

Online-Batch Strongly Convex Multi Kernel Learning

Francesco Orabona*
Università degli Studi di Milano
orabona@dsi.unimi.it

Luo Jie
Idiap Research Institute
EPF Lausanne
jluo@idiap.ch

Barbara Caputo
Idiap Research Institute
bcaputo@idiap.ch

Abstract

Several object categorization algorithms use kernel methods over multiple cues, as they offer a principled approach to combine multiple cues, and to obtain state-of-the-art performance. A general drawback of these strategies is the high computational cost during training, that prevents their application to large-scale problems. They also do not provide theoretical guarantees on their convergence rate.

Here we present a Multiclass Multi Kernel Learning (MKL) algorithm that obtains state-of-the-art performance in a considerably lower training time. We generalize the standard MKL formulation to introduce a parameter that allows us to decide the level of sparsity of the solution. Thanks to this new setting, we can directly solve the problem in the primal formulation. We prove theoretically and experimentally that 1) our algorithm has a faster convergence rate as the number of kernels grow; 2) the training complexity is linear in the number of training examples; 3) very few iterations are enough to reach good solutions. Experiments on three standard benchmark databases support our claims.

1. Introduction

Categorization is one of the most challenging problems in computer vision today. Object categories present a wide visual variability within each class. This, coupled with robustness issues (e.g. changes in illumination, occlusion, clutter), makes it unclear how to build general models suitable for all categories. Because of this, a dominant approach is to learn instead what distinguishes them, by using highly discriminative and robust features combined with machine learning techniques [8, 9, 13, 16, 24, 25]. In particular this has been recently translated into Support Vector Machine (SVM) based classifiers combined with kernels over multiple cues [2, 8, 9, 13, 24, 25]. Results obtained by these methods on various benchmark databases represent the current state-of-the-art in object categorization. Among them, Multi Kernel Learning (MKL) approaches have attracted

considerable attention [8, 13, 24]. However most emphasis has been put so far on their accuracy, and recent findings seem to indicate that current MKL algorithms do not improve much over the naive baseline of averaging all the kernels [9].

Almost every interesting categorization problems have more than two classes, and most of the MKL algorithms [12, 18, 22] solves the multiclass problem by decomposing it into multiple independent binary classification tasks (except [27]). However, recent evidence [9] seems to suggest that a principled multiclass formulation (such as those in [9, 23, 27]) achieves better performance, at least on sparse problems using l_1 regularization. Moreover, to our knowledge, none of the MKL algorithms [12, 18, 22] provides theoretical guarantees on their convergence rate. In practice, the learning process is usually stopped early, before reaching the optimal solution, based on the common assumption that it is enough to have an approximate solution of the optimization function. Considering the fact that current MKL algorithms are solved based on their dual representation, this might mean being stopped far from the optimal solution [10]. Last but not least, scalability is also very important for many real world applications.

The contribution of this paper is a Multiclass MKL algorithm that has a guaranteed and fast convergence rate to the optimal solution. We also generalize the MKL learning problem, adding a parameter to tune the level of sparsity in the kernel domain. We show experimentally that aiming at sparsity, as in the original MKL formulation, is not always the optimal strategy. Our algorithm has a training time that depends linearly on the number of training examples, with a convergence rate sub-linear in the number of kernels used. At the same time, it achieves state-of-the-art performance on standard benchmark databases. The algorithm is based on a stochastic sub-gradient descent algorithm in the primal objective formulation. Minimizing the primal objective function directly results in a convergence rate that is faster and provable, rather than optimizing the dual objective. We show that by optimizing the primal objective function directly, we are able to solve the multiclass formulation effi-

*Work done while at Idiap Research Institute, Martigny, Switzerland

ciently, with a running time which is linear to the number of classes. We can stop the algorithm after few iterations, while still retaining a performance close to the optimal one. We call this algorithm OBSCURE, Online-Batch Strongly Convex mUlti keRnel lEarning.

1.1. Multiple Cues and Kernels

Consider the task of image classification with M classes, F different cues and N training instances $\{\mathbf{x}_i\}_{i=1}^N$ drawn from an unknown fixed probability distribution. We want to learn a score function $s(\cdot, \cdot)$ that best predicts the class $\hat{y}$ for any future sample $\mathbf{x}$ drawn from the same distribution, where the predicted class is the one with the highest score

$$\hat{y}(\mathbf{x}) = \operatorname{argmax}_{y \in \mathbb{Y}} s(\mathbf{x}, y). \quad (1)$$

This score function should be learned using all the F different cues, to gain robustness and performance.

Some of the methods addressing this task are based on a two-layers structure [9, 16]. A classifier is trained for each cue and then their outputs are combined by another classifier. Even if this strategy has recently received attention in the computer vision community, this kind of approach is the oldest and dates back to the seminal work of Wolpert [26]. They use Cross-Validation (CV) methods to create the training set for the second layer [9, 26]. Hence they have a run-time of about $K+1$ times the training of a single classifier, such as support vector machine (SVM), where K is the number of folds of the CV. This method is currently considered the state-of-art method for image classification tasks [9].

Another interesting strategy uses a one-layer architecture, such as the MKL [14, 18, 22, 24, 27]. Using the theory of *kernels*, one solves a joint optimization problem while also learning the optimal weights for combining the kernels, with each cue corresponding to a kernel. The optimization problem is similar in all these approaches. This approach is theoretically founded, plus it consists of a unique optimization problem. However solving it is more complex than training, *e.g.*, a single SVM classifier. Another issue is that current MKL approaches do not scale well to the number of training examples and number of classes. For example, the SILP algorithm [22, 27] depends polynomially on the number of training examples and number of classes with an exponent of ~ 2.4 and ~ 1.7 respectively. For the other algorithms these dependencies are not clear.

