

Deep Scene Image Classification with the MFAFVNet

Yunsheng Li Mandar Dixit Nuno Vasconcelos
University of California, San Diego
La Jolla, CA 92093

yl554@ucsd.edu mdixit@ucsd.edu nvasconcelos@ucsd.edu

Abstract

The problem of transferring a deep convolutional network trained for object recognition to the task of scene image classification is considered. An embedded implementation of the recently proposed mixture of factor analyzers Fisher vector (MFA-FV) is proposed. This enables the design of a network architecture, the MFAFVNet, that can be trained in an end to end manner. The new architecture involves the design of a MFA-FV layer that implements a statistically correct version of the MFA-FV, through a combination of network computations and regularization. When compared to previous neural implementations of Fisher vectors, the MFAFVNet relies on a more powerful statistical model and a more accurate implementation. When compared to previous non-embedded models, the MFAFVNet relies on a state of the art model, which is now embedded into a CNN. This enables end to end training. Experiments show that the MFAFVNet has state of the art performance on scene classification.

1. Introduction

Scene image classification is an important problems for applications of computer vision such as robotics, image search, geo-localization, etc. It is also a challenging problem, e.g. object recognition methods do not necessarily work for scenes. This is because scenes include both a holistic component, the *gist* of the scene, and an object-based component. Furthermore, the object vocabulary is usually open-ended and it does not suffice to recognize objects, as most scenes are collections of objects in characteristic spatial layouts. There is also a need to model relationships between objects. Historically, this motivated different approaches to scene classification, including holistic gist descriptors [20] and descriptors based on local features.

Localization approaches included descriptors of contextual relationships between objects [14] or semantics properties such as the objects in the scene [15] or scene attributes [25]. Many of these approaches were based on the

bag-of-features (BoF) representation, using local features such as SIFT or HOG [19, 3], combined through a pooling operator. Pooling has a critical role in scene classification task. For two reasons. First, there is a need to integrate object information across the scene. Second, this integration must be invariant to the locations of individual objects, which can change drastically within the same scene class. Eventually, sophisticated pooling strategies, such as the vector of locally aggregated descriptors (VLAD) [12] or the Fisher vector (FV) [23] emerged as the dominant pooling mechanisms for scene classification.

In recent years, CNNs have become the feature extractors of choice for scene classification. In fact, the implementation of a holistic scene classifier with a CNN is almost trivial. It suffices to train the CNN on whole scene images. The main challenges are the assembly of a large dataset of such images and the standard difficulties of training a deep network. These problems have been addressed in [29], through the assembly of the Places dataset and its use to train a deep scene classifier. The use of deep object-based representations for scene classification has, however, proven more challenging. This is, in great part, because object-based scene classification is a difficult transfer learning problem. While the ease with which CNNs can be transferred across datasets, by simple finetuning, is one of their greatest assets for vision, this procedure has its limitations. The transfer from a localized representation (needed to recognize objects) to a holistic task (classification of whole scenes) cannot, in general, be solved by simple finetuning.

The previous success of pooling on this type of transfer created a resurgence of interest in the topic within the deep learning realm. Several authors proposed Fisher vectors or similar pooling mechanisms for features extracted by object recognition CNNs. Early methods adopted a BoF-like approach, based on the extraction of features from intermediate CNN layers, which were then fed to dictionary learning methods such as clustering [11] or sparse coding [17] and finally used to implement descriptors such as the VLAD [11] or Fisher vector [17]. Later, [5] proposed the semantic Fisher Vector, which extracts features from the

softmax layer at the top of the CNN, converting features from probability space to the natural parameter space.

All these methods suffer from two drawbacks. First, the Fisher vector structure is not easy to integrate in a CNN. This is because the Fisher vector is defined with respect to the probability distribution of the CNN features, usually estimated with a mixture learned by maximum likelihood. The Fisher vector is then a complex expression of the mixture parameters, which changes when these change. In result, the Fisher vector cannot be learned by simply back-propagating the output of the scene classifier. All methods above avoid this difficulty by using the CNN to extract features and learning the Fisher vector independently. This, however, prevents end to end training and, consequently, the finetuning of the object network to the scene classification task. In result, the transfer between the two tasks relies solely on the Fisher vector, which is sub-optimal.

