



IJITCE

ISSN 2347- 3657

International Journal of Information Technology & Computer Engineering

www.ijitce.com



Email : ijitce.editor@gmail.com or editor@ijitce.com

Review of Medical Image Retrieval Systems and Future Directions

¹CHEEKATYALA ANJAN KUMAR

²ACHINA VENKATARAMANA

³JOGU PRAVEEN,

⁴BURRI NARENDER REDDY

Abstract:

The goal of this study is to present an overview of online systems for content-based medical image retrieval, with a focus on the United States (CBIR). The authors of this study hope to identify the advantages and disadvantages of these systems, as well as approaches to improving the relevance of multi-modal (text and picture) information retrieval in the I Medline system, which is currently under development at the National Library of Medicine, by the end of the study (NLM). A total of seven medical information retrieval systems were investigated in this study, including Figuresearch, BioText, GoldMiner, Yale Image Finder, Yottalook, Image Retrieval for Medical Applications (IRMA), and I Medline. Figuresearch was the most popular system among participants. The systems were assessed in accordance with the system of gaps described in [1]. However, not all of these systems make advantage of the visual information supplied in biological literature in the form of figures and drawings, but a significant number of them do. All, on the other hand, make an attempt to extract image information from the full-text of the articles and to acquire figures and photos in response to a search query, which is a common practise. It is the purpose of iMedline to advance the state-of-the-art in multimodal information retrieval by merging image and text data in the calculation of relevance, a goal that has so far been accomplished. In this work, we discuss the shortcomings of current medical image retrieval systems, as well as future directions and next phases in the development of iMedline's context-based medical image retrieval system.

1. Introduction

In addition to online literature databases like PubMedCentral® [2] and BioMedCentral® [3], there is a wealth of biomedical information available, including case studies from patient records kept in electronic health records (EHRs). The recovery of this information may be valuable to physicians, patients, and those who teach or study medical sciences since it may aid in better diagnosis, treatment planning, classroom learning, and medical research. While traditional bibliographic or full-text databases provide a substantial quantity of textual information, online biomedical literature includes a significant amount of visual information in the form of figures and drawings that is not generally available

through these traditional resources. Captions and full-text excerpts, while helpful in explaining the meaning of figures and photographs, fall short of correctly representing the semantic information included in medical imagery, which is best understood visually by human specialists. By going beyond traditional text-based searching and adding both text and visual aspects in search queries, we expect to identify more effective techniques of extracting information from a variety of different sources than are now available. As a starting point, we examine and evaluate seven medical information retrieval systems: FigureSearch, BioText, GoldMiner, Yale

,M.Tech Assistant Professor , anjankumarcheekatla@gmail.com

, M.Tech Assistant Professor, venkatachina5@gmail.com

M.Tech Assistant Professor, jpraveen.pec@gmail.com

M.Tech Assistant Professor, narendarburri@gmail.com

Department: ECE

Pallavi Engineering College Hyderabad, Telangana 501505

Image Finder, York University's Yottalook (York), Image Retrieval for Medical Applications (IRMA), and the National Library of Medicine's iMedline. FigureSearch is a medical information retrieval system developed by the National Library of Medicine (National Library of Medicine). In order to determine how to increase the relevance of multi-modal (text and picture) information retrieval in iMedline by incorporating the lessons learned from these initiatives, we will examine the inadequacies and capabilities of these systems once they have been completed and evaluated. The goal is to make it easier to identify and obtain biomedical literature by focusing on its visual content. We also want to make it easier to retrieve photographs that are semantically related to each other. This will aid in differential diagnosis, clinical decision support, research, and educational endeavours, among other things. The following diagram depicts the overall structure of this work. The first half of this paper presents an overview of the many methods of retrieving medical information from the internet that are currently available. Following that, an evaluation of these systems is conducted in accordance with the semantic gaps identified in [1]: content, feature, usability, and performance gaps, among other factors. In the following section, we discuss potential research directions that could be pursued in order to close these gaps and carry out context-based medical image retrieval in the medical imaging domain, respectively.

