

Predicting the Performance of Learning Algorithms Using Support Vector Machines as Meta-Regressors

Silvio B. Guerra, Ricardo B. C. Prudêncio and Teresa B. Ludermir

Center of Informatics, Federal University of Pernambuco
Pobox 7851 - CEP 50732-970 - Recife (PE) - Brazil
silvio.guerra@gmail.com, rbcp@cin.ufpe.br
tbl@cin.ufpe.br

Abstract. In this work, we proposed the use of Support