The second problem is that the Fisher vector is usually learned with respect to a Gaussian mixture model (GMM). Since CNN features are high dimensional, it is impractical to rely on Gaussians of full covariance. Instead, the mixture components are chosen to have diagonal covariance. This creates problems when the feature manifold is non-linear. In this case, a large number of diagonal GMM components are required and the Fisher vector is very high dimensional. This is indeed the norm for computer vision applications, where Fisher vectors usually have several thousand dimensions. While the CNN is trained to produce linearly separable responses to the different classes, there is no guarantee CNN feature distributions are prone to modeling with the diagonal GMM. On the contrary, given the highly non-linear feature transformation implemented by a deep CNN, this is unlikely. Hence, a Fisher vector based on the diagonal-GMM is likely to be very high dimensional and potentially sub-optimal for scene classification.

Recently, some works have attempted to solve these problems. One possibility is to bypass the Fisher vector pooling. For example, [7] proposed a compact bilinear pooling (CBP) mechanism that enables end-to-end training by simple backpropagation. While this was shown applicable to scene classification, the performance of CBP is inferior to those of previous approaches, such as the semantic Fisher vector of [5] or the sparse coding methods of [18], for equivalent object CNNs. Another possibility is to embed the Fisher vector in the CNN architecture, by deriving a neural network implementation of its equations. [1] proposed the NetVLAD, an embedded implementation of the VLAD descriptor, and [26] proposed the Deep FisherNet, an embedded implementation of the GMM Fisher vector. However, to avoid the difficulties of the complete Fisher vector, these methods make approximations, such as disregarding covariance structure (VLAD) or using crude approximations of posterior mixture probabilities (Deep FisherNet). These ap-

proximations can be quite sub-optimal.

An alternative strategy, proposed by [6], is to use a better model of CNN feature statistics than the diagonal GMM. Instead, this work proposed a Fisher vector based on the mixture of factor analyzers (MFA). This has the advantage of accounting for the covariance information of each mixture component, which is modeled through factor analysis. Under the MFA model, good results can be achieved with mixtures of few components and Fisher vectors of reduced dimension. However, while the MFA-Fisher vector (MFA-FV) currently holds state of art results for scene classification, it is not an integrated model, i.e. it is learned independently of the network. This raises all the reservations discussed above.

In this work, we derive an embedded implementation of the MFA-FV, and use it to design a network architecture, the MFAFVNet, which can be trained in an end to end manner. This involves the derivation of a MFA-FV layer that implements a statistically correct version of the MFA-FV, through a combination of network computations and regularization. The computations replicate those of the MFA-FV, regularization guarantees that all parameters have a statistically valid interpretation. When compared to previous embedded implementations, the MFAFVNet relies on a more powerful statistical model, which accounts for covariance information, and a more accurate implementation. This results in significant performance gains for scene classification. When compared to previous non-embedded models, the MFAFVNet relies on a state of the art model, which is embedded. This enables end to end training and better scene classification performance. Extensive experiments on the MIT Indoor and SUN datasets show that the MFAFVNet achieves state of the art performance for scene classification.

The mixture of factor analyzers (MFA) is a mixture model whose components follow the factor analysis model. A MFA of C components is defined by the distributions

$$p(c) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

2. The MFA Fisher Vector

In this section, we review the main ideas behind the MFA-FV.

2.1. Mixture of Factor Analyzers

The factor analysis (FA) model is a probabilistic extension of principal component analysis (PCA) [23]. Given a vector $x \in \mathbb{R}^D$ of D observations, it explains its covariance structure by assuming that the variability of the observations can be explained by a small number $d < D$ of hidden or latent factors, usually represented as a factor vector $z \in \mathbb{R}^d$. Observations x are assumed to be sampled according to the model

$$x - \mu = \Lambda z + \epsilon, \quad (1)$$

where μ is the mean value of x , $\Lambda \in \mathbb{R}^{D \times d}$ is a factor loading matrix, and ϵ is a noise term. Factors z and noise ϵ are independent of each other and Gaussian, and the

