

Image Retrieval in Medical Applications: The IRMA-Approach*

Thomas Lehmann¹, Berthold Wein², Daniel Keysers³, Jörg Bredno¹, Mark Oliver Güld¹, Christian Thies¹, Henning Schubert², and Michael Kohnen²

¹ Institute of Medical Informatics, Medical Faculty

² Department of Diagnostic Radiology, Medical Faculty

³ Chair of Computer Science VI, Computer Science Department
Aachen University of Technology (RWTH), D-52057 Aachen, Germany
<mailto:lehmann@computer.org>, <http://www.irma-project.org>

Abstract. The growing need of organizing and administrating large archives of images has lead to substantial progress in the field of image retrieval in the past few years. Among all others, the approaches concerning information extracted immediately from a picture have shown encouraging results. Entire systems for content-based image retrieval (CBIR) have been implemented and tested on different areas of application. However, those systems still fail when applied to medical image databases. This problem is caused by the complex descriptive semantics of medical knowledge that is extracted from images by human experts. Common CBIR mainly concerns global information such as histograms, texture or conspicuous shapes. These simple features are insufficient to describe multi-layered information in medical images such as radiographs.

In this paper, we present a novel multi-step approach, which is specially designed for content-based image retrieval in medical applications (IRMA). From a theoretical and practical point of view, the aspects of knowledge representation and system implementation are taken into account, respectively. IRMA is based on a conceptual and algorithmic separation of: (a) image categorization, (b) registration, (c) local feature extraction, (d) local feature selection, (e) index generation by hierarchical blob representations, (f) object identification of blobs and (g) image querying on the abstract blob-level. This concept is implemented as a distributed development and runtime system, where the methods computing the steps (a) to (g) as well as the processed image data can be accessed from multiple users at different physical sites. This enables co-located groups of developers to contribute to the system's functional range at the same time, as the system is in production state for medical exertion.

Keywords: Content-Based Image Retrieval (CBIR), Content-Based Indexing, Medical Image Databases, Image Content, Registration, Classification, Categorization, Picture Archiving and Communication Systems (PACS).

* This work is supported by the Deutsche Forschungsgemeinschaft (DFG), grant Le 1108/4.

1 Introduction

The importance of retrieval techniques increases in the emerging fields of medical imaging and picture archiving and communication systems (PACS). Information contained in large databases can be used to improve health care by supporting diagnosis tasks. The retrieval of assured diagnoses or a patient's case history offers not obvious information to the physician. Up to now, textual index entries are mandatory to retrieve medical images from hospital archives or other sources [1]. This also holds for DICOM archives [2]. In clinical routine, the acquisition of textual information has high data entry costs, i.e., it is cumbersome and time consuming. The medical expert enters his report in a standardized form which is appended to the image. However, this input cannot be adjusted with the actual content of the image. As a result, information contained in medical images differs considerably from that residing in alphanumeric format [3]. This leads to incomplete recall, or even worse, to low precision, if the expert enters only a short description or something just similar to the actual diagnosis, respectively. Hence, a textual image database cannot be considered as a reliable source of information. To close the gap between abstract alphanumeric description and actual intrinsic content of an image, a content-based information extraction approach has been claimed.

Due to the complex and unsharp character of medical image content, this information has to be extracted on different levels of abstraction. A coarse categorization of the image due to its acquisition domain CT, x-ray, etc. has to be followed by recognition of contained objects such as thorax, mamma, etc. and their main orientation. Present systems for content-based image retrieval (CBIR) in medicine work with a semantic restriction to dedicated fields of application. For example, ASSERT operates on high-resolution CTs of the lung [4] while GEMINI is focussed on tumors in mammograms [5]. The architecture of these systems is optimized to their special field of application and cannot evolve with the requirements and improvements in diagnostic technologies and procedures.

Consequently, an open concept of medical CBIR must support the radiologist by transferring his data and knowledge into the database and the developer by introducing new algorithmic methods of content description. The latter is mainly necessary to offer new descriptions of image content as the physician's task at hand may change. Hence, a flexible software architecture is required providing the user with an easy to use interface and the programmer with an application programming interface that enables him to add methods of image analysis and knowledge representation in a simple way. Both rely on a database holding raw data, processed images, content information and transformation methods.