From a theoretical point of view, if we consider a two-layers architecture with the first layer composed by kernel classifiers, and a linear classifier in the second stage, the two approaches are very similar. In both cases the final prediction function is written as

$$\hat{y}(\mathbf{x}) = \operatorname{argmax}_{y \in \mathbb{Y}} \sum_{j=1}^F \beta_y^j s^j(\mathbf{x}, y), \quad (2)$$

where β_y^j are the weights learned by the one-layer or two-layers framework, and s^j is the score function for each kernel. Therefore the two formulations are essentially equivalent, with differences given only by the specific training procedures used. In both methods a regularizer that favors the selection of only a subset of the kernels is used [1, 9, 22, 24].

The main contribution of this paper is showing that the one-layer formulation, beside being more principled, can also achieve a comparable performance and a considerably lower training time than state-of-the-art two-layers architectures. We propose a p -norm version of the standard MKL algorithm, and we minimize it with a two stages algorithm. The first one is an online initialization procedure that determines quickly the region of the space where the optimal solution lives. The second stage refines the solution found by the first stage. Differently from the other methods, our algorithm solves the optimization problem directly in the primal formulation, in both stages. Using recent approaches in optimization theory, the algorithm takes advantage of the abundance of information to reduce the training time [21]. In fact, we show that the presence of a large number of kernels helps the optimization process instead of hindering it, obtaining, theoretically and practically, a faster convergence rate with more kernels.

The rest of the paper presents the theory and the experimental results supporting our claims. Section 2 revises the basic definitions of MKL and generalizes it to p -norm formulation. Section 3 presents the theory and algorithm of OBSCURE, while Section 4 reports experiments on categorization tasks.

2. p -norm Multi Kernel Learning

In this section we first introduce formally the MKL framework and its notation, then its p -norm generalization.

2.1. Definitions

Notations. Let $\{\mathbf{x}_i, y_i\}_{i=1}^N$, with $N \in \mathbb{N}$, $\mathbf{x}_i \in \mathbb{X}$ and $y_i \in \mathbb{Y} = \{1, \dots, M\}$, $M > 2$, be the training set. We indicate matrix and vectors with bold letters. A bar, *e.g.* $\bar{\mathbf{w}}$, denotes the vector formed by the concatenation of the F vectors $\mathbf{w}^j$, hence $\bar{\mathbf{w}} = [\mathbf{w}^1, \mathbf{w}^2, \dots, \mathbf{w}^F]$.

Multi-class Classifier. A common approach to multiclass classification is to use joint feature maps $\phi(\mathbf{x}, y)$ on data $\mathbb{X}$ and labels $\mathbb{Y}$ [23]. The function s^j will be defined as

$$s^j(\mathbf{x}, y) = \mathbf{w}^j \cdot \phi^j(\mathbf{x}, y), \quad (3)$$

where $\mathbf{w}^j$ is a hyperplane¹. The functions $\phi^j(\mathbf{x}, y)$ map the samples into a high, possibly infinite, dimensional space. With multiple cues, we will have F functions

¹For simplicity we will not use the bias, it can be easily added by modifying the kernel definition.

$\phi^j(\cdot, \cdot), j = 1, \dots, F$. This will also define F kernels $K^j((\mathbf{x}, y), (\mathbf{x}', y'))$ as $\phi^j(\mathbf{x}, y) \cdot \phi^j(\mathbf{x}', y')$. This definition includes the case of training M different hyperplanes, one for each class. Indeed $\phi^j(\mathbf{x}, y)$ can be defined as

$$\phi^j(\mathbf{x}, y) = [\mathbf{0}, \dots, \mathbf{0}, \underbrace{\phi'^j(\mathbf{x})}_y, \mathbf{0}, \dots, \mathbf{0}], \quad (4)$$

where $\phi'^j(\cdot)$ is a transformation that depends only on data. Similarly $\mathbf{w}$ will be composed by M blocks, $[\mathbf{w}^1, \dots, \mathbf{w}^M]$. Hence, by construction, $\mathbf{w} \cdot \phi^j(\mathbf{x}, r) = \mathbf{w}^r \cdot \phi'^j(\mathbf{x})$. According to the defined notation, $\bar{\phi}(\mathbf{x}, y) = [\phi^1(\mathbf{x}, y), \dots, \phi^F(\mathbf{x}, y)]$.

Loss Function. We define a multi-class loss function [23]

$$\ell(\mathbf{w}, \mathbf{x}, y) = \max_{y' \neq y} |1 - \bar{\mathbf{w}} \cdot (\bar{\phi}(\mathbf{x}, y) - \bar{\phi}(\mathbf{x}, y'))|_+, \quad (5)$$

where $|t|_+$ is $\max(t, 0)$. This loss function is convex and it upper bounds the multi-class misclassification loss.

Norms and dual norms. A generic norm of a vector $\mathbf{w}$ is indicated by $\|\mathbf{w}\|$, its *dual norm* is indicated by $\|\mathbf{w}\|_*$. For $\mathbf{w} \in \mathbb{R}^d$ and $p \geq 1$, we denote by $\|\mathbf{w}\|_p$ the p -norm of $\mathbf{w}$, i.e., $\|\mathbf{w}\|_p = (\sum_{i=1}^d |w_i|^p)^{1/p}$. The dual norm of $\|\cdot\|_p$ is $\|\cdot\|_q$, where p and q satisfy $1/p + 1/q = 1$. In the following p and q will always satisfy this relation.

Group Norm. It is possible to define a $(2, p)$ group norm $\|\bar{\mathbf{w}}\|_{2,p}$ on $\bar{\mathbf{w}}$ as

$$\|\bar{\mathbf{w}}\|_{2,p} := \left\| [\|\mathbf{w}^1\|_2, \|\mathbf{w}^2\|_2, \dots, \|\mathbf{w}^F\|_2] \right\|_p, \quad (6)$$

that is the p -norm of the vector of F elements, formed by 2-norms of the vectors $\mathbf{w}^j$. The dual norm of $\|\cdot\|_{2,p}$ is $\|\cdot\|_{2,q}$ [11].

