

(19) World Intellectual Property Organization
International Bureau



(43) International Publication Date
13 April 2006 (13.04.2006)

PCT

(10) International Publication Number
WO 2006/039686 A2

(51) International Patent Classification:
G06F 7/00 (2006.01)

(21) International Application Number:
PCT/US2005/035612

(22) International Filing Date: 3 October 2005 (03.10.2005)

(25) Filing Language: English

(26) Publication Language: English

(30) Priority Data:
60/615,085 1 October 2004 (01.10.2004) US

(71) Applicant (for all designated States except US): **UNIVERSITY OF SOUTHERN CALIFORNIA** [US/US]; 3716 S. Hope St., Suite 313, Los Angeles, California 90007-4344 (US).

(72) Inventors; and

(75) Inventors/Applicants (for US only): **XIE, Hua** [—/US]; 3740 McClintock Avenue, Los Angeles, California 90089 (US). **ORTEGA, Antonio** [ES/US]; 806 S. Stanley Avenue, Los Angeles, California 90036 (US).

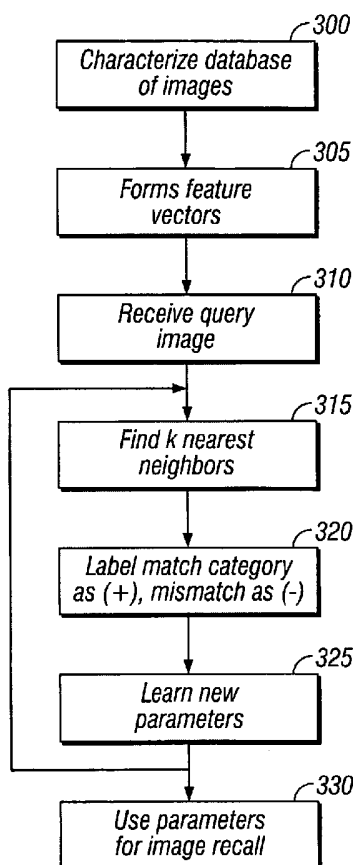
(74) Agent: **HARRIS, Scott, C.**; FISH & RICHARDSON P.C., P.O. Box 1022, Minneapolis, Minnesota 55440-1022 (US).

(81) Designated States (unless otherwise indicated, for every kind of national protection available): AE, AG, AL, AM, AT, AU, AZ, BA, BB, BG, BR, BW, BY, BZ, CA, CH, CN, CO, CR, CU, CZ, DE, DK, DM, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, HR, HU, ID, IL, IN, IS, JP, KE, KG, KM, KP, KR, KZ, LC, LK, LR, LS, LT, LU, LV, LY, MA, MD, MG, MK, MN, MW, MX, MZ, NA, NG, NI, NO, NZ, OM, PG, PH, PL, PT, RO, RU, SC, SD, SE, SG, SK, SL, SM, SY, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, YU, ZA, ZM, ZW.

(84) Designated States (unless otherwise indicated, for every kind of regional protection available): ARIPO (BW, GH, GM, KE, LS, MW, MZ, NA, SD, SL, SZ, TZ, UG, ZM, ZW), Eurasian (AM, AZ, BY, KG, KZ, MD, RU, TJ, TM), European (AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HU, IE, IS, IT, LT, LU, LV, MC, NL, PL, PT, RO, SE, SI, SK, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, ML, MR, NE, SN, TD, TG).

[Continued on next page]

(54) Title: USER PREFERENCE TECHNIQUES FOR SUPPORT VECTOR MACHINES IN CONTENT BASED IMAGE RETRIEVAL



(57) Abstract: Searching multimedia information which allows determining preferences based on very little amounts of data. The preferences are nonparametrically determined. Each preference is quantized into one of a plurality of bins. By doing the quantization, the distances between positive and negative samples are increased. The quantization amount may change depending on the number of samples which are used. The quantization can be used in a support vector machine or the like.

WO 2006/039686 A2



Published:

— without international search report and to be republished
upon receipt of that report

For two-letter codes and other abbreviations, refer to the "Guidance Notes on Codes and Abbreviations" appearing at the beginning of each regular issue of the PCT Gazette.

