



Graph regularized nonnegative sparse coding using incoherent dictionary for approximate nearest neighbor search



Yupei Zhang, Ming Xiang*, Bo Yang

Department of Computer Science and Technology, Xi'an Jiaotong University, China

ARTICLE INFO

Article history:

Received 4 August 2016

Revised 13 January 2017

Accepted 30 April 2017

Available online 1 May 2017

Keywords:

Nonnegative sparse coding

Incoherent dictionary learning

Laplacian regularization

Approximate nearest neighbor searching

Image retrieval

ABSTRACT

In this paper, we consider the problem of approximate nearest neighbor (ANN) retrieval with the method of sparse coding, which is a promising tool of discovering compact representation of high-dimensional data. A new study, exploiting the indices of the active set of sparse coded data as retrieval code, exhibits an appealing ANN route. Here our work aims to enhance the method via considering its shortages of the local structure of the data. Our primary innovation is two-fold: We introduce the graph Laplacian regularization to preserve the local structure of the original data into reduced space, which is indeed beneficial to ANN. And we impose non-negativity constraints such that the retrieval code can more effectively reflect the neighborhood relation, thereby cutting down on unnecessary hash collision. To this end, we learn an incoherent dictionary with both rules via a novel formulation of sparse coding. The resulting optimization problem can be provided with an available solution by the widely used iterative scheme, where we resort to the feature-sign search algorithm in the sparse coding step and exploit the method that uses a Lagrange dual for dictionary learning step. Experimental results on publicly available image data sets manifest that the rules are positive to promote the classification and ANN accuracies. Compared with several state-of-the-art ANN techniques, our methods can achieve an interesting improvement in retrieval accuracy.

© 2017 Elsevier Ltd. All rights reserved.

1. Introduction

Nearest neighbor retrieval aims to, from a large-scale data set, find a data subset that is most similar to a query sample. It is always a fundamental component in a wide range applications, including dimensionality reduction [1], pattern classification [2] and image retrieval [3,4]. Nowadays high-dimensional data collection, however, results in expensive computational cost and deteriorative performance of traditional routes due to the curse of dimensionality [5]. Recent studies are mainly devoted to hunting for approximate nearest neighbor (ANN) instead of exact nearest neighbor, since the meaning of “exact” is highly subjective [6]. The main idea of ANN is to find the nearest neighbor with high probability “only” rather than probability 1, and thereby permits us to develop various effective and efficient approaches. Note that we in this paper evaluate similarity using the Euclidean distance, which is relevant to many applications.

The sketchy treatment scheme for ANN search from a large-scale data set, illustrated in Fig. 1, can be partitioned into the fol-

lowing steps: First it is prohibitive to directly manipulate the original data set and sampling a training set is indeed necessary, since the original data set is enormous and constantly changing. Hence, a model of feature mapping capable of generalization in second step should be constructed for dealing with unseen data and the curse of dimensionality. And then, all samples in the original data set and query set are projected into new representation space in which ANN search is easier to be carried out, according to the constructed model. Finally, ANN searches of all query data are carried out via one-to-one matching where one can exploit various methods of indexing for speeding up. Among these steps, most attention is paid to learning an effective model of feature transformation and knitting an efficient indexing structure based on the retrieval codes.

Recall that the theory of sparse representation and dictionary learning (SRDL) is to partition data space into multiple non-orthogonal subspaces and represent each sample into a subspace [7]. Therefore, in this paper, we consider to tackle the problem of ANN search based on this promising technique. The most appealing advantages of utilizing SRDL for ANN are: the high generalization ability to unseen sample [8] and the concise data representation for the compact space-partition [9]. Concretely, the dictionary atoms learned from data are high-level such that all samples,

* Corresponding author.

E-mail addresses: xjtuiezhyp@stu.xjtu.edu.cn (Y. Zhang), mxiang@mail.xjtu.edu.cn (M. Xiang), yangboo@stu.xjtu.edu.cn (B. Yang).