2. Medical Information Retrieval Systems

2.1 GoldMiner

Goldminer® [4] allows you to search figure captions to retrieve images from over 11000 free, open-access peer-reviewed journal articles from websites of organisations such as the American Roentgen Ray Society (ARRS), the American Society of Neuroradiology (ASN), the British Institute of Radiology (BIR), and the Radiological Society of North America (RSNA) (RSNA). Figure caption terms are linked to concepts in the Unified Medical Language System (UMLS®) metathesaurus and/or Medical Subject Headings (MeSH®) terms in the National Library of Medicine's Unified

Medical Language System (UMLS®) metathesaurus using this tool. The following examples of list and grid displays of search results for various search criteria are shown in FIGURE 1. Users can search for images based on their age, modality, and gender, all of which are determined by the caption text in the image. In addition, it enables you to search for a large number of terms at the same time.

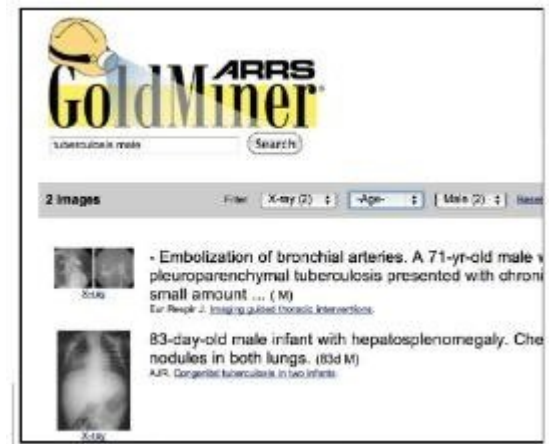


Figure 1. The ARRS GoldMiner search engine

2.2 FigureSearch

As illustrated in Figure 2, the FigureSearch search engine created at the University of Wisconsin at Milwaukee is a component of the askHermes system [5], which is a tool designed to improve the quality of patient care by giving information to physicians at the point of care. It searches for medical articles on the internet using the Lucene® text indexing and search technology and displays the results in a list format. The images are presented on the left, while the title, authors, figure caption, and summary are displayed on the right side of the page, as shown below. The ability of the search engine to automatically construct summaries from articles (including the purpose, experimental approach, outcome, and conclusion) utilising sentences from the main text distinguishes it from the competition.



Figure 2. The FigureSearch search engine.

BioText (2.3) The BioText search engine [6], created at the University of California at Berkeley and depicted in Figure 3, makes use of the Lucene search engine to index over 300 open access journals and extract figures and text from online articles. It is also based on the Lucene search engine. Users have the option of conducting searches on either the full-text or abstracts of journal articles. It differs from other search engines in that it can search table captions and extract part/expanded views of tables from web articles, something that other search engines cannot do. Results can also be filtered by relevance and date, among other criteria.



Figure 3. BioText search engine.

2.4 Yottalook

Yottalook [7] is a multilingual search engine that retrieves images from peer-reviewed journal papers published on the Internet. It supports thirty-three languages and does multilingual searches. It makes

use of Google's indexing technology as well as a proprietary software package known as iVirtuoso for natural query analysis, semantic ontology building, and relevance determination, among other things. Using natural query analysis, you can produce keywords from search query data. It employs an upgraded version of the RSNA's RadLex® medical ontology to find relationships between terms or concepts that are synonymous with one another. Semantic ontology generation is the term used to describe this process. The relevance of a search result is determined automatically by a relevance algorithm (which is included in the iVirtuoso software) and is used to rank the results that are returned. Result views include grid and list views, as shown in Figure 4. Users can save their searches using their myRSNA account credentials, which is also displayed in Figure 4.

3.1 Deficit in performance By assessing whether a search engine can be utilised to acquire information from a narrow or vast range of datasets, the performance gap may be utilised to evaluate the scope of an application. In addition to accessing material from online journal databases, all systems discussed here, with the exception of the IRMA system, have a comprehensive range of operating capabilities. Although there are techniques for undertaking a thorough quantitative evaluation of the retrieval performance of these systems on online biomedical journals, there is currently no technique for doing so.