USER PREFERENCE TECHNIQUES FOR SUPPORT VECTOR MACHINES IN
CONTENT BASED IMAGE RETRIEVAL

CLAIM OF PRIORITY

[0001] This application claims priority under 35 USC §119(e) to U.S. Patent Application Serial No. 60/615,085, filed on October 1, 2004 the entire contents of which are hereby incorporated by reference.

BACKGROUND

[0002] Multimedia information is often stored on the Internet. Conventional search engines are limited in their ability to access this information. Content based information retrieval systems have been used for automatically indexing and accessing this kind of information. These systems access and index large amounts of information. Multiple features including color, texture, shape and the like are extracted from the query signals. Retrieval is then performed using a similarity matching, where the different features are matched against similar patterns. Given an input feature pattern, the matching attempts to search for similar patterns within the database.

[0003] Content based image retrieval systems leave a semantic gap between the low level features that they index, and the higher-level human concepts. Many different attempts have been made to design techniques that introduce the user into the searching loop, to enable the system to learn a user's particular preferences of query.

[0004] Relevance feedback can be used to allow the user to interactively tune the system to their own interest. This kind of feedback can be used to assess whether certain proposed images are relevant to their query or not relevant. The system learns from the examples using a machine learning technique, which is used to tune the parameters of the search. It returns a new set of similar images, and iteratively repeats the process until the user is satisfied with the result. The action is a query updating scheme, and hence can be regarded as a machine learning task.

[0005] Techniques of relevance feedback in content based information retrieval systems have conventionally used feature re-weighting. The weights associated with each feature for a K nearest neighbor classifier are adjusted based on feedback. Those features that are the best at discriminating between positive and negative samples receive a more significant weight for the distance computation.

[0006] Another technique is to set up an optimization problem as a systematic formulation to the relevance feedback problem. The goal of the optimization problem is to find the optimal linear transformation which maps the feature space into a new space. The new space has the property of clustering together positive examples, and hence makes it easier to separate those positive examples from the negative examples.

[0007] Support vector machines may be used for the relevance feedback problem in a content based retrieval system. The support vector machines or SVMs may be incorporated as an automatic tool to evaluate preference weights of the relative images. The weights may then be utilized to compute a query refinement. SVMs can also be directly used to derive similarity matching between different images.

[0008] Different techniques have been used in the context of support vector machine methods. The kernel function of such a machine usually has a significant effect on its discrimination ability.

SUMMARY

[0009] A technique is disclosed that enables searching among multimedia type information, including, for example, images, video clips, and other.

[0010] An embodiment describes a kernel function which is based on information divergence between probabilities of positive and negative samples that are inferred from user preferences for use in relevance feedback in content based image retrieval systems. A special framework is also disclosed for cases where the data distribution model is not known a priori, and instead is inferred from feedback. The embodiment may increase the distance between samples non-linearly, to facilitate learning. An embodiment uses probabilistic techniques, e.g., where underlying probabilistic models are learned based on quantization and used within a machine, such as a support vector machine, that determines distances between multimedia content, such as images.

DESCRIPTION OF DRAWINGS

[0011] FIG. 1 is a block diagram of an exemplary system;

[0012] Figure 2 shows a technique used by the support vector machine of an embodiment;

[0013] Figure 3 shows a flowchart of operation.

DETAILED DESCRIPTION

[0014] A block diagram of the overall system is shown in figure 1. A server 100 stores a plurality of multimedia information 106 in its memory 105. The multimedia information such as 106 may be indexed, or may be addressed without

indexing. The server is connected to a channel 125, which may be a private or public network, and for example may be the Internet. At least one client 110 is connected to the channel 125 and has access to the content on the server.

