

Incremental Dictionary Learning for Unsupervised Domain Adaptation

Boyu Lu¹

bylu@umiacs.umd.edu

Rama Chellappa¹

rama@umiacs.umd.edu

Nasser M. Nasrabadi²

nasser.m.nasrabadi.civ@mail.mil

¹ Department of Electrical and Computer Engineering
University of Maryland
College Park, MD, USA

² U.S. Army Research Lab.
Adelphi, MD, USA

Domain adaptation (DA) methods attempt to solve the domain mismatch problem between source and target data. In this paper, we propose an incremental dictionary learning method where some target data called supportive samples are selected to assist adaptation. The idea is partially inspired by the bootstrapping-based methods [1, 3], which choose from the target domain some samples and add them into source domain for retraining the classifier. However, the suitable sample selection and stopping criteria for DA setting is a tricky problem. For the sample selection criteria, we choose supportive samples that are close to the source domain, so that they act as a bridge to connect the two domains and reduce the domain mismatch. More specifically, given the source dictionary D , we select the target samples that minimize the reconstruction error when represented by D . Then we augment the source domain by adding supportive samples and retrain the dictionary. For the stopping criteria, we guarantee that the domain mismatch decreases monotonically during adaptation. This is realized by checking whether adding new supportive samples will reduce the domain dissimilarity after each iteration. The proposed approach is shown in Fig. 1.

Supportive Samples Selection: We select the supportive samples using $W^{(k+1)}$ by solving the following optimization problem:

$$\begin{aligned} W_j^{(k+1)} &= \underset{W_j}{\operatorname{argmax}} \quad \operatorname{tr}(W_j P_j^{(k+1)}) \\ \text{s.t.} \quad W_j^{(k+1)} \cdot \sum_{l=1}^k W_j^{(l)} &= 0, \quad \|W_j^{(k+1)}\|_0 = Q, \quad j = 1, \dots, C \end{aligned} \quad (1)$$

where $W_j \in \mathbb{R}^{N_t \times N_t}$ are diagonal matrices with each element in the j^{th} column of W on the diagonal, e.g., $W_j = \operatorname{diag}\{w_{1j}, w_{2j}, \dots\}$ and similarly $P_j = \operatorname{diag}\{p_{1j}, p_{2j}, \dots\}$. $p_{ij} \in [0, 1]$ represents the probability that target sample x_i^t belongs to the class j and $w_{ij} \in \{0, 1\}$ indicates whether the target sample x_i^t is selected as supportive samples for class j . Q is the number of supportive samples for each class. This objective function (1) maximizes the confidence of the selected supportive samples. The first constraint requires that we keep adding new supportive samples to the source domain. The second constraint ensures that the number of supportive samples for each class is balanced.

Augmented Source Domain Update: After selecting the supportive samples, we augment source data by adding weighted supportive samples to existing source data:

$$X_j^{(k+1)} = [X_j^{(k)} | X_j^t W_j^{(k+1)} P_j^{(k+1)}] \quad j = 1, \dots, C \quad (2)$$

Since the labels of the supportive samples may be erroneous, each selected supportive sample is weighted by its confidence. The weights indicate the reliability of the labels of the supportive samples and highly confident supportive samples will contribute more to the model.

Dictionary Update: Dictionary is updated by solving the following dictionary learning problem:

$$D_j^{(k+1)} = \underset{D_j, Z_j}{\operatorname{argmin}} \|X_j^{(k+1)} - D_j \cdot Z_j\|_F^2 + \lambda \|Z_j\|_1 \quad j = 1, \dots, C. \quad (3)$$

We solve (3) using the online dictionary learning method [2]. The dictionary obtained in the previous iteration is used as the initial dictionary in the next iteration. In this way, the computational cost is relatively low.

Stopping criterion: One trivial stopping criterion is to stop when there is no new supportive samples for one of the classes. But our goal is to guarantee that the adaptation monotonically reduce the domain divergence. So we design a domain similarity measure and perform adaptation only when the domain similarity increases after each iteration.

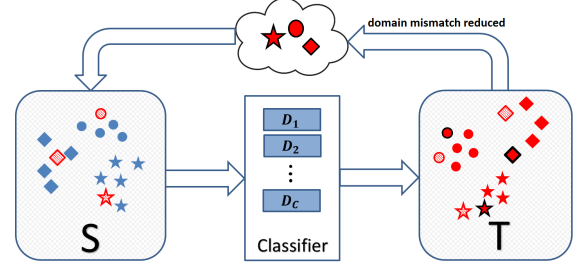


Figure 1: Scheme of the proposed method. The original source data is colored in *blue* and the target data is colored in *red*. Different shapes represent different classes. The red samples with shadow indicate the previously selected supportive samples that have been added to the source domain. The red samples with black border represent the supportive samples selected in the current iteration.

In order to quantify domain similarity, we introduce a simple domain similarity measure for X^s and X^t :

$$\rho(X^s, X^t) = \sqrt{\frac{1}{N_s N_t} \sum_i \sum_j (x_i^s x_j^t)^2} = \sqrt{\frac{\operatorname{tr}((X^s)^T X^t (X^t)^T X^s)}{N_s N_t}}. \quad (4)$$

Since the classification accuracy on supportive samples is good, the main reason for the performance to drop in the target domain is that the source classifier behaves poorly on the non-supportive samples. It indicates that domain mismatch mainly lies between the source samples and the non-supportive samples. If the distance between supportive samples and non-supportive samples is smaller than the distance between the source domain and the non-supportive samples, selecting supportive samples can help reduce the domain mismatch.

Theorem 1. We divide the target samples into two part, supportive samples X_f and non-supportive samples X_n with N_f and N_n samples, respectively. With the definition of ρ above, and if $\rho(X_f, X_n) > \rho(X^s, X_n)$, then the domain similarity (or mismatch) will increase (or decrease) when we add some supportive samples to the source domain:

$$\rho(X_{new}^s, X^t) > \rho(X_{old}^s, X^t) \quad (5)$$

where $X_{old}^s = X^s$ and $X_{new}^s = [X^s | X_f]$.

We evaluate the proposed method for object classification and face recognition and compare with several state-of-the-art unsupervised DA methods. Experimental results show that our method outperforms other approaches significantly in most cases.

- [1] Lorenzo Bruzzone and Mattia Marconcini. Domain adaptation problems: A DASVM classification technique and a circular validation strategy. *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, 32(5):770–787, 2010.
- [2] Julien Mairal, Francis Bach, Jean Ponce, and Guillermo Sapiro. Online dictionary learning for sparse coding. In *Proceedings of the 26th Annual International Conference on Machine Learning*, pages 689–696. ACM, 2009.
- [3] Tatiana Tommasi and Barbara Caputo. Frustratingly easy nbnn domain adaptation. In *Computer Vision (ICCV), 2013 IEEE International Conference on*, pages 897–904. IEEE, 2013.