In this paper, we describe the knowledge representation framework of content-based image retrieval in medical applications (IRMA) and its implementation in a distributed development and runtime environment. Based on general requirements for CBIR (Sec. 2), some of the unique challenges in medical imaging and the resulting constraints to medical image retrieval are stressed (Section 3). In Sections 4 and 5, we describe the resulting multi-step approach of IRMA and main aspects of software architecture, respectively. With respect to system ar-

chitecture, incorporation of medical knowledge, and modeling of semantic levels, advantages and limits of the IRMA concept are discussed in the last section of this paper.

2 General Requirements for Image Retrieval

Extraction of non-textual information from images leads to the complex problem of automated description of image content, i.e. image understanding. With regard to the retrieval task, this has to be done in a normalized way to provide comparability of the machine representations of images. Thus, finding a normalized computable description of image content is one of the key-tasks when building a CBIR system. In the early 1990s, general concepts and implementations of CBIR systems have been developed [6], where two basic principles became important.

Feature Extraction. For a set of images, numerical features are extracted automatically from

- color or grayscale histograms
- texture
- shape

and used to describe the entire image. This provides a dramatic reduction of data, which is required to perform online retrieval. There is always a trade-off between computational cost and discriminative precision of image describing features. Common CBIR systems [7] have low data entry costs, i.e., they require minor user interaction when archiving an image and consequently, only a rudimentary understanding of image content, i.e., they compute a simple set of features concerning mainly global image information. Such systems do not distinguish important and unimportant features within the image. The features used for automated indexing characterize the entire image rather than local information such as regions or objects. In contrast, queries of medical or diagnostic relevance include searching for organs, their relative locations, and other distinct features such as morphological appearances. Therefore, common CBIR systems cannot guarantee a meaningful query completion when used within the medical context [8].

Feature Comparison The numerical features describing an image are combined to a normalized feature vector. Finding images similar to another image or an user sketched prototype becomes a comparison of the new images' feature vector with those stored in the database (Fig. 1). Based on the query by example (QBE) paradigm, a CBIR system equals a classification task, where simple classifiers such as nearest neighbor serve as a measure of equality. The semantic model, which is used for stepwise description of knowledge consists of only two layers, the raw data layer and the feature layer. Further knowledge on image

content is not considered during QBE processing. Later generations of CBIR systems add a third layer to the semantic model. For the content-based retrieval engine (CORE), objects and spatial relationships are described by “concepts” within the so called interpretation layer [9]. In Blobworld, this layer is build from regions of uniform color or texture on an abstract level of interpretation [10].

3 Characteristics of Medical Image Retrieval

The problem of image retrieval in medical applications has been systematically analyzed by TAGARE et al. [3]. They have postulated four basic properties:

1. Information contained in medical images such as local textures or shape cannot be expressed by words in a precise and, above all, unique way. Thus, based on automatic generated and reduced data, content-based access is required for medical image retrieval.
2. Medical image information is complexly linked and nested. Local features depend on mutual orientation in geometry and time of the represented image objects. This correlation and order has to be conserved in the feature vector representation. Thus, a local approach is necessary.
3. Semantic interpretation of medical images is context-dependent with respect to the diagnostic query. However, the context is unknown at the moment of image acquisition or data entry. Furthermore, the context may change during the course of examination. Thus, the kind and number of stored features of an image must stay flexible.
4. Medical image interpretation is a complex and poorly understood process. Diagnostic inferences derived from images rest on an incomplete, continuously evolving model of normality. Thus, the usage of secure knowledge about the images, i.e. a-priori knowledge, is inevitable.

Following these properties, it becomes obvious why the results are rather poor when common CBIR systems are used to retrieve medical images [8, 11]. They perform only a global color, texture, or shape analysis and represent their knowledge on two or three levels of abstraction. However, the modeling of medical

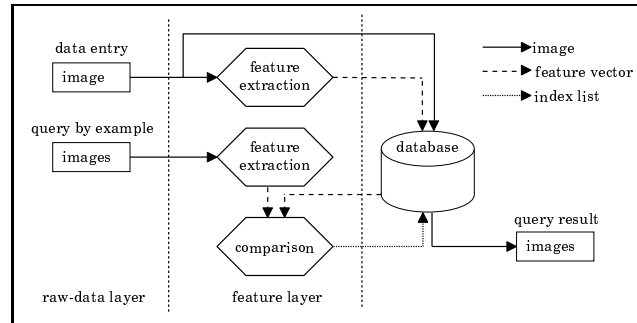


Fig. 1. Flow of operations and semantic layers of a simple CBIR system.

knowledge for image retrieval requires a more differentiated view on information.