[0015] A query is sent from the client 110 to the server 100. The query may be modeled and modified using the techniques described in this application. The query may be, for example, a request for multimedia information. For example the first query could be a "query by example", i.e., the user would first identify an image that is representative of what the user is looking for, the system would then search for images that exhibit similarity in the feature space to this "example" image. Alternatively the user may request multimedia data based on a high level concept, for example the user may request an image of a certain type, e.g., an image of "beach with a sunset". Yet another possibility is for the user to provide a non-textual description of what the user is searching for, e.g., a sketch that represents the kind of image that the user is looking for. Similar techniques can be used to search for other kind of multimedia information, e.g., videos and animations of conventional kinds (.avi's, flash, etc), sounds of compressed and uncompressed types, and others. In any of these cases, techniques are disclosed to improve the quality of the information retrieved (where quality in this context is assessed by how close the resulting query responses are to the user's interest). These improvements are achieved by successive

user relevance feedback about each of the successive query responses provided by the system. The server and client may be any kind of computer, either general purpose, or some specific purpose computer, such as a workstation. The computer may be a Pentium class computer, running Windows XP or Linux, for example, or may be a MacIntosh computer. The programs may be written in C, or Java, or any other programming language. The programs may be resident on a storage medium, e.g., magnetic or optical, e.g. the computer hard drive, a removable disk or other removable medium. The programs may also be run over a network.

[0016] The techniques disclosed herein employ an empirical model to capture probabilistic information about the user's preferences from positive and negative samples. In an embodiment, a kernel is derived. The kernel is called the user preference information divergence kernel. This scheme is based on no prior assumptions about data distribution. The kernel is learned from the actual data distribution. Relevance feedback iterations are used to improve the kernel accuracy.

[0017] An embodiment uses a support vector machine that operates using a non-parametric model - that is, one where the model structure is not specified a priori, but is instead determined from data. The term nonparametric does not require that the model completely lacks parameters; rather, the number and nature of the parameters is flexible and not fixed in advance. The machine may be running on either or both of the

client 110 or server 100, or on any other computer that is attached to the channel 125. Other machines, such as neural networks, may alternatively be used.

[0018] Support vector machines operate based on training sets. The training set contains L observations. Each observation includes a pair: a feature vector

[0019] Where $x_i \in R^n$, i can extend between 1 and L , and an associated semantic class label y_i . The class label can be 1, to represent relevance, or -1 to represent irrelevance. The vector x can be a random variable drawing from a distribution, which may include probabilities:

$\{P(x|y = +1), P(x|y = -1)\}.$

[0020] In general there can be more classes, and the feedback can take other forms. For example rather than just providing feedback as relevant or irrelevant, the user could provide a number to quantify the degree of relevance of each object.

[0021] The goal of the relevance training is to learn the mapping g , between x and y , based on the training data.

[0022] The optimal mapping may be considered as a maximum likelihood classifier:

$$g(x) = \arg \max_i P(x|y = i) \dots (1)$$

[0023] in a typical scenario of this type, there is a lot of data to carry out the training. It has been found, however, that the number of training samples in practical applications

of content-based retrieval may not be sufficiently large relative to the dimensionality of the feature vector to estimate a probabilistic model of the data that can support a conventional maximum likelihood classifier. Moreover, the lack of sufficient training data may make it unrealistic to use traditional density estimation techniques. Techniques are disclosed herein to estimate the probabilistic model even with this problem. These techniques use non-parametric methods for density estimation and the resulting models are used to introduce a modified distance in the feature space, which has the property that it increases the distance between positive and negative samples.

[0024] The labeled training set is denoted as:

$\{x_i, y_i\}, i=1, \dots, L, y_i \in \{-1, +1\}, x_i \in R^n$. Support Vector Machines are classification tools that separate data by using hyper-planes such that a maximum margin is achieved in the separation of the training data. Figure 2 illustratively shows how all vectors labeled +1 lie on one side and all vectors labeled -1 lie on the other side of the hyperplane. Mathematically:

$$\begin{aligned} w \cdot x_i + b &\geq +1 \text{ for } y_i = +1 \\ w \cdot x_i + b &\leq -1 \text{ for } y_i = -1 \end{aligned} \tag{2}$$

[0025] where w is normal to the hyperplane H . The training vectors that lie on hyperplanes $H_0: w \cdot x_i + b = 1$ and $H_1: w \cdot x_i + b = -1$, are referred to as support vectors. It can

be shown that the margin between the two hyperplanes H_0 and H_1 is simply $\frac{2}{\|w\|}$. Thus, searching for the optimal separating

hyperplane becomes a constrained optimization problem:

minimizing $\|w\|^2$ subject to the constraints. Lagrange multipliers are used to maximize the Lagrangian objective function with respect to positive Lagrange multipliers $\alpha_i, i=1, \dots, L$, subject to constraints $\sum_i \alpha_i y_i = 0$.