2.2. Multi Kernel Learning

The MKL optimization problem was first proposed in [1] and extended to multiclass in [27]. It can be written as

$$\begin{aligned} \min_{\mathbf{w}_j} \quad & \frac{\lambda}{2} \left(\sum_{j=1}^F \|\mathbf{w}^j\|_2 \right)^2 + \frac{1}{N} \sum_{i=1}^N \xi_i \\ \text{s.t.} \quad & \bar{\mathbf{w}} \cdot (\bar{\phi}(\mathbf{x}_i, y_i) - \bar{\phi}(\mathbf{x}_i, y)) \geq 1 - \xi_i, \forall i, y \neq y_i. \end{aligned} \quad (7)$$

This same formulation is used in [1, 22], while in [18] the proposed formulation is slightly different, although it is proved to be equivalent. Note that we weight the regularization term by λ and divide the loss term by N , instead of the more common formulation with only the loss term weighted by a parameter C . Our choice greatly simplifies the math of our algorithm. The two formulations are fully equivalent when setting $\lambda = \frac{1}{CN}$.

We will now generalize this formulation to group-norms. Using the notation defined above, we can rewrite (7) as

$$\min_{\bar{\mathbf{w}}} \quad \frac{\lambda}{2} \|\bar{\mathbf{w}}\|_{2,1}^2 + \frac{1}{N} \sum_{i=1}^N \ell(\bar{\mathbf{w}}, \mathbf{x}_i, y_i), \quad (8)$$

where $\bar{\mathbf{w}} = [\mathbf{w}^1, \mathbf{w}^2, \dots, \mathbf{w}^F]$. The $(2, 1)$ group norm is used to induce sparsity in the domain of the kernels. This means that the solution of the optimization problem will select a subset of the F kernels. However, even if sparsity can be desirable for specific applications, it could bring to a decrease in performance. Moreover the problem in (8) is not strongly convex [11], so its optimization algorithm is rather complex and its rate of convergence is usually slow [1, 22].

We propose to generalize the optimization problem, using a generic group norm

$$\min_{\bar{\mathbf{w}}} \quad \frac{\lambda}{2} \|\bar{\mathbf{w}}\|_{2,p}^2 + \frac{1}{N} \sum_{i=1}^N \ell(\bar{\mathbf{w}}, \mathbf{x}_i, y_i), \quad (9)$$

where $1 < p \leq 2$. We define $f(\bar{\mathbf{w}}) = \frac{\lambda}{2} \|\bar{\mathbf{w}}\|_{2,p}^2 + \frac{1}{N} \sum_{i=1}^N \ell(\bar{\mathbf{w}}, \mathbf{x}_i, y_i)$ and $\bar{\mathbf{w}}^*$ equals to the optimal solution of (9), $\bar{\mathbf{w}}^* = \arg \min_{\bar{\mathbf{w}}} f(\bar{\mathbf{w}})$. The added parameter p will allow us to decide the level of sparsity of the solution. In fact it is known that the 1-norm favors sparsity, and here the 1-norm favors a solution in which only few hyperplanes have a norm different from zero. Moreover this new formulation has the advantage of being λ/q -strongly convex [11]. Strong convexity is a key property to design fast batch and online algorithms: the more a problem is strongly convex the easier it is to optimize it [11, 19]. Many optimization problems are strongly convex, as the SVM objective function. When p tends to 1, the solution gets close to the sparse solution obtained solving the problem in (7), but the strong convexity decreases. When p equals to 2, it is equivalent to using a single kernel equal to the sum of all the kernels. Recently a different p -norm MKL problem has been also proposed in [12], that allows non-sparse solutions. However, their algorithm did not take advantage of nice properties of the strong convexity for the optimization process. In the next section, we will show how to use the strong convexity to design a fast algorithm to solve (9).

3. The OBSCURE Algorithm

Our basic optimization tool is the framework developed in [19, 20]. It is a general framework to design and analyze stochastic sub-gradient descent algorithms for any strongly convex function. At each step the algorithm takes a random sample of the training set and calculates a sub-gradient of the objective function evaluated on the sample. Then it performs a sub-gradient descent step with decreasing learning rate, followed by a projection of the solution inside the space where the solution lives. The algorithm Pegasos,

Algorithm 1 OBSCURE stage 1 (online)

```

1: Input:  $q, \eta$ 
2: Initialize:  $\bar{\theta}_1 = \mathbf{0}, \bar{w}_1 = \mathbf{0}$ 
3: for  $t = 1, 2, \dots, T$  do
4:   Sample at random  $(x_t, y_t)$ 
5:    $\hat{y}_t = \operatorname{argmax}_{y \neq y_t} \bar{w}_t \cdot \bar{\phi}(x_t, y)$ 
6:    $\bar{z}_t = \bar{\phi}(x_t, y_t) - \bar{\phi}(x_t, \hat{y}_t)$ 
7:   if  $\ell(\bar{w}_t, x_t, y_t) > 0$  then  $\bar{\theta}_{t+1} = \bar{\theta}_t + \eta \bar{z}_t$ 
8:   else  $\bar{\theta}_{t+1} = \bar{\theta}_t$ 
9:    $w_{t+1}^j = \frac{1}{q} \left( \frac{\|\theta_{t+1}^j\|_2}{\|\bar{\theta}_{t+1}\|_{2,q}} \right)^{q-2} \theta_{t+1}^j, \forall j = 1, \dots, F$ 
10: end for
11: return  $\bar{\theta}_{T+1}, \bar{w}_{T+1}$ 
12: return  $R = \sqrt{\|\bar{w}_{T+1}\|_{2,p}^2 + \frac{2}{\lambda N} \sum_{i=1}^N \ell(\bar{w}_{T+1}, x_i, y_i)}$ 