4 Knowledge Representation in IRMA

Image retrieval in medical applications has to fulfill these statements without restrictions on either image category or query content. Evolving from the aspects of medical knowledge representation and technical needs, we present the IRMA approach for medical CBIR [12], which is based on the conceptual and algorithmic separation of the following steps (Fig. 2):

1. *Categorization* using global image features.
2. *Registration* in geometry and contrast for each likely category.
3. *Feature extraction* using local image descriptions.
4. *Feature selection* and combination with respect to category and query.
5. *Indexing* resulting in a hierarchical multi-scale blob representation.
6. *Identification* of blobs by linking a-priori knowledge to image content.
7. *Retrieval* processed on the abstract blob-level.

4.1 Categorization

For high-level image analysis, basic image characteristics must be already known. For example, different filters are needed to extract local information from ultrasound images than from skeletal radiographs. Within the IRMA system, basic

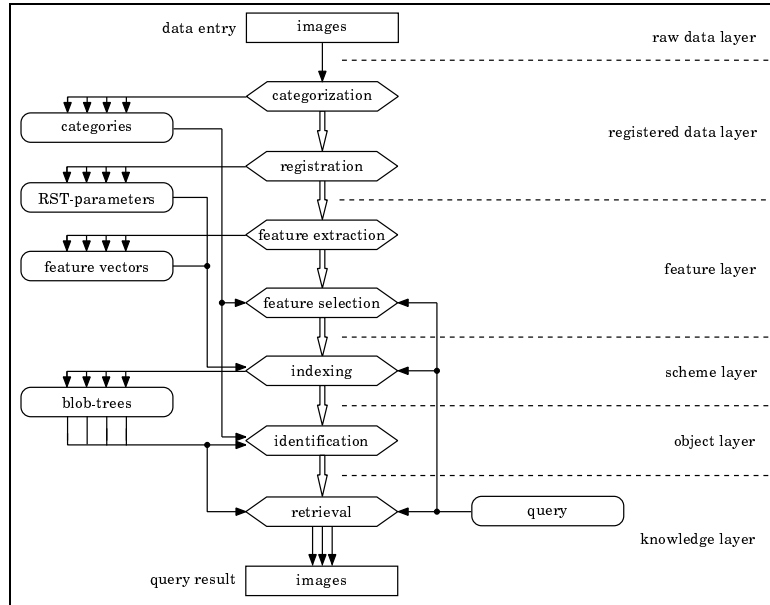


Fig. 2. Processing steps and semantic layers of IRMA. Multiple arrows indicate one-to-many relations.

image categories are determined in the first step of processing. It is important to note that this categorization is based only on global features characterizing the entire image such as histograms, cooccurrence matrices, median of gray-scales or colors, and texture parameters. Also, the IRMA categories were related to the classification scheme introduced by the USINT working group of the EurIPACS AIM project [13]. Hence, categories are defined a-priori from medical knowledge but not from data-dependent cluster analysis of the feature space. Each image is classified with respect to

- image modality and technique,
- body region examined,
- orientation of imaging, and
- functional system imaged.

The categorization of the images according to feature vectors must not necessarily be unique. The IRMA approach allows each image to be linked to several categories, where the likelihood is also stored within the IRMA database. Strictly speaking, this leads to a fuzzy approach when describing the membership of one feature vector to a dedicated image category.

The number of entries in the global feature vector is extensible. Whenever a new global feature extraction algorithm is incorporated to the IRMA system, the corresponding vector components are calculated for all images of the IRMA database. This is done automatically by parallel batch processing on all IRMA workstations [14], see also Sec. 5.2.

4.2 Registration

Medical a-priori knowledge must be incorporated into the IRMA system in order to derive diagnostic conclusions from images. Diagnostic conclusions are deduced from an incomplete but continuously evolving model of normality. Therefore, the IRMA system provides flexible prototypes for each category, which are defined by a medical expert. For each likely category, parameters are determined describing geometric registration, particularly, rotation, scaling, and translation (RST), and contrast adjustment. These parameters enable further processing for object identification and especially orientation-independent retrieval. They are stored together with the probabilities for the most likely categories given by the current image. To keep the costs of computation and storage low, the parameters are stored without actually transforming the image. Also note, that the RST parameters still describe global properties of an image.