$$\max \left(\sum_i \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j y_i y_j x_i \cdot x_j \right) \quad (3)$$

[0026] The training samples might not be linearly separable in the original space. In an embodiment, the data is first mapped to some other Euclidean space H (possibly infinitely dimensional) using a mapping $\Phi: \mathcal{X} \mapsto H$. The training algorithm only depends on the inner products between sample vectors. A kernel function K can be defined such that $K(x_i, x_j) = \Phi(x_i) \cdot \Phi(x_j)$. Then the inner product $x_i \cdot x_j$ is replaced by $K(x_i, x_j)$ everywhere in the training algorithm. This avoids the need to explicitly compute the mapping Φ . The resulting classifier takes the form of $g(x) = \sum_{i=1}^{N_s} \alpha_i y_i K(x_i, x) + b$. $\{\alpha_i, i=1, \dots, N_s\}$ and b are the parameters that can be learned using quadratic programming. N_s is the number of support vectors.

[0027] Much of the flexibility and classification power of support vector machines resides in the kernel function, since these make it possible to discriminate within challenging data sets, e.g., those where linear discrimination may be suboptimal. Typical kernel functions include linear, polynomial and radial basis functions ("RBF"):

$$\text{Linear: } K(x, z) = x \cdot z \quad (4)$$

$$\text{Polynomial: } K(x, z) = (Ax \cdot z + B)^p \quad (5)$$

$$\text{Radial Basis: } K(x, z) = e^{-\gamma \|x - z\|^2} \quad (6)$$

[0028] Where z is another vector of the same dimension as x , and $(\cdot)$ denotes the inner product of two vectors. A , B , p and γ are constants that are set a priori. These kernels are generic and do not explicitly take into account the statistics of user-provided feedback information available in content-based retrieval systems. Thus, if using a support vector machines in such a system, one would have to select a kernel a priori and then the performance of the system will depend significantly on the nature of the feedback provided by the user.

[0029] In the embodiment, user feedback can be exploited in order to create a modified kernel function.

[0030] Content based retrieval systems rely on a low level similarity metric. However, the information obtained from the user through relevance feedback is often at a high level, since those perceptual interpretations of an image depend greatly on subjective feelings of the user, usage context, and the

application. Generic models are not applicable to all scenarios. Accordingly, the techniques disclosed herein assist in non-parametrically learning the user's preferences empirically based on probabilistic information obtained from training data obtained from the user's feedback. This information is then used to derive a kernel that is customized for the specific user and task.

[0031] For mathematical purposes, each image is assumed to be represented by one feature vector $x \in R^n$.

[0032] The relevance feedback problem can then be regarded as a machine learning task whose goal is to infer the user's preference based on the information that is learned from the user labeled data. The information that is learned is arranged into a multi-dimensional feature vector

$$x = (\chi_1, \chi_2, \dots, \chi_n)^T.$$

[0033] The marginal probability of each label for each component of the feature vector is defined as

$$\{P(y=+1|\chi_1), P(y=-1|\chi_1)\}.$$

[0034] Where x_1 is the 1-th component of the feature vector. These marginal distributions can be empirically estimated from the training data obtained from successive user feedback samples. Difficulties in the estimation, however, may be expected because of two reasons: first x_1 can in general take values in either a large discrete set or over a continuous

range. In addition, limited amounts of training data are available in any specific situation. The techniques described herein preferably use a non-parametric probability estimation approach obtained by accumulating successive iterations of user feedback.