```

Algorithm 2 OBSCURE stage 2 (batch)

```

1: Input:  $q, \bar{\theta}_1, \bar{w}_1, R, \lambda$ 
2: Initialize:  $s_0 = 0$ 
3: for  $t = 1, 2, \dots, T$  do
4:   Sample at random  $(x_t, y_t)$ 
5:    $\hat{y}_t = \operatorname{argmax}_{y \neq y_t} \bar{w}_t \cdot \bar{\phi}(x_t, y)$ 
6:   if  $\ell(\bar{w}_t, x_t, y_t) > 0$  then  $\bar{z}_t = \bar{\phi}(x_t, y_t) - \bar{\phi}(x_t, \hat{y}_t)$ 
7:   else  $\bar{z}_t = \mathbf{0}$ 
8:    $d_t = \lambda t + s_{t-1}$ 
9:    $s_t = s_{t-1} + 0.5 \left( \sqrt{d_t^2 + q \frac{(\frac{\lambda}{q} \|\bar{\theta}_t\|_{2,q} + \|\bar{z}_t\|_{2,q})^2}{R^2}} - d_t \right)$ 
10:    $\eta_t = \frac{q}{\lambda t + s_t}$ 
11:    $\bar{\theta}_{t+\frac{1}{2}} = (1 - \frac{\lambda \eta_t}{q}) \bar{\theta}_t + \eta_t \bar{z}_t$ 
12:    $\bar{\theta}_{t+1} = \min \left( 1, qR / \|\bar{\theta}_{t+\frac{1}{2}}\|_{2,q} \right) \bar{\theta}_{t+\frac{1}{2}}$ 
13:    $w_{t+1}^j = \frac{1}{q} \left( \frac{\|\theta_{t+1}^j\|_2}{\|\bar{\theta}_{t+1}\|_{2,q}} \right)^{q-2} \theta_{t+1}^j, \forall j = 1, \dots, F$ 
14: end for

```

based on this framework, is the current state-of-art solver for linear SVM [20, 21].

Given that the $(2, p)$ group norm is strongly convex, we could use this framework to design an efficient MKL algorithm. It would inherit all the properties of Pegasos [20, 21]. In particular the convergence rate, and hence the training time, would be proportional to $\frac{q}{\lambda}$. Although in general this convergence rate can be quite good, it becomes slow when λ is small and/or q is big. Moreover it is common knowledge that in many real-world problems, particularly in visual learning tasks, the best setting for λ is very small, or equivalently C is big (the order of $10^2 - 10^3$). Notice that this is a general problem. The same problem also exists in the other SVM optimization algorithms such as SMO and similar approaches [10], as their training time also depends on the value of the parameter C .

Do *et al.* [6] proposed a variation of the Pegasos algorithm called proximal projected sub-gradient descent. This

formulation has a better convergence rate for small values of λ , while retaining the fast convergence rate for big values of λ . A drawback is that the algorithm needs to know in advance an upper bound on the norm of the optimal solution. In [6] the authors proposed an algorithm that estimates this bound while training, but it gives a speed-up only when the norm of the optimal solution $\bar{w}^*$ is small. This is not the case in most of the MKL problems for categorization tasks.

Our OBSCURE algorithm takes the best of the two solutions. We first extend the framework of [6] to the generic non-Euclidean norms. Then we solve the problem of the upper bound of the norm of the optimal solution using an new online algorithm. This takes advantage of the characteristic of the MKL task and quickly converges to a solution close to the optimal one. Hence OBSCURE is composed by two steps: the first step is a fast online algorithm (Algorithm 1), used to quickly estimate the region of the space where the optimal solution lives. The second step (Algorithm 2) starts from the approximate solution found by the first stage, and exploiting the information on the estimated region, it uses a stochastic proximal projected sub-gradient descent algorithm.

The following theorem² gives a theoretical guarantee on the convergence rate of OBSCURE to the solution of (9).

Theorem 1. *Suppose that $\|\phi^j(x_t, y_t)\|_2 \leq 1, \forall j = 1, \dots, F, t = 1, \dots, N$. Let $1 < p \leq 2, \delta \in (0, 1), R$ the value returned by the first stage, and $c = \sqrt{2}F^{1/q} + \lambda R$. Then with probability at least $1 - \delta$ over the choices of the random samples we have that after T iterations of the 2nd stage of the OBSCURE algorithm, the difference between $f(\bar{w}_T)$ and the optimal solution of (9), $f(\bar{w}^*)$, is less than*

$$\frac{c\sqrt{q}\sqrt{1+\log T}}{\delta} \min \left(\frac{c\sqrt{q}\sqrt{1+\log T}}{\lambda T}, \frac{4R}{\sqrt{T}} \right).$$

Moreover if the problem is linearly separable by a hyper-plane $\bar{u}$ and the first stage is run until convergence, R is less than $2(1 + \eta F^{\frac{2}{q}}) \|\bar{u}\|_{2,p}$.

The theorem first shows that a good estimate of R can speed-up the convergence of the algorithm. In particular if the first term is dominant, the convergence rate is $\mathcal{O}(\frac{q \log T}{\lambda T})$. If the second term is predominant, the convergence rate is $\mathcal{O}(\frac{R\sqrt{q \log T}}{\sqrt{T}})$, so it becomes independent from λ (i.e. independent from C). The algorithm will always optimally interpolate between these two different rates of convergence. As said before, the rate of convergence depends on p , through q . When p tends to 1, the solution tends to the sparse one of (7), with a worst rate. However in the experiment section we show that the best performance is not always given by the sparsest solution. Moreover Theorem 1

²For an extended version of this paper with proofs, see <http://francesco.orabona.com/papers/obscure-proofs.pdf>.

also shows that, when p is close to 1, the convergence rate has a sublinear dependency on the number of kernels, F , and if the problem is linearly separable it can have a faster convergence rate using more kernels. We will explain this formally in Section 3.2.

The training time of OBSCURE is proportional to the number of steps given by Theorem 1 multiplied by the complexity of each step. This in turn is dominated by the prediction (line 5 in Algorithms 1 and 2), that has complexity $\mathcal{O}(NFM)$. Note that this complexity is common to any other similar algorithm, and it can be reduced using methods like kernel caching [5].

In the following we introduce the necessary mathematical tools to be able to derive OBSCURE and its theorem.