4.3 Feature Extraction

Up to now, the images are regarded from a global point of view. However, a large part of information in medical images is geometric [3]. Hence, the processing of semantic queries drawn from medical routine requires local features. For instance, characteristic objects such as organs require local analysis. Therefore,

local features are now connected to an image on a pixel level resulting in feature images rather than feature vectors. The number of local feature images is not limited but expandable according to diagnostic requirements. If new methods are introduced the required subset of feature images is computed in a parallel batch process as mentioned above.

The feature images strongly extend the data volume although they are calculated only once for each input image but are suitable for all likely categories (category-free features). With respect to the restriction of special category-dependent feature extraction methods, also category-specific features are used. These category-specific features include segmentation using active contours or active shapes. The selection of those methods is possible according to the use of a-priori shape knowledge from the categories.

4.4 Feature Selection

The IRMA concept allows the de-coupling of feature extraction and selection. With this abstraction, the feature extraction can be made retrieval-dependent. A set of proper feature vectors for a query can be determined by information gained from the category as well as medical knowledge of the context. For example, the retrieval of radiographs with respect to bone fractures is done using edge-based feature images (e.g. form set). If the query concerns bone tumors, the same radiograph is described by another set of feature images (e.g. texture set). However, queries which operate on a certain subset of feature images require processing of the subsequent indexing procedure for all images in the database. This consumes too much time with respect to online queries. Hence, the IRMA concept comes with a few pre-defined sets combining distinct feature images for each category. Finding these feature sets will be done using statistical feature reduction methods such as Fisher's linear discriminant analysis, principle component analysis (PCA) or Karhunen-Loève-transform (KLT).

4.5 Indexing

A drastic blow-up of information results from the previous steps. This has to be reduced again for query processing. Therefore, some characteristic information is computed from the local feature sets. In particular, the multi-spectral feature images are segmented into relevant regions on hierarchical levels of resolution. Those regions are described by invariant moments resulting in structures we call blobs [10]. For example, these blobs equal ellipses if only first and second moments are used for description. Each blob is characterized by a N -dimensional feature vector formed by the average of the N feature images within the region. For each feature set, a blob structure is generated and stored within the IRMA database. In addition, each blob representation is registered with respect to the parameters for each likely category. The blob concept is supplemented by a multi-scale approach to support different levels of granularity when concerning regions of interest.

4.6 Identification

In the previous step, blob structures have been registered with respect to the prototype in each category. Therefore, the blob structure can serve as a general description of image content where certain blobs may correspond to well defined morphological structures in the image, such as organs, bone, or others. This correspondence is obtained from the discriminant properties of the local features. Also, each category's prototype is represented by a blob structure, which models the high-level medical a-priori knowledge on categorized images. The entities of the prototype blob structure contain properties of the corresponding image region, such as location, size, shape, and texture, as well as labels from a medical nomenclature. Thus, the feature vector values serve for blob identification and the added labels can be used for semantic queries. By means of a sophisticated interface, this labeled prototype blob trees are a powerful data structure enabling a medical expert to introduce high-level knowledge into the IRMA system. Knowledge on physiological occurrence of morphological structures within the human body is of significant value to allow the processing of medico-diagnostic queries. Hence, blob identification provides content understanding at a morphological level. However, an identification is not forced for all blob entities. Unlabeled blobs are also handled by the system.

4.7 Retrieval

The actual retrieval task is performed by a comparative search in the hierarchical blob structures. According to the image representation in IRMA a query is build from a list of possible categories of the recall images, a QBE blob structure on the optimal scale to process the query and the set of local features that best describe the significant properties for the query (e.g., texture set, form set). Thus, IRMA contains and extends the techniques implemented in all known CBIR systems. There are two general applications for automated image retrieval:

1. *Primitive queries* cover follow-up studies within a PACS. For instance, a physician wants to compare all images of a patient's abdominal region. Such queries can be performed on the global features without regarding the blob structure.
2. *Semantic queries* are based on a throughout diagnostic procedure. For example, a physician searches for representative images of known diseases during reporting. Such queries are processed by means of a user-described blob structure either by identifying blobs and, therefore, assigning properties that are related to morphological structures, or, by providing a pre-processed example image with a choice of significant blobs for the query.

In other words, the search is performed by comparison of blob structures with respect to the blobs of interest and selected features. It can be accelerated because only the blob structures for images of the assigned categories have to be compared. Data structures and distance measures that have already been described for image retrieval [15] need refining to process only significant blobs.

The multi-scale approach permits object-oriented processing, where dedicated regions of interest, such as labeled organs, can be addressed on different levels of detail.