[0035] For each feature vector component χ_l , a quantizer A_l is defined that includes B_l reconstruction levels with B_l-1 decision boundaries denoted as $\{b_1, \dots, b_{B_l-1}\}$. The probabilities $\{P(y=+1|\chi_l), P(y=-1|\chi_l)\}$ are estimated by counting the number of samples that fall in each bin:

$$P(y=\pm 1|\chi=r_{lk}) = \frac{\sum_{i=1}^L 1(y_i=\pm 1) 1(|\chi_{il}-r_{lk}|\leq \Delta_{lk})}{\sum_{i=1}^L 1(|\chi_{il}-r_{lk}|\leq \Delta_{lk})} \quad (7)$$

[0036] where the indicator function $1(\cdot)$ takes value one when its argument is true and zero otherwise. L is the number of labeled training data. χ_{il} is the l -th component of training vector $\mathbf{x}_i$. Δ_{lk} is the size of the quantization interval along dimension l centered at reconstruction value r_{lk} . For those quantization bins where there is no training data, the probability can be set to zero since they make no contribution to differentiating classes. Obviously, the design of quantizers A_l s plays an important role in probability estimation. This embodiment uses a simple uniform quantization scheme where all quantization bins in a given feature

dimension have the same size Δ_{lk} , which is computed from the dynamic range of the data $[\max(\chi_l), \min(\chi_l)]$. This range changes from iteration to iteration, where the size of the range gets smaller as the number of iterations increases. The the number of quantization levels applied B_l is applied as:

$$\Delta_{lk} = \Delta_l = \frac{\max(\chi_l) - \min(\chi_l)}{2xB_l} \quad (8)$$

[0037] In addition to these techniques, other techniques such as K-nearest-neighbor techniques, least squares estimation, and others can be used.

[0038] Moreover, with successive relevance feedback iterations, the amount of available training data increases. Thus, it is possible to estimate more reliably a larger number of model characteristics. For example, as the amount of available training data increases, it is possible to increase the number of quantization bins used to represent the model, thus forming smaller quantization bins. In an embodiment, the number of bins is increased in successive relevance feedback iterations, in order to provide an increasingly accurate representation of the underlying user preference model.

[0039] The probability model described above can view a feature $x = (\chi_1, \chi_2, \dots, \chi_n)'$ as a sample drawn from a random source, which has relevance statistics given by $P^+(x) = (p_1^+, \dots, p_n^+)$ and $P^-(x) = (p_1^-, \dots, p_n^-)$. $p_i^\pm = P(y = \pm 1 | \chi_i)$ are estimated by quantizing the

component χ_1 using A_1 based on the training data obtained from relevance feedback.

[0040] The distance between x and z , another feature vector with probability vectors $Q^+ = (q_1^+, \dots, q_n^+)$ and $Q^- = (q_1^-, \dots, q_n^-)$. A distance is defined based on the Kullback-Leibler divergence of their probability vectors P and Q :

$$D(x\|z) = \sum_{i=1}^n p_i^+ \log\left(\frac{p_i^+}{q_i^+}\right) + \sum_{i=1}^n p_i^- \log\left(\frac{p_i^-}{q_i^-}\right) \quad (9)$$

$0 \times \log(0) = 0$ is bounded by continuity arguments. Since the KL divergence is not symmetric, equation (9) can be used to form a symmetric distance measure $D_s(x, z)$, as:

$$D_s(x, z) = D(x\|z) + D(z\|x). \quad (10)$$

[0041] The proposed user preference information divergency (UPID) kernel function in the generalized form of RBF kernels with the original Euclidean distance $d()$ replaced by the proposed distance of (10):

$$K(x, z) = e^{-pD_s(x, z)} \quad (11)$$

[0042] In an embodiment, the proposed distance can also be combined with other kernel forms (such as linear, polynomial, etc.) in such a way that the formed kernel satisfies Mercer's condition.

[0043] The distance (11) is a positive definite metric, thus the proposed UPID kernel satisfies Mercer's condition. As the model parameters α_i, b and N_s are learned from the training set,

we evaluate the likelihood that an unknown object $\mathbf{x}$ is relevant to the query by computing its score $f(\mathbf{x})$:

$$f(x) = \sum \alpha_i y_i K(x, x_i) + b \quad (12)$$

[0044] Where $\mathbf{x}_i$ is the i -th support vector and there are a total of N_s support vectors, which are obtained from the learning process. Larger scores make it more likely that the unknown object belongs to the relevant class and thus should be returned and displayed to the user.