3.1. Batch p -norm MKL

We first state a Lemma that is a generalization of Theorem 1 in [6] to general norms, using the framework in [19]. We need two additional definitions. Given a convex function $f : S \rightarrow \mathbb{R}$, its Fenchel conjugate $f^* : S \rightarrow \mathbb{R}$ is defined as $f^*(\mathbf{u}) = \sup_{\mathbf{v} \in S} (\mathbf{v} \cdot \mathbf{u} - f(\mathbf{v}))$. A vector $\mathbf{x}$ is a sub-gradient of a function f at $\mathbf{v}$, indicated with $\partial f(\mathbf{v})$, if $\forall \mathbf{u} \in S, f(\mathbf{u}) - f(\mathbf{v}) \geq (\mathbf{u} - \mathbf{v}) \cdot \mathbf{x}$.

Lemma 1. *Let $h(\cdot) = \frac{\alpha}{2} \|\cdot\|^2$ be a 1-strongly convex function w.r.t. a norm $\|\cdot\|$ over S . Assume that for all t , $g_t(\cdot)$ is a σ -strongly convex function w.r.t. $h(\cdot)$, and $\|\mathbf{z}_t\|_* \leq L_t$. Then for any $\mathbf{u} : \|\mathbf{u} - \mathbf{w}_t\| \leq 2R$, and for any sequence of non-negative $\xi_1, \dots, \xi_T$, Algorithm 3 achieves the following bound for all $T \geq 1$,*

$$\sum_{t=1}^T (g_t(\mathbf{w}_t) - g_t(\mathbf{u})) \leq \sum_{t=1}^T \left[4\xi_t R^2 + \frac{L_t^2}{\sigma t + \frac{\sum_{i=1}^t \xi_i}{\alpha}} \right].$$

With this Lemma we can now design stochastic sub-gradient algorithms. In particular, setting $\|\cdot\|_{2,p}$ as norm, $h(\bar{\mathbf{w}}) = \frac{q}{2} \|\bar{\mathbf{w}}\|_{2,p}^2$, and $g_t(\bar{\mathbf{w}}) = \frac{\lambda}{q} h(\bar{\mathbf{w}}) + \ell(\bar{\mathbf{w}}, \mathbf{x}_t, y_t)$, we obtain Algorithm 2 that solves the p -norm MKL problem in (9). In particular lines 6-7 correspond to the calculation of the sub-gradient of the multiclass loss function (5). The updates are done on the dual variables $\bar{\theta}_t$, in lines 11-12, that are transformed into $\bar{\mathbf{w}}_t$ in line 13.

Note also that Algorithm 2 can start from any vector, while this is not possible in the Pegasos algorithm where at the very first iteration the starting vector is multiplied by 0 [20]. The parameter R is basically an upper bound on the norm of the optimal solution, i.e. $R \geq \|\bar{\mathbf{w}}^*\|_{2,p}$. In the next Section we show how to initialize this algorithm and to calculate R in an efficient way.

3.2. Initialization through an online algorithm

In Theorem 1 we saw that if we have a good estimate of R , the convergence rate of the algorithm can be much faster.

Algorithm 3 Proximal projected sub-gradient descent

```

1: Input:  $R, \sigma, \mathbf{w}_1 \in S$ 
2: Initialize:  $s_0 = 0$ 
3: for  $t = 1, 2, \dots, T$  do
4:   Receive  $g_t$ 
5:    $\mathbf{z}_t = \partial g_t(\mathbf{w}_t)$ 
6:    $s_t = s_{t-1} + \frac{\sqrt{(\alpha\sigma t + s_{t-1})^2 + \frac{\alpha L_t^2}{R^2}} - (\alpha\sigma t + s_{t-1})}{2}$ 
7:    $\eta_t = (\sigma t + s_t/\alpha)^{-1}$ 
8:    $\mathbf{w}_{t+1} = \nabla h^*(\nabla h(\mathbf{w}_t) - \eta_t \mathbf{z}_t)$ 
9: end for

```

Moreover starting from a *good* solution could speed-up the algorithm even more.

We propose to initialize Algorithm 2 with an online algorithm. Algorithm 1 is the online version of problem (9) and it is derived using Corollary 7 in [11]. It is similar to the $2p$ -norm matrix Perceptron in [3], but it overcomes the disadvantage of being used with the same kernel on each feature. As in [3], for Algorithm 1 it is possible to prove a relative mistake bound. We omit the details for lack of space, a future longer version of this work will include it.

We can run it just for few iterations and then evaluate its norm and its loss. In Algorithm 1 R is then defined as

$$R := \sqrt{\|\bar{\mathbf{w}}_{T+1}\|_{2,p}^2 + \frac{2}{\lambda N} \sum_{i=1}^N \ell(\bar{\mathbf{w}}_{T+1}, \mathbf{x}_i, y_i)} \geq \sqrt{\|\bar{\mathbf{w}}^*\|_{2,p}^2 + \frac{2}{\lambda N} \sum_{i=1}^N \ell(\bar{\mathbf{w}}^*, \mathbf{x}_i, y_i)} \geq \|\bar{\mathbf{w}}^*\|_{2,p}. \quad (10)$$

So at any moment we can stop the algorithm and obtain an upper bound on $\|\bar{\mathbf{w}}^*\|_{2,p}$.

If the dimension of the space induced by the F kernels is big enough, it is very likely that the classification problem is linearly separable. When this is the case, we can prove that Algorithm 1 will converge to a solution which has null loss on each training sample, in a finite number of steps. More specifically we can state the following Theorem.