Browsing the IRMA database requires mechanisms for query refinement and relevance feedback. Also, methods must be provided to extend this choice with qualitative descriptors like “larger”, “more intense”, or “between these two examples”, and combinations of these. The formulation of such queries requires intense work on the design of the user interface. However, the current state of development of the IRMA system is not focused here.

5 The IRMA Development and Runtime Environment

The core of the IRMA system consists of a collection of database driven applications supporting the processing steps described in Section 4. Resulting from its interdisciplinary nature, distributed development and querying is supported within the IRMA approach. Hence, efficient organization of algorithms, data structures, processed data and their persistency and consistency is a key challenge of the IRMA system. All IRMA development sites access a central database, which holds administrative information for all distributable resources such as images, methods, feature values, query views, and jobs. Each resource is identified by a unique system-wide identification (ID).

5.1 Design Goals

In IRMA, processing of images is split into a chain of *methods* applied to a feature vector representing the output from the previous method. Methods are executables embedding a user-programmed function (implementing the core feature processing functionality) into a pre-compiled framework provided by the IRMA system.

Distribution. Since IRMA is an interdisciplinary research project involving several institutes at the Aachen University of Technology, the development of feature extraction and classification methods will be done at several sites. Furthermore, it is desirable to access the computation power of a whole network cluster, especially for tasks like image processing. Therefore, the IRMA system provides

- automated method transfer to keep all sites up-to-date and
- automated job distribution to balance the load of computations.

Transparency. The IRMA system has to offer transparency regarding aspects from the viewpoint of a distributed system as well as from the user’s point of view, e.g., to allow the programmer to concentrate on the medico-diagnostic aspect of his algorithm, in particular

- location and access transparency for images, methods, and feature values,
- replication transparency for methods in development,
- concurrency transparency for worklist processing and feature generation,
- system transparency at method implementation time, and
- job distribution transparency when issuing a query.

5.2 Distributable Resources

The IRMA database contains information about all hosts that are part of the system. This information is stored in the table *computer*, mapping the host's Internet address to its IRMA-ID. Large data objects are kept outside the database only storing their location in a pair consisting of the host's IRMA-ID and the file system path (which has to be accessible via NFS or FTP for all other IRMA hosts). The pair is stored in the database (DB) table *source* and each entry referring to a file object links to a source entry.

Images. Beside DICOM data, secondary digitized images have to be processed. Using the IRMA location management, image data can be stored on various file servers. Image access is provided via an image class, hiding all file format specific I/O from the user (Fig. 3). The name attribute holds the filename of the image. Additional administrative data regarding the image format is stored in separate table entries to perform checks (e.g. whether a certain method is applicable to the image) without the need to access the image file itself.

Methods. In IRMA, a method can either transform a feature vector (derived from an image) into a new feature vector (feature extraction) or evaluate one or a set of feature vectors for classification or query purposes. The programmer using the IRMA system has to implement only the transformation or evaluation

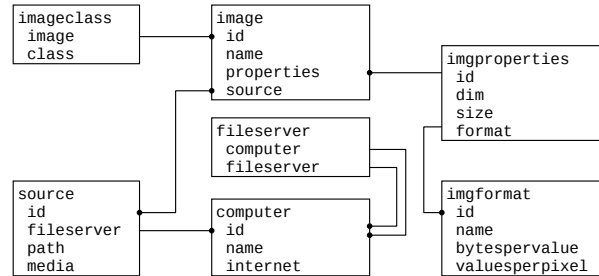


Fig. 3. IRMA tables that contain image-related data: The table *image* stores the image name and links to data entries in *imgproperties* and in *source*. The table *source* stores information about the physical file location, while entries in *imageclass* contain class information. Circles indicate that the attribute can be referenced by multiple entries inside the table it is connected to.

algorithm itself – all other parts of the method are system-generated and automatically linked to the algorithm. At the present state, six method types have been defined:

- *Local methods* typically produce image features on a per-pixel basis. A local method can be applied to each image (or feature vector) isolated from its application to other images.
- *Global methods* produce a constant number of features per image. Like local methods, they are applied to each image independently.
- *Universal methods* extract features calculated from a set of images. Statistically driven feature reduction like the KLT or the LDA are examples for this type of methods.
- *Distance methods* calculate a float type value representing a distance measure between two feature vectors.
- *Training methods* calculate a feature vector representing parameters for a classifier method. They are basically universal methods since they take into account features from the query view’s reference images. This method type was introduced to allow the separation between feature extraction and classification: a generic training method, for example an expectation maximization training for a mixture density based statistical classifier, can be applied to a variety of different features.
- *Classifier methods* apply a decision rule based on trained parameters to a sample feature vector.