3.2 Deficiency in a feature The sort of text and/or picture features employed by a search engine is listed in the Feature Gap section. GoldMiner, Figuresearch, YIF, Yottalook, and BioText are among the programmes that use text-only features at the present time. IRMA and iMedline, on the other hand, make advantage of global image qualities such as colour and texture to communicate the visual content of a picture. These systems, on the other hand, have not addressed concerns such as deriving multi-scale features, feature extraction from local regions of interest, region labelling, and other issues that need to be addressed.

3.3 Unsatisfactory Usability The usability gap is used to compare systems against different forms of enquiries, and feedback is applied to improve the relevance of search results by making them more relevant. Text-only information retrieval systems already in use execute retrieval by using keywords, phrases, and/or multiple keywords to find relevant data. Text and image queries are utilised by systems

that make advantage of picture features. Flexible query refining, as well as enhanced relevance feedback, are two areas where there are still holes to be filled in this sector. In the future, query refinement techniques such as union, intersection, and negation of searches, as well as hybrid inquiries that mix text and picture information, will be researched in order to boost the relevancy of search results and to improve user experience. Furthermore, these systems must take into account user feedback and participation, which can be accomplished through relevance feedback.

3.4 Inconsistency in the information presented It is the phrase used to refer to the conceptual disparities that exist between textual and visual concepts that are used in information retrieval. A clear demonstration is provided by the table, which clearly shows that systems that rely solely on textual information, such as captions and full-text extracts, may only represent concepts that explain the content of an image to the extent that they are synonymous with MeSH words or UMLS concepts, as demonstrated in the example. It is impossible for them to accurately depict the visual information contained in images, such as anatomical and pathological information, disease severity, and so on, unless this information is provided in the text that is accompanying the photograph. Using medical images to extract visual information, and then mapping that information to high-level textual medical concepts, it is possible to overcome this constraint. This ability may be demonstrated by the IRMA system on chest radiographs and mammograms, both of which are pictures with a small picture domain. We want to increase the functionality of such a system by combining text and picture data, and to enable context-based image retrieval in the process.

This includes an investigation of the following topics of interest: Pre-filtering photos based on modality, body part, and orientation is used to automatically categorise images and decrease the search space in order to limit the search space. The process of identifying regions of interest in medical images by utilising annotation markers inside figures, such as arrows, letters, or symbols, that have been retrieved from the image and correlating them with concepts in the accompanying text (2) Developing effective methods for measuring regions of interest/image patches is also critical in order to allow them to be indexed and compared in order to do similarity retrieval is another essential goal. Detailed descriptions of our current and future research activities in these areas, as well as our findings, will

be provided in the following section. Fourteenth, work in the future 4.1 Image Categorization Using Machine Learning is the fourth chapter. A pre-filtering process known as automatic picture categorization can be used to reduce the search space, allowing for faster and more effective similarity matching on large image collections. It is particularly beneficial in the field of radiology. Before it is possible to establish whether two medical images are visually comparable to one another, the images in the database must first be classed according to modality, body part, and orientation before being compared to one another in the database. By extracting image features from the query image and comparing them against an index of image features, it is possible to construct a ranked list of images based on their likeness to the query image. 4.2 Automatic image annotation and ROI extraction methods are used in this section. When working with major regions of interest or critical spots within an image, it is necessary to derive features over these areas in order to extract meaningful information from them in order to extract relevant information. This problem is solved by first separating subfigures from composite figures and then searching for useful "pointers" or annotations (arrows, symbols, or text labels) that point to the ROI [15], which we refer to as the ROI pointer. As part of this research, we are looking into techniques of automatically identifying and recognising images' annotations (arrows, text labels) as a means of linking image ROIs with concepts obtained from image captions [16], which is now under investigation.