Theorem 2. *Suppose that $\|\phi^j(\mathbf{x}_t, y_t)\|_2 \leq 1, \forall j = 1, \dots, F, t = 1, \dots, N$, and $1 < p \leq 2$. If the problem (9) is linearly separable by a hyperplane $\bar{\mathbf{u}}$, then the Algorithm 1 will converge to a solution in a finite number of steps less than $2q(1/\eta + F^{\frac{2}{q}})\|\bar{\mathbf{u}}\|_{2,p}^2$. Moreover the returned value of R will be less than $2(1 + \eta F^{\frac{2}{q}})\|\bar{\mathbf{u}}\|_{2,p}$.*

From the theorem it is clear the role of η : a bigger value will speed up the convergence, but will decrease the quality of the estimate of R . So η governs the trade-off between speed and precision of the first stage. If p is close to 1, the dependency on the number of kernels in this theorem is strongly sublinear, moreover increasing the number of kernels to $F' > F$, we have that $\|\bar{\mathbf{u}}\|_{2,p}^2$ will most likely decrease or remain constant. This means that we expect Algorithm 1 to converge, in a number of steps that is almost

independent on F and in some cases even *decreasing* in F . The same consideration holds for the value of R returned by the algorithm, that can decrease when we increase the number of kernels. A smaller value of R will mean a faster convergence of the second stage. We will confirm this statement experimentally in Section 4.

4. Experiments

In this section we test OBSCURE on the Oxford flowers [17], Caltech-101 [7] and MNIST [15] datasets. Although our MATLAB implementation of the algorithm³ is not optimized for speed, it is already possible to observe the advantage of the low runtime complexity. This is particularly evident when training on datasets containing large numbers of categories and lots of training samples. In all our experiments, the parameter η is fixed at 2, and p is chosen from the set $\{1.01, 1.05, 1.10, 1.25, 1.50, 1.75, 2\}$. The regularization parameter λ is set through CV, as $\frac{1}{CN}$, where $C \in \{1, 10, 100, 1000\}$.

4.1. Oxford flowers

The Oxford flowers dataset [17] contains 17 different categories of flowers. Each class has 80 images with three predefined splits (train, validation and test). The authors also provide seven precomputed distance matrices⁴. These distance matrices are transformed into kernel using $\exp(-\gamma^{-1} \cdot d)$, where γ is the mean of the pairwise distances and d is the distance between two examples. We used a value of p equal to 1.05, found through CV.

We have implemented an extended version of the original Pegasos algorithm [20, 21] for problem (9). We first compare the running time performance between OBSCURE and Pegasos. Their generalization performance on the testing data (Figure 1(Left)) as well as the value of the objective function (Figure 1(Right)) are shown in Figure 1. In the same Figure, we also present the results obtained using other combination methods: SILP [22], SimpleMKL [18] and LP- β [9]. The cost parameter is selected from the range $C \in \{1, 10, 100, 1000\}$ for MKL methods. We see that OBSCURE converges much faster compared to Pegasos. This proves that, as stated in Theorem 1, OBSCURE has a better convergence rate than Pegasos. All the feature combination methods achieve similar results on this dataset. LP- β is order of magnitudes faster as it uses an efficient standard SVM solver [5].

4.2. Caltech-101 datasets

The Caltech-101 [7] dataset is a standard benchmark dataset for object categorization. Here we followed the experimental setup originally proposed and widely used in the

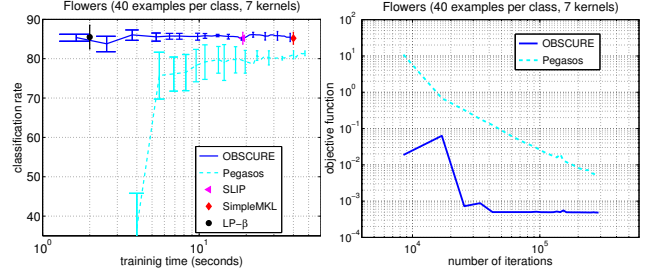


Figure 1. Comparison of performance on Oxford flowers dataset.

literature. In our experiments, we used the pre-computed features and kernels of [9], with the same training and test split⁵. This allows us to compare against them directly. Following that, we report results using all 102 classes of the Caltech-101 dataset using five splits. There are five different image descriptors, using different setup of parameters and computed at different scales. It results in a total of 39 kernels. Note that, as they are derived from 5 features only, some of them might be redundant. For brevity, we omit the details of the features and kernels that can be found in [9].

Figure 2 shows the behavior of our algorithm using different values of the parameter p (Figure 2(Left)), different number of kernels (Figure 2(middle)) and the running time under different size of training examples (Figure 2(right)). The dashed line in Figure 2(left & middle) corresponds to the results obtained by the first online stage of the OBSCURE algorithm. It can be observed from the figures that:

- [Figure 2(Left)] The online step of OBSCURE achieves a performance close to the optimal solution in a training time order of magnitudes faster (10^1 to 10^3). When p is large (*i.e.* q is small) the online stage converges even faster. This is consistent with Theorem 2.
- [Figure 2(Left)] By changing p , it is possible to improve performance. As stated before, when p tends to 1, the solution gets close to the sparse solution. In particular here 3 $\|w^j\|_2$ (out of 4) approach 0. When p equals 2, we obtain a dense solution, that corresponds to use the sum of all the kernels. Although some of the kernels may contain redundant information, all of them may be informative for classification. Thus imposing sparsity on them does not always help increasing performance. Hence the optimal p here is 1.10 – 1.25.
- [Figure 2(Middle)] OBSCURE has a better converges rate when there are more kernels, as stated in Theorem 2. That is, the algorithm achieves a given accuracy in less iterations when more kernels are given.
- [Figure 2(Right)] We can see that the algorithm converges quite fast to the optimal solution. Using 15 examples per class, the run time is similar to the runtime of LP- β (about 24 mins). When the number of training

