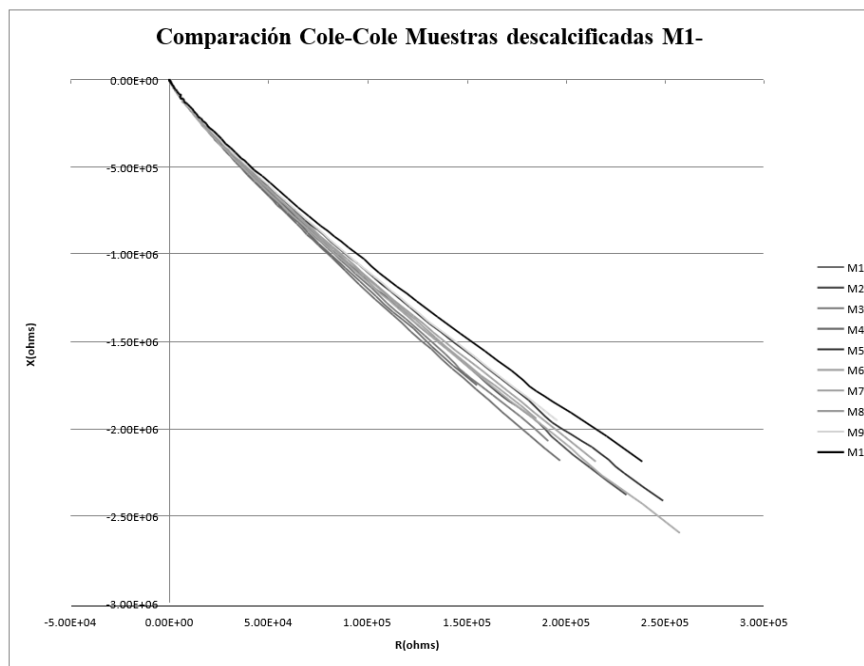


Fig. 12. Comparación de resultados de los diagramas Cole-Cole obtenidos de la prueba de espectroscopia de impedancia eléctrica en las 10 muestras descalcificadas de hueso.

diagrama Cole-Cole y obtener el circuito equivalente del circuito de prueba, los resultados se muestran en la figura 9.

Y de acuerdo a cálculo del circuito equivalente por medio del programa Z-view se obtienen los siguientes valores (fig10) del circuito RC en donde $R_s=0.28349$, que representa la resistencia de contacto en los electrodos, $R_1=5132\Omega$ y $C_1=1.217E-10$ o $C_1=121.7E-12$ que es equivalente a 121.7 pF.

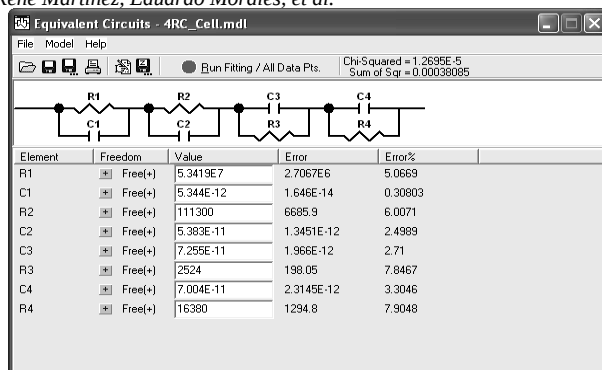


Fig. 13. Circuito Equivalente hueso descalcificado.

De acuerdo a los resultados tenemos que el error en R es 0.62% y en C 0.41%, acorde lo medido con el puente y calculado con el programa Z-view, con respecto a los valores RC de prueba utilizados por tanto consideramos que el puente Hp4294A funciona correctamente.

En base a la anterior prueba se procedió a realizar la prueba de Espectroscopia de impedancia Eléctrica en las 10 muestras de tejido descalcificado en la figura 11 se aprecia la celda de prueba conectada al puente Hp4294A para la realización de la prueba de Espectroscopia, dicha prueba se realizó para un intervalo de frecuencia de 50 hz a 110Mhz, y con la opción del puente para incrementar la frecuencia de forma logarítmica, los resultados de las mediciones se observan en la figura 12 (mostrada en la última página), en donde se comparan los diagramas Cole-Cole [6] de las 10 muestras de hueso de 1 a 10 días de descalcificación, esto con el fin de contar con una muestra ósea descalcificada controlada que sería la forma de simular osteoporosis. De los datos del diagrama Cole-Cole obtenido de una de las muestras, se simula en el programa Z-view y se obtiene el siguiente circuito equivalente (fig13), que relaciona el comportamiento de la muestra ósea descalcificada.

4. Conclusiones

Conforme a los resultados obtenidos y mostrados en la figura 12 se puede apreciar un cambio en el diagrama Cole-Cole, la muestra M1 corresponde a un día de descalcificación y la M10 a 10 días, lo que indica que hay una variación en la impedancia de la muestra ósea descalcificada, por tanto se puede relacionar un cambio de impedancia con la pérdida de calcio, que sería una forma de simulación de osteoporosis en hueso.

Por tanto, se cumple plenamente el objetivo del proyecto que era comprobar si es posible medir descalcificación ósea por medio de espectroscopia de Impedancia Eléctrica, lo que estos resultados dan la pauta para considerar sea esta una nueva técnica de medición para descalcificación ósea.

A su vez en un trabajo futuro sería obtener el comportamiento de los 10 circuitos equivalentes y relacionarlos con las cantidades de calcio obtenidos en la prueba de descalcificación y observar que parámetros se mantienen casi sin variación, y

determinar en el circuito equivalente obtenido a que corresponden los cuatro circuitos RC obtenidos.

Agradecimientos. El estudio realizado formo parte del proyecto “Estudio Espectroscópico de Tejidos Biológicos” financiado por Concyteg Guanajuato proyecto NoCI30380407.

Referencias

1. Macdonald, J. R.: Impedance Spectroscopy. *Annals of Biomedical Engineering*, 20, pp. 289–305 (1992)
2. Jonscher, A.K.: Dielectric Relaxation in Solids. Chelsea Dielectrics Press, London (1983)
3. Duck, F.A.: Physical Properties of Tissues. A Comprehensive Reference Book (1990)
4. Duck, F.A.: A comprehensive Reference Book. Ed. Academic Press: San Diego (1990)
5. Bourne, J.R., Rigaud, B., Morucci, J.P., Chauveau, N.: Bioelectrical Impedance Techniques in Medicine, *Physiol. Meas.*, 24, pp. 545–554 (2003)
6. Zehe, A., Ramírez, A.: Efectos electrocinéticos de células biológicas y partículas coloidales en la espectroscopia dieléctrica a bajas frecuencias *Revista Mexicana de Ingeniería Biomedica*, XXV(1), pp. 16–24 (2004)
7. Bragos, R.: Biomass monitoring using impedance spectroscopy, *Annals of the New York Academy of Sciences*, 873, pp. 299–305 (1999)

Electronic edition
Available online: <http://www.rcs.cic.ipn.mx>



ISSN: 1870-4069
<http://rscs.cic.ipn.mx>



Centro de Investigación
en Computación