[0045] In operation, the system may operate according to the flowchart of figure 3. The flowchart may be carried out in dedicated hardware or on a computer. The computer described herein may be any kind of computer, either general purpose, or some specific purpose computer such as a workstation. The computer may be a Pentium class computer, running Windows XP or Linux, or may be a MacIntosh computer. The programs may be written in C, or Java, or any other programming language. The programs may be resident on a storage medium, e.g., magnetic or optical, e.g., the computer hard drive, a removable disk or other removable medium. The programs may also be run over a network.

[0046] The database of images is characterized at 300. The features may be extracted by feature extraction algorithms such as those described in the literature. For example, the images may be represented in terms of color, texture and shape. Color features may be computed as histograms, texture features may be

formed by applying the Sobel operator to the image and histogramming the magnitude of the local image gradient. The shape feature may then be characterized by histogramming the angle of the edge. Using eight bins for each histogram, a 72 dimensional feature vector is formed at 305.

[0047] In the embodiment, the system maintains a database of images, where the image set is divided into different categories, and that database is characterized and vectored in 300 and 305. The embodiment divides the image set into the categories of sunsets, coasts, flowers, exotic cars, Maya and Aztec, fireworks, skiing, owls, religious stained-glass, Arabian horses, glaciers and mountains, English country Gardens, divers and diving, pyramids, and oil paintings. Of course, different categories could be used, but these are exemplary of the information that might be obtained. In some instances, e.g., in the retrieval of images from the Internet, there may not exist predefined semantic image categories; the system would have access to large unstructured collections of images or other multimedia objects, where each object would have an associated feature vector, but the set of feature vectors would not be otherwise structured.

[0048] Query feedback is based on the actual image categories. Quality of that retrieval result may be measured by precision and recall. Precision is the percentage of relevant objects that are to be retrieved to the query image. This may measure the purity of the retrieval. Recall is a

measurement of completeness of the retrieval, for example computed as a percentage of retrieved relevant objects in the total relevant set in the databases.

[0049] This query feedback is used for a first embodiment, which allows experimental evaluation of system. In another embodiment, which, the feedback would be whatever the user provides as feedback to the system. Either and/or both of these elements, can be provided as "feedback".

[0050] At 310, a query image is received. At 315, the image is compared. In an embodiment, a first iteration is carried out where the nearest neighbors are computed based on Euclidean distance, and in later iterations, proposed modified distance is used with SVMs to compute similarity between images. At 320, the positive matches are labeled as a positive, while the negative matches are labeled as negative. In the first embodiment, the nearest neighbors with the same category as the query image are labeled as positive, and success is measured when the matching image is in the same category as a query image. In the second embodiment, the user feedback is used to mark the images that are returned as positive or negative. In effect, therefore, this system is obtaining images, and being told to find images that are like the query image.

[0051] At 325, the system learns new model parameters from these images, and then returns to repeat the process beginning at 315.

[0052] For the techniques described above, the parameters α_i s and b are learned using equation 3 on the labeled images, and by using the classifier in equation 12 with the new set of parameters. The images with the highest score are the most likely ones to be the target images for the user. The new parameters are used to compute the similarity between any image and the query image using Equation 12. Finally, at 330, the learned parameters are used for image recall.

[0053] Although only a few embodiments have been disclosed in detail above, other embodiments are possible and the inventors intend these to be encompassed within this specification. The specification describes specific examples to accomplish a more general goal that may be accomplished in other way. This disclosure is intended to be exemplary, and the claims are intended to cover any modification or alternative that might be predictable to a person having ordinary skill in the art. For example, while the above has described this being carried out with a support vector machine, any kind of device that can carry out similar types of processing can alternatively be used.

[0054] Also, the inventors intend that only those claims which use the words "means for" are intended to be interpreted under 35 USC 112, sixth paragraph. Moreover, no limitations from the specification are intended to be read into any claims, unless those limitations are expressly included in the claims.

WHAT IS CLAIMED IS:

1. A method, comprising:
obtaining data indicative of positive matches between multimedia content and queries for multimedia content and indicative of negative matches between said multimedia content and said queries to multimedia content;
increasing a distance between said positive matches and said negative matches and using said increased distance to train a system for future queries.
2. A method as in claim 1, wherein said using said increased distance comprises compiling statistics of a user's specific matches based on said queries, and using said statistics for training said system.
3. A method as in claim 1, wherein said distance is a Euclidean distance.
4. A method as in claim 1 wherein said increasing a distance comprises quantizing said positive matches and said negative matches to obtain a non-parametric model of the statistics of a user's specific matches and defining a modified distance by applying said model in a way that results in increasing the distance between the positive matches and the negative matches.

5. A method as in claim 1, wherein said increasing a distance increases the distance nonlinearly.

6. A method as in claim 1,
wherein said increasing a distance comprises using a kernel function that weights positive matches between a multimedia search criteria and multimedia content with a first vector in a first plane, and weights negative matches between said search criteria and said content with a second vector in a second plane, and increasing a margin between said first plane and said second plane.

7. A method as in claim 1, wherein said method is carried out in a support vector machine.

8. A method as in claim 1, wherein said multimedia content is one of an image, a video and/or a sound.

9. A method as in claim 1, wherein said increasing a distance comprises quantizing positive and negative matches, and further comprising decreasing a size of quantization where additional samples are obtained.

10. A method as in claim 9, where the quantization is a vector quantization.

11. An apparatus, comprising:

a processor, obtaining data indicative of positive matches between multimedia content and queries for multimedia content and indicative of negative matches between said multimedia content and said queries to multimedia content, increasing a distance between said positive matches and said negative matches and using said increased distance in a search process for future queries.

12. An apparatus as in claim 11, wherein said processor forms a support vector machine.

13. An apparatus as in claim 11, wherein said processor forms a neural network.

14. An apparatus as in claim 11 further comprising a memory which stores statistics of a user's specific matches based on said queries as training information.

15. An apparatus as in claim 11, wherein said distance is a Euclidean distance.

16. An apparatus as in claim 11, wherein said processor quantizes said positive matches and said negative matches to obtain a non parametric model of the statistics of a user's specific matches and defining a modified distance by applying said model in a way that results in increasing the distance between the positive matches and the negative matches.

17. An apparatus as in claim 11, wherein said increasing a distance increases the distance nonlinearly.

18. An apparatus as in claim 11, wherein said processor is operative to receive further queries for multimedia information, and retrieve said multimedia information based on said queries and said training.

19. An apparatus as in claim 11, wherein said multimedia content is at least one of an image, a video and/or a sound.

20. An apparatus as in claim 11, wherein said processor increases a distance comprises by quantizing positive and negative matches, and further decreases a size of quantization when additional samples are obtained.

21. An apparatus as in claim 11 where the processor carries out vector quantization.

22. A method, comprising:

obtaining data indicative of positive and negative matches between multimedia content and queries for multimedia content, where the multimedia content includes at least one of an image, a sound and/or a video and indicative of negative matches;

increasing a distance between said positive matches and said negative matches by quantizing each of the matches to one of a plurality of quantization states;

increasing a number of quantization states based on an increase in a number of matches, to reduce a size of said quantization state; and

using quantized information to train a system for future queries.

23. A method as in claim 22, wherein said quantizing is a vector quantizing.

1/2

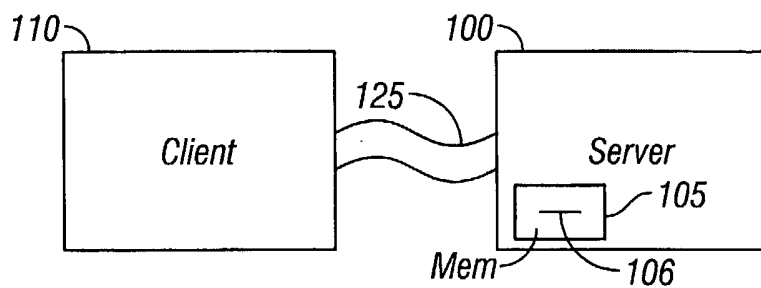


FIG. 1

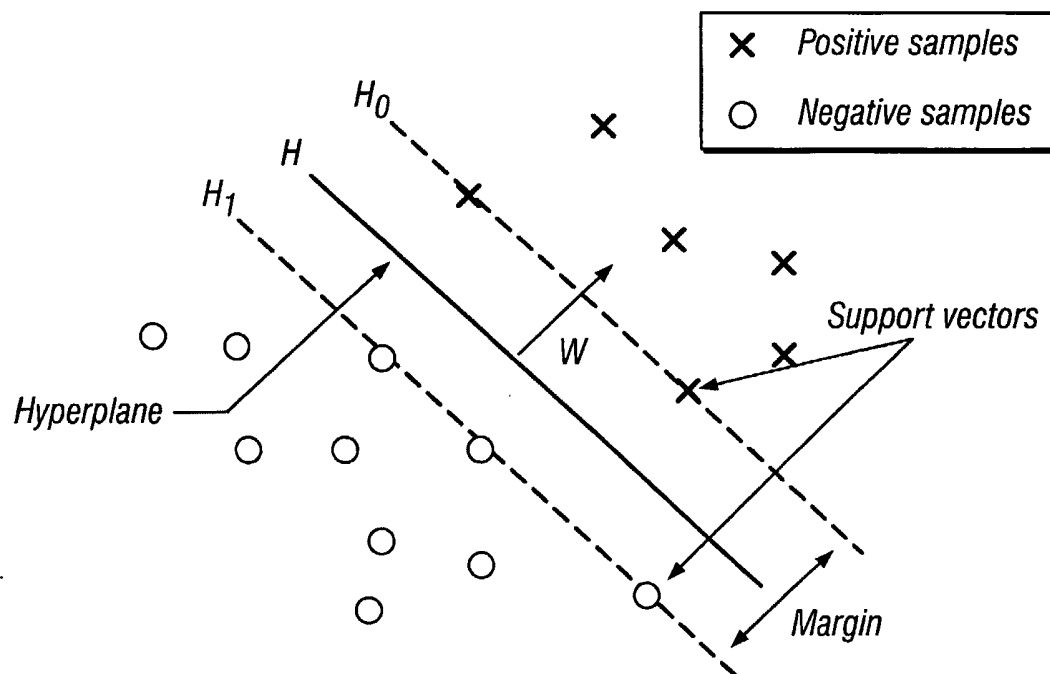


FIG. 2

2/2

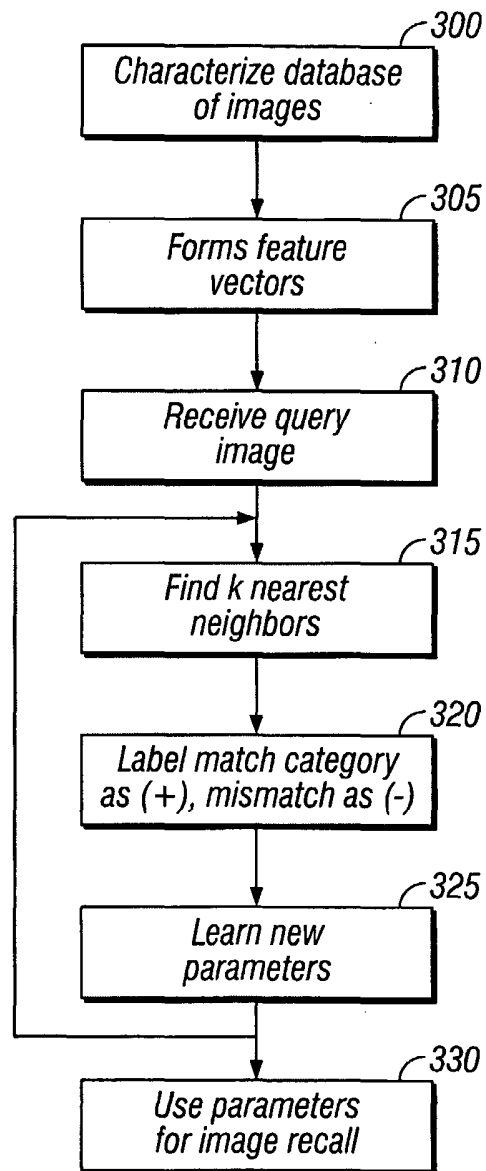


FIG. 3