³Code available at <http://dogma.sourceforge.net/>

⁴www.robots.ox.ac.uk/~vgg/research/flowers/

⁵www.vision.ee.ethz.ch/~pgehler/projects/iccv09/

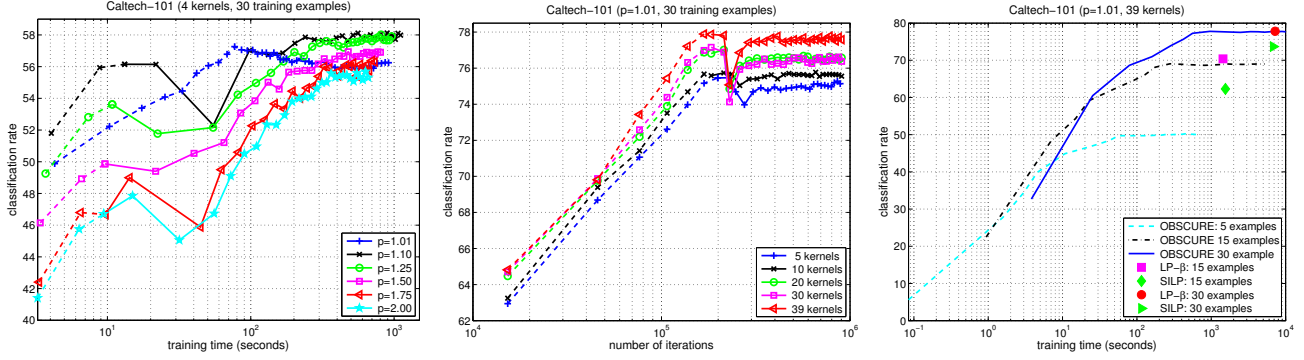


Figure 2. Behaviors of the OBSCURE algorithm on Caltech-101 dataset: (Left) the effect of different value of p , using four PHOG [2] kernels computed at different spatial pyramid level, as similar experiment performed in [9]; (Middle) the effect of different number of kernels randomly sampled from the 39 kernels; (Right) running time for different number of training examples using all the 39 kernels.

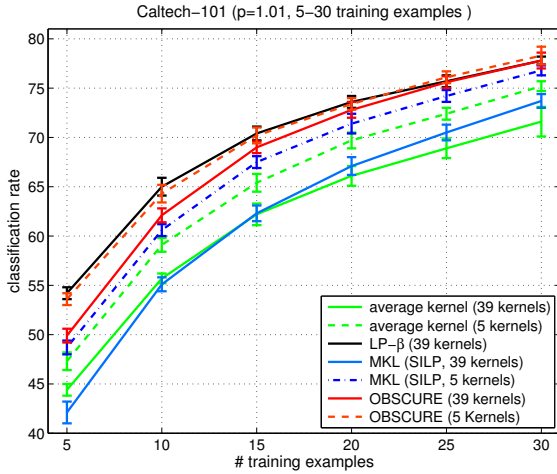


Figure 3. Performance comparison on Caltech-101 using different combination methods.

examples increases to 30, our algorithm has an advantage over LP- β , that takes about 121 mins.

In Figure 3 we report the results obtained using different combination methods. The best results for OBSCURE were obtained when p is at the smallest value (1.01). This is probably because among these 39 kernels many were redundant or not discriminative enough. For example, the worst single kernel achieves only an accuracy of $13.5\% \pm 0.6$ when trained using 30 images per category, while the best single kernel achieves $69.4\% \pm 0.4$. Thus, sparser solutions are to be favored. We see also that our method achieves performance comparable to the state-of-art (LP- β , [9]), and outperforms the other MKL (SILP) methods. One possible reason may be the one-vs-all multiclass extension used in the MKL algorithm. The sparse MKL algorithm may choose different subset of kernels in different independent binary classification tasks, which may introduce a bias on some classes in the final decision process. However, note that although our algorithm obtains a solution close to the

sparse one, it will never reach a completely sparse solution. This may be one of the reasons for the gap in performance between OBSCURE and LP- β [9]. However, this may not be critical, since usually in practice all used features/kernels are informative. Non-informative/duplicate features are unlikely to be included in a real system. We did a simple test by selecting five kernels from the five different families of features [9] which achieve low leave-one-out (LOO) error using 30 training examples per class. It can be done automatically using LS-SVM, which has a closed form solution for LOO error estimation [4]. The results as well as the performance of the averaging of these five kernels are also shown in Figure 3. We see that the algorithm improves slightly over the previous one. This suggests that OBSCURE as well as SILP, when provided with discriminative features, could increase performance even further. It also seems to indicate that there is a margin to improve the regularization used in MKL methods, as currently more kernels do not necessarily transform into better accuracy.

4.3. MNIST

In the last experiment we use the MNIST [15] dataset of handwritten digits. The dataset has a training set of 60,000 gray-scale 28×28 pixel digit images for training and 10,000 images for testing. We cut the original digit image into four square blocks (14×14) and obtained an input vector from each block. We used three kernels on each block: a linear kernel, a polynomial kernel and a RBF kernel, resulting in 12 kernels. Figure 4 shows the generalization performance on the test set achieved by OBSCURE over time, for various sizes of training set. We see that OBSCURE quickly converges to the best performance. It also shows that the time to reach the optimum is approximately linear in the number of training samples. The SVM performance using averaging kernel and the best kernel is also plotted. Notice that in the figure we only show the results of up to 20,000 training samples for the sake of comparison, otherwise we

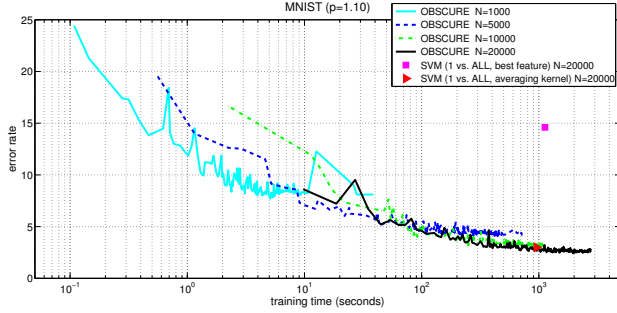


Figure 4. The generalization performance of MNIST dataset over different size of training samples.

could not cache all the 12 kernels in memory. However, by computing the kernel “on the fly” we are able to solve the MKL problem using the full 60,000 examples efficiently.