Regarding the nature of the method types, it is easy to define a suitable programming interface for each of them. Consistent database access (fetching and storing features) and net-wide job distribution are handled automatically and are fully transparent to the method programmer. Via the definition of method dependencies, the system can automatically determine which intermediate features have to be calculated. Since the IRMA environment is heterogeneous, methods are distributed as source code. When a new method needs to be installed, the IRMA system fetches the required source files and automatically initializes the make process to create the binary executable on the actual host (Fig. 4). Note that an executable can be used by more than one method, e.g. by using different method parameters in each method. The parameter string also allows variable place-holders (via a C format-string like syntax).

Feature values. In the IRMA system, feature values are uniquely identified by a 3-tuple consisting of

- the image, where the feature value was derived from,
- the method, that was used to calculate the feature value and
- the query view, which represents the context of the feature (see below).

Beside basic data types, which are stored inside the DB, IRMA offers the file feature to efficiently access a large number of features. The definition of the file format is left to the programmer for largest flexibility. In addition, the data type

image feature can be used for query refinement and relevance feedback during browsing. The image and file features' physical locations are handled the same way as for images or source files.

Query views. By means of pre-defined and offline combined feature sets, the query can incorporate the medical background. Query views define

- the methods that are combined to a certain set of features as well as their parameters. This enables to apply the same method for local feature extraction in several query views with different parameters avoiding multiple storage;
- the method to calculate the query result. This is based on either evaluation or comparison of feature values.
- the set of images the query is based on. These *reference images* allow to train a classifier and also enable category-specific processing within the IRMA architecture.

If the QBE paradigm is used, a *sample image* has to be specified. In this case, either a distance measure (which will be embedded into a nearest neighbor classifier by the system) or a pair consisting of a training method and a classifier method can be used to calculate the query result (Fig. 5). The usage of query views enables easy experiment verification, and concurrent use of the system in production while other parts are being tested or evaluated.

Jobs. A single method invocation for feature extraction or evaluation (depending on the method type) is processing one or multiple feature vectors. Pending jobs are stored inside the *worklist* table that is accessed by all daemons (see Sec. 5.3) of IRMA (Fig. 6). Note the relationship between the *worklist* table and the *feature* table (which holds a 3-tuple entry once the corresponding job was processed).

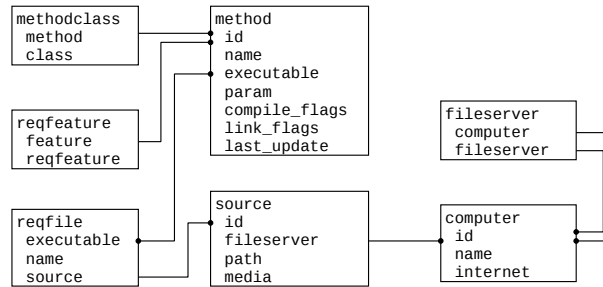


Fig. 4. IRMA tables that contain method-related data: The table *method* stores information about method parameters and the executable. The source files for the executable are stored in *reqfile*, their location is handled the same way as for images. The table *reqfeature* allows the definition of method dependencies.

5.3 IRMA Daemon

The IRMA daemon integrates its host into the IRMA cluster processing pending jobs, which are hold in the worklist table. Each daemon polls this table for new jobs. Jobs can be generated explicitly or implicitly (e.g. when a user initiates a query). Through the use of method dependencies, a job may spawn other jobs, e.g. if a required feature has not been computed yet. For each accepted job, the IRMA daemon checks whether the corresponding method is already running on this host. If required, it initiates the method transfer and installs the method. Thereafter, the daemon starts a child process running the method. The child takes over the job, fetches the required images or features (or initiates their generation by adding jobs to the worklist) and builds the method parameters according to the query view context (Fig. 6). Upon job completion, the child will poll the worklist table for jobs requesting the same method.

6 Preliminary Results

Currently, the database system in use is PostgreSQL. Implementation was done using C++. The file I/O relies on several free image libraries like “libjpeg”, “libtiff”, and “DICOM3-tools”. This provides easy portability across a wide range of UNIX environments.