4.3 Visual Keywords

Using visual features to quantify images, we have developed "visual keywords" [17], which are local image attributes that are used to generate a bag of concepts that is similar to the bag of words representation often used in information retrieval from text materials. Image "patches" are created by uniformly splitting a single image into non-overlapping parts, and visual keywords are used to model the colour and texture attributes produced from these patches. In order to automatically classify photographs into several modality categories, this approach has been employed in the past. Visual keywords can also be used to improve the relevance of visually comparable photos when utilising text-based image retrieval (IR) algorithms to find them in the first place. We are currently investigating unsupervised image segmentation algorithms to define gross image regions in order to enhance the current framework's ability to recognise the gross

anatomy of pictures [18]. Briefly said, medical graphics play an extremely significant part in the process of clinical decision-making. Because of this, approaches for effectively mining medical images for information utilising textual descriptions, image features, and a mix of text and image features must be investigated. The text associated with photographs, such as captions and full-text snippets, is indexed by the majority of currently accessible search engines (Table 1), which allows them to run searches in response to a user query. We are now working on conducting semantic retrieval of biomedical pictures utilising a combination of text and image features applied to generating regions of interest, representing regions or image patches as visual keywords, and increasing relevant feedback in order to improve performance.

References

- [1] T. M. Deserno, S. Antani, R. Long, "Ontology of Gaps In Content-Based Image Retrieval", *Journal of Digital Imaging*, vol. 22, no. 2, pp. 202-15, April, 2009.
- [2] PubMedCentral: <http://www.ncbi.nlm.nih.gov/pmc/>
- [3] BioMedCentral: <http://www.bmc.org>
- [4] ARRS GoldMiner: <http://goldminer.rrs.org/>
- [5] Y. Hong, C. Yong-gang, "Automatically Extracting Information Needs from Ad Hoc Clinical Questions", *Proceedings of AMIA Symposium*, pp. 96-100, 2008.
- [6] M. A. Hearst, A. Divoli, H. Guturu, A. Ksikes, P. Nakov, M. A. Wooldridge, J. Ye, "BioText Search Engine: beyond 2197, 2007.
- [7] Yottalook: http://www.yottalook.com/index_img.php
- [8] S. Xu, J. McCusker, M. Krauthammer, "Yale Image Finder (YIF): a new search engine for retrieving biomedical 2008.
- [9] IRMA: http://ganymed.imib.rwth-aachen.de/irma/veroeffentlichungen_en.php
- [10] iMedline: <http://archive.nlm.nih.gov/iti/>
- [11] MedGIFT: <http://medgift.hevs.ch/silverstripe/>
- [12] J. Lim, J. Chevallet, "VisMed: A Visual Vocabulary Approach for Medical Image Indexing and Retrieval", in *proceedings of Asia Information Retrieval Symposium*, pp. 84-96, 2005.
- [13] ASSERT: <http://cobweb.ecn.purdue.edu/RVL/Research/CBIR/index>
- [14] BRISC: [BRISC: http://brisc.sourceforge.net/](http://brisc.sourceforge.net/) [15] B. Cheng, J.R. Stanley, S. Antani, G.R. Thoma, "A Novel Computational Intelligence-based Approach for Medical Image Artifacts Detection", *Proceedings of the 2010 International Conference on Artificial Intelligence and Pattern Recognition*, pp.113-20, 2010.

[16] D. You, S. Antani, D. Demner-Fushman, M.M. Rahman, V. Govindaraju, G. R. Thoma, "Biomedical Article Retrieval Using Multimodal Features and Image Annotations In Region-based CBIR," *Proceedings of the Document Recognition and Retrieval Conference*. pp. 1-10, 2010.

[17] M.M. Rahman, S. Antani, G. R. Thoma, "A Classification-Driven Similarity Matching Framework For Conference on Multimedia Information Retrieval, pp.147-154, 2010.

[18] P. Ghosh, S. K. Antani, L. R. Long, G. R. Thoma, "Automata for Region-based Image Retrieval", accepted for the IEEE Conference on Healthcare Informatics, Imaging and Systems Biology, 2011.