5. Conclusions and Discussion

This paper presents OBSCURE, a novel and efficient algorithm for solving p -norm MKL. It uses a hybrid two-stages online-batch approach, optimizing the objective function directly in the primal with a stochastic subgradient descent method. Experiments show that OBSCURE achieves state-of-art performance on multiclass classification problems. Furthermore, the solution found by the on-line stage is close to the optimal one for various tasks, while being computed several orders of magnitude faster. Our approach is general, hence it can be applied to any other algorithm with a strongly convex regularizer [11]. For example the framework can be very easily extended to solve other problems such as *structure output prediction* [23], to have an MKL algorithm for structured output.

OBSCURE has a faster convergence rate as the number of cues/kernels grows. Thus we expect to achieve better performance with more discriminative features. A simple feature selection technique such as cross-validation could already be beneficial. On the other hand, our results show that non-sparse models might get better performance (in the sense of accuracy and speed). This is in agreement with recent findings in [12]. As a last remark, we notice that the disadvantageous results of MKL methods, reported in [9], may be because those algorithms does not have a proper multiple class formulation for the object categorization problems. By using our method, MKL can still be an efficient machine learning tool for cue combination tasks.

Acknowledgments

The kernel matrixes of Caltech-101 were kindly provided by Peter Gehler, who we also thank for his useful comments. This work was sponsored by the EU project DIRAC IST-027787. FO also acknowledges partial support by the PASCAL2 NoE FP7-216886.

References

- [1] F. R. Bach, G. R. G. Lanckriet, and M. I. Jordan. Multiple kernel learning, conic duality, and the SMO, algorithm. In *Proc. ICML*, 2004. 2, 3
- [2] A. Bosch, A. Zisserman, and X. Munoz. Representing shape with a spatial pyramid kernel. In *Proc. CIVR*, 2007. 1, 7
- [3] G. Cavallanti, N. Cesa-Bianchi, and C. Gentile. Linear algorithms for online multitask classification. In *Proc. COLT*, 2008. 5
- [4] G. Cawley. Leave-one-out cross-validation based model selection criteria for weighted LS-SVMs. In *Proc. IJCNN*, 2006. 7
- [5] C. C. Chang and C. J. Lin. *LIBSVM: A Library for Support Vector Machines*, 2001. Software available at www.csie.ntu.edu.tw/~cjlin/libsvm. 5, 6
- [6] C. B. Do, Q. V. Le, and C.-S. Foo. Proximal regularization for online and batch learning. In *Proc. ICML*, 2009. 4, 5
- [7] L. Fei-Fei, R. Fergus, and P. Perona. One-shot learning of object categories. *IEEE PAMI*, 28(4):594–611, 2004. 6
- [8] P. Gehler and S. Nowozin. Let the kernel figure it out: Principled learning of pre-processing for kernel classifiers. In *Proc. CVPR*, 2009. 1
- [9] P. Gehler and S. Nowozin. On feature combination for multiclass object classification. In *Proc. ICCV*, 2009. 1, 2, 6, 7, 8
- [10] D. Hush, P. Kelly, C. Scovel, and I. Steinwart. QP algorithms with guaranteed accuracy and run time for support vector machines. *JMLR*, 7, 2006. 1, 4
- [11] S. Kakade, S. Shalev-Shwartz, and A. Tewari. On the duality of strong convexity and strong smoothness: Learning applications and matrix regularization. Technical report, TTI, 2009. 3, 5, 8
- [12] M. Kloft, U. Brefeld, S. Sonnenburg, P. Laskov, K.-R. Müller, and A. Zien. Efficient and accurate lp-norm multiple kernel learning. In *Proc. NIPS*, 2009. 1, 3, 8
- [13] A. Kumar and C. Smorchescu. Support kernel machines for object recognition. In *Proc. ICCV*, 2007. 1
- [14] G. Lanckriet, N. Cristianini, P. Bartlett, and L. E. Ghaoui. Learning the kernel matrix with semidefinite programming. *JMLR*, 5, 2004. 2
- [15] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86:2278–2324, 1998. 6, 7
- [16] M. E. Nilsback and B. Caputo. Cue integration through discriminative accumulation. In *Proc. CVPR*, 2004. 1, 2
- [17] M.-E. Nilsback and A. Zisserman. A visual vocabulary for flower classification. In *Proc. CVPR*, 2006. 6
- [18] A. Rakotomamonjy, F. R. Bach, S. Canu, and Y. Grandvalet. SimpleMKL. *JMLR*, 9:2491–2521, November 2008. 1, 2, 3, 6
- [19] S. Shalev-Shwartz and Y. Singer. Logarithmic regret algorithms for strongly convex repeated games. Technical Report 2007-42, The Hebrew University, 2007. 3, 5
- [20] S. Shalev-Shwartz, Y. Singer, and N. Srebro. Pegasos: Primal estimated sub-gradient solver for SVM. In *Proc. ICML*, 2007. 3, 4, 5, 6
- [21] S. Shalev-Shwartz and N. Srebro. SVM optimization: inverse dependence on training set size. In *Proc. ICML*, 2008. 2, 4, 6
- [22] S. Sonnenburg, G. Rätsch, C. Schäfer, and B. Schölkopf. Large scale multiple kernel learning. *JMLR*, 7, 2006. 1, 2, 3, 6
- [23] I. Tschantz, T. Hofmann, T. Joachims, and Y. Altun. Support vector machine learning for interdependent and structured output spaces. In *Proc. ICML*, 2004. 1, 2, 3, 8
- [24] M. Varma and D. Ray. Learning the discriminative power-invariance trade-off. In *Proc. ICCV*, 2007. 1, 2
- [25] A. Vedaldi, V. Gulshan, M. Varma, and A. Zisserman. Multiple kernels for object detection. In *Proc. ICCV*, 2009. 1
- [26] D. Wolpert. Stacked generalization. *Neural Networks*, 5(2), 1992. 2
- [27] A. Zien and C. S. Ong. Multiclass multiple kernel learning. In *Proc. ICML*, 2007. 1, 2, 3