In first practical experiments, the categorization step was evaluated on 1,617 images of six classes taken from daily routine. The best classification error rate of 8.0% was achieved using invariant distance measures within a statistical framework, which means a relative improvement of 42% with respect to the baseline statistical system with 14.0% error rate and a relative improvement of 56% with respect to the Euclidean distance nearest neighbor error rate of 18.1% [16].

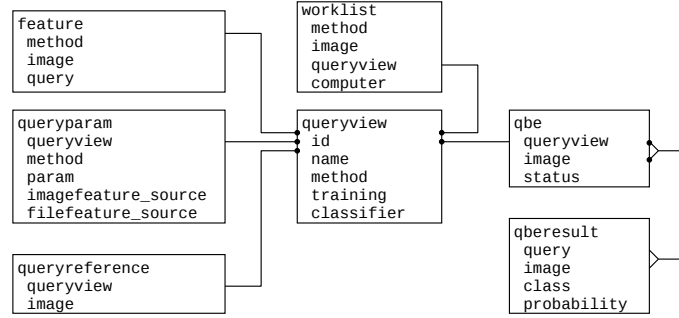


Fig. 5. IRMA tables used to store query view-related parameters: The table *queryreference* holds the reference set, *queryparam* stores the method parameters. When using QBE, it is possible to use a query view for multiple samples. This separation is done via multiple entries in *qbe*, which reference the same query view. For each QBE, *qberesult* stores the result.

7 Discussion

The IRMA approach introduces multi-step CBIR for medical use, which models abstract expert knowledge. All images are processed from the *raw data layer* (see Fig. 1). The information gained here is added to an image in the form of global image features and results in a second semantic layer called *registered data layer*. In contrast to other CBIR systems, this IRMA layer allows queries across all kinds of medical images regardless of modality, region or orientation. The next layer is called *feature layer* and provides the main advantage when comparing IRMA with other approaches. The separation of feature extraction and selection offers category- and query-dependent information extraction. According to CHU et al. [17], the layer obtained from the indexing step is called *scheme layer*. It provides a data structure describing spatial and hierarchical relationship of blobs. With respect to the later retrieval task, this normalization of content description is necessary. The blob structure is a general abstraction of the image content which has to be linked to medical terminology and knowledge. This is done in the *object layer* where medical objects are identified with subsets of the hierarchical blob structure. Thus an IRMA query mainly consists of finding corresponding subsets in different blob structures representing the images.

This framework is a general approach subsuming common CBIR systems. It is not focused on a certain modality or medical aspect and, therefore, it requires a flexible development platform to introduce specific methods on each layer, and to deal with advances in diagnostic insight. Consequently, aspects of application have to be considered. The image archive consists of large files and in a big hospital, this files might spread over several hosts. Hence, basic architecture of IRMA is based on a database system and has been implemented as a distributed development and runtime platform. The database administrates image process-

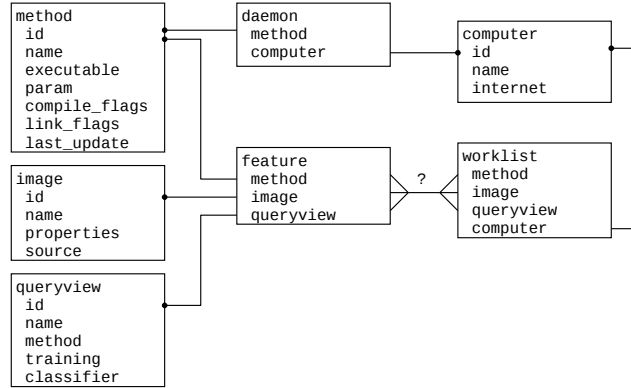


Fig. 6. IRMA tables used for job processing. The table *daemon* contains all started methods (an its host), *worklist* holds pending jobs. For each job, a daemon first checks whether the feature is already present (table *feature*). If not, the method is started with the context defined by the corresponding *queryview* entry.

ing methods as well as data structures. It offers different views on stored data according to each semantic layer of the IRMA framework. The distribution of all database objects (methods as well as data) demands sophisticated solutions for persistence and consistence of those objects but offers a great degree of freedom for development and application. The system provides automated method transfer, distributed computation and transparent data access. Consistence is achieved by timestamped views, where universal methods operate on worklists based on the state of the system at the moment of the methods invocation. Persistence is ensured by an image-based indexing scheme, where the key of methods or raw data serves as a universal index. The development group can be distributed over several physical sites. And finally, the users can extend the functionality of the system without communication overhead via a standardized method interface.