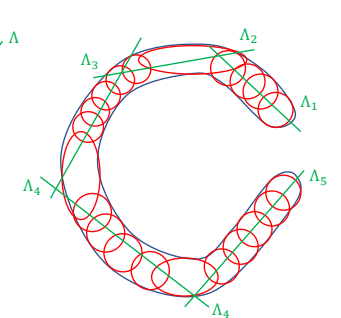


Figure 1. Probability distributions defined on linear subspaces (left) or non-linear manifolds (right) can require many mixture components to approximate, when covariances are diagonal (shown in red). However, they can frequently be approximated by a few MFA components (in green).

noise variables are assumed uncorrelated, i.e. distributed as $\mathcal{N}(\epsilon; 0, I)$. The factors can be dependent, i.e. they are distributed as $\mathcal{N}(z; 0, \psi)$, but the matrix ψ is sometimes assumed diagonal. It follows from the linearity of (1) that x is Gaussian with covariance

$$\Sigma = \text{cov}(x - \mu) = \text{cov}(\Lambda z + \epsilon) = \Lambda \Lambda^T + \psi. \quad (2)$$

Hence, the factor loading matrix Λ has a role similar to the principal component matrix of PCA.

The mixture of factor analyzers (MFA) is a mixture model whose components follow the factor analysis model. A MFA of C components is defined by the distributions

$$p(c) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

2. The MFA Fisher Vector

In this section, we review the main ideas behind the MFA-FV.

2.1. Mixture of Factor Analyzers

The factor analysis (FA) model is a probabilistic extension of principal component analysis (PCA) [23]. Given a vector $x \in \mathbb{R}^D$ of D observations, it explains its covariance structure by assuming that the variability of the observations can be explained by a small number $d < D$ of hidden or latent factors, usually represented as a factor vector $z \in \mathbb{R}^d$. Observations x are assumed to be sampled according to the model

$$x - \mu = \Lambda z + \epsilon, \quad (1)$$

where μ is the mean value of x , $\Lambda \in \mathbb{R}^{D \times d}$ is a factor loading matrix, and ϵ is a noise term. Factors z and noise ϵ are independent of each other and Gaussian, and the

Fisher vector reduces to the score $\mathcal{G}(\theta)$. As is common in computer vision, we adopt this practice in this work.

The Fisher vector is commonly used with the standard Gaussian mixture model, defined by

$$p(c) = \pi_c \quad (6)$$

$$p(x|c) = \mathcal{N}(x; \mu_c, \Sigma_c). \quad (7)$$

However, because vision data tends to be high-dimensional, it is usually difficult to use full covariance Gaussians in this model, and the covariances Σ_c are assumed as diagonal. This removes a lot of the expressiveness of the mixture model. In general, many components are needed to achieve a good approximation of the distribution $p(x)$. As illustrated in Figure 1, this is particularly true when the data lives on correlated low-dimensional subspaces or non-linear manifolds that can be approximated by a set of low-dimensional subspaces. The MFA is a substantially better model for this type of data since, in this case, only a few mixture components and a small number d of hidden factors are required to estimate the covariance structure of (2). This is illustrated in Figure 1 as well. This observation, motivated the introduction of the MFA-Fisher vector (MFA-FV) in [6], which was shown to have the form

$$\mathcal{G}_{\mu_c}(x) = \sum_i p(c|x_i; \theta) \psi^{-1/2} (I - \Lambda_c \Gamma_c) (x_i - \mu_c) \quad (8)$$

$$\mathcal{G}_{\Lambda_c}(x) = \sum_i p(c|x_i; \theta) \psi^{-1/2} (\Lambda_c \Gamma_c - I) \quad (9)$$

$$\Gamma_c = \Lambda_c^T \Sigma_c^{-1} \Lambda_c \quad (10)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$

$$p(z|c) = \mathcal{N}(z; 0, I) \quad (4)$$

$$p(x|z, c) = \mathcal{N}(x; \Lambda_c z + \mu_c, \Psi_c) \quad (5)$$

$$p(c|x) = \pi_c \quad (3)$$