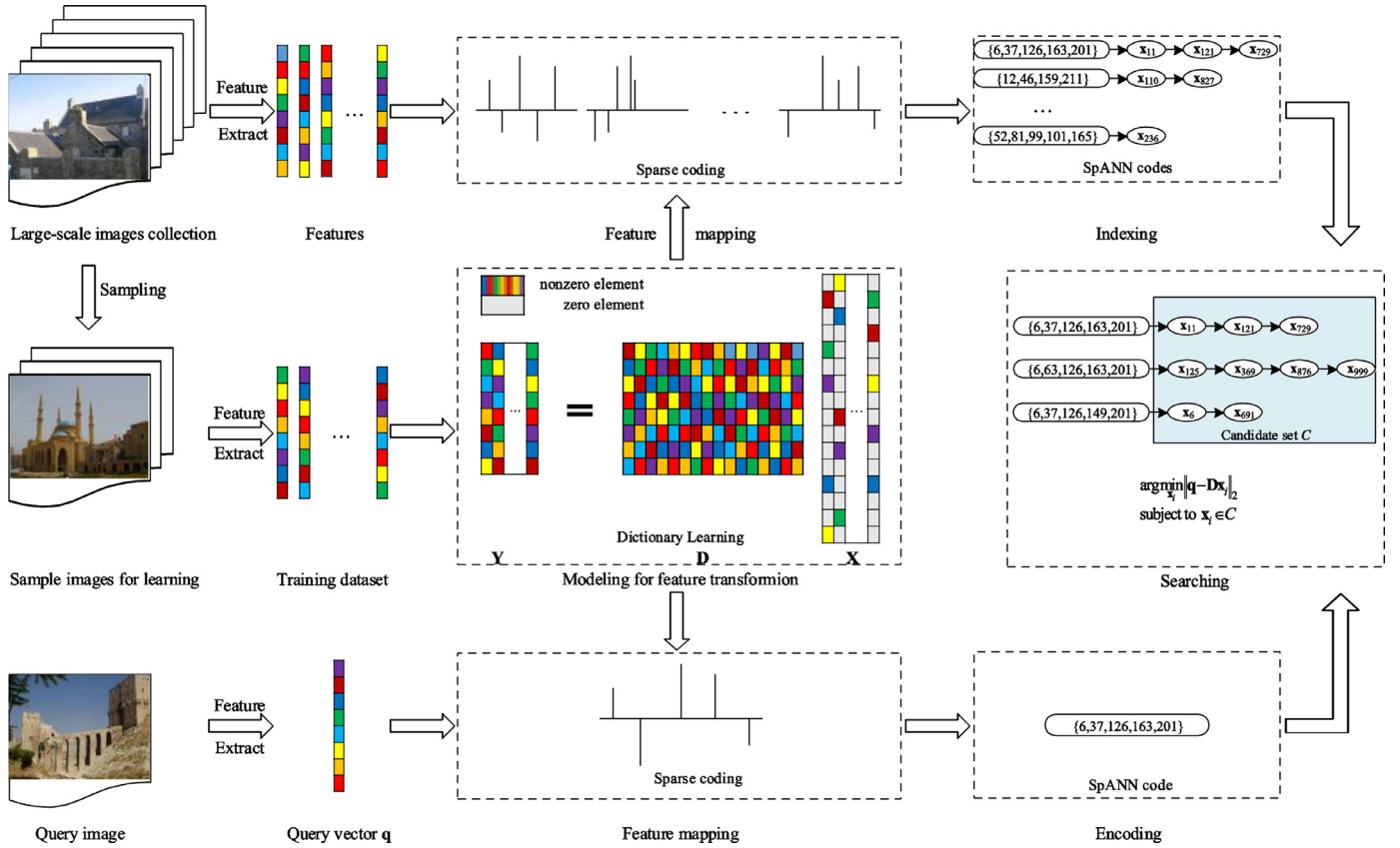


Fig. 1. Sketchy ANN retrieval route based on sparse representation and dictionary learning.

including training data and unseen data even other different data types, can be well described [10,11]. Thus, the dictionary model for feature transformation can effectively avoid the issue of over-fitting which is often neglected in existing retrieval methods. And the sparse codes of all samples are yielded by automatically discovering their inherent subspaces in the learned dictionary, giving rise to the very low-dimensional representation of data and the distinguishing chosen set of atoms [12,13]. Hence, we can gain low storage requirement and natural space partition for ANN.

The recent research, put forth by Cherian, develops an interesting way for ANN search based on SRDL [4]. Through this paper, we call it Cherian's method. The significant innovation of that method is to employ the indices of the active set of sparse coded data as compact ANN code called SpANN code. Note that the SpANN codes yielded on the dictionaries learned using the traditional dictionary learning algorithms, e.g. K-SVD [7], are often found to be sensitive to small differences in data points. That is, two similar points with small variation may be described by two completely different set of atoms, shown in [14,15]. This difficulty is theoretically analyzed and related to the coherence of the dictionary in Cherian's method: small coherence implies better performance on ANN search [4]. Hence, an incoherent dictionary learning algorithm is proposed for robust SpANN codes which are next exploited to index an inverted file table for fast retrieval.

Although Cherian's method has achieved high retrieval accuracy, it fails to consider the local structure of the data, which is proven to be very helpful for ANN search [16,17]. Towards this end, in this paper, we discuss two regularizations inspired by local structure of data for improvement of Cherian's method. One is graph Laplacian regularization, which aims to preserve the locality and the similarity of the original data into SpANN codes [18]. The other is non-negativity constraints, which aims to encourage the

similar samples to be described by a similar set of atoms, meanwhile cutting down on hash collision [19,20]. Both of them can enhance the robustness of SpANN code. The main contributions of this paper can be summarized as follows: 1) With the graph Laplacian regularization, we propose the graph regularized incoherent dictionary learning algorithm, referred to as GIDL. Via imposing non-negativity constraints, we put forth the non-negative sparse coding using incoherent dictionary, dubbed NIDL. Finally integrating GIDL and NIDL, we reap the graph regularized non-negative sparse coding using incoherent dictionary, called GNIDL. In experiments, we conduct the comparisons of the three approaches to survey the improvements resulted from the two regularizations. 2) Due to the fact that GNIDL can degrade into GIDL and NIDL, we directly tackle the optimization problem of GNIDL via iterating the two following steps: sparse coding and dictionary learning. For the former, we solve the graph regularized non-negative sparse coding by modifying the feature-sign search algorithm [21]; and we exploit the Lagrange dual problem of the resulting optimization problem in the latter to learn an incoherent dictionary [4]. And experimental results show that the iterative algorithm can reap a convergence solution in less than 100 iterations. In addition, we in all experiments also adopt the inverted file and retrieval settings that are used in Cherian's method. The experimental results manifest that the two rules can effectively improve the performance of Cherian's method and arrive at higher retrieval accuracy than the state-of-the-art methods, with low storage requirement and acceptable cost of query time.

The remainder of this paper is organized as follows: In Section 2, we review related works of ANN search, the theories of SRDL and Cherian's ANN approach. In Section 3, we introduce the motivations of the two regularizations and present the methods: GIDL, NIDL and GNIDL. Section 4 exhibits the optimization algorithms

Download English Version:

<https://daneshyari.com/en/article/4969526>

Download Persian Version:

<https://daneshyari.com/article/4969526>

[Daneshyari.com](https://daneshyari.com)