There still remains the problem of introducing new data by a physician. Here, an interface is required which supports the entry of prototypes and blob hierarchies for the *registered data layer* and the *object layer*. Furthermore, the context specific feature sets as stated in Sec. 4.4 have to be specified. An adequate solution for relevance feedback between the semantic object definition in the *object layer* and the *feature layer* is also needed.

References

1. S. T. C. Wong and D. A. Tjandra, "A Digital Library for Biomedical Imaging on the Internet," *IEEE Comm Mag*, 1999; 1:84–91.
2. S. L. Lou, D. R. Hoogstrate, and H. K. Huang, "An Automated PACS Image Acquisition and Recovery Scheme for Image Integrity Based on the DICOM Standard," *Comp Med Imag Graph*, 1997; 21(4):209–218.
3. H. D. Tagare, C. C. Jaffe, and J. Dungan, "Medical Image Databases: A Content-Based Retrieval Approach," *JAMIA*, 1997; 4:184–198.
4. C.-R. Shyu, C. E. Brodley, A. C. Kak, A. Kosaka, A. M. Aisen, and L. S. Broderick, "ASSERT: A Physician-in-the-Loop Content-Based Retrieval System for HRCT Image Databases," *Comp Vis Imag Understand*, 1999; 75(1/2):111–132.
5. P. Korn, N. Sidiropoulos, C. Faloutsos, E. Siegel, and Z. Protopapas, "Fast and Effective Retrieval of Medical Tumor Shapes," *IEEE Trans KDE*, 1998; 10(6):889–904.
6. W. Niblack, R. Barber, W. Equitz, M. Flickner, P. Yanker, and J. Ashley, "The QBIC-Project: Querying Images by Content Using Color, Texture and Shape," *Procs SPIE*, 1993; 1908:173–187.
7. A. Smeulders, M. Worring, S. Santini, A. Gupta, and R. Jain, "Content Based Image Retrieval at the End of the Early Years," *IEEE Trans PAMI*, 2000; 22(12):1349–1380.
8. M. Nappi, G. Polese, and G. Tortora, "FIRST: Fractal Indexing and Retrieval System for Image Databases," *Imag Vis Comp*, 1998; 6:1019–1031.
9. J. K. Wu, A. D. Narasimhalu, B. M. Mehtre, C. P. Lam, and Y. J. Gao, "CORE: A Content-Based Retrieval Engine for Multimedia Information Systems," *Multimed Sys*, 1995; 3:25–41.

10. C. Carson, M. Thomas, S. Belongie, J.M. Hellerstein, and J. Malik, "Blobworld: A System for Region-Based Image Indexing and Retrieval," Procs 3rd Int Conf on Visual Information Systems, Springer-Verlag, 1999.
11. A. Pentland, R. W. Picard, and S. Scarloff, "Photobook: Content-Based Manipulation of Image Databases," Procs SPIE, 1994; 2185:34–47.
12. T. M. Lehmann, B. Wein, J. Dahmen, J. Bredno, F. Vogelsang, M. Kohnen, "Content-based Image Retrieval in Medical Applications: A Novel Multi-Step Approach," Procs SPIE, 2000; 3972:312–320.
13. T. Wendler, B. Wein, and R. van den Broek, "Deliverable USINT in the EurIPACS-Project, Technical Report, AIM," Ref. XIII, European Commission, Brussels, Feb. 1994.
14. J. Bredno, M. Kohnen, J. Dahmen, F. Vogelsang, B. Wein, and T. M. Lehmann, "Synergetic Impact Obtained by a Distributed Developing Platform for Image Retrieval in Medical Applications (IRMA)," Procs SPIE, 2000; 3972: 321–331.
15. J. Goldstein, R. Ramakrishnan, U. Shaft, and J.-B. Yu, "Using Constraints to Query R*-Trees", Technical Report No. 1301, Dept. of Computer Science, University of Wisconsin, Madison, Feb. 1996.
16. D. Keysers, J. Dahmen, H. Ney, B. Wein, T. M. Lehmann, "A Statistical Framework for Model-Based Image Retrieval in Medical Applications," J Electron Imag, 2002; accepted for publication.
17. W. W. Chu, C.-C. Hsu, A. F. Crdenas, and R. K. Tiara, "Knowledge-Based Image Retrieval with Spatial and Temporal Constructs," IEEE Trans KDE, 1998; 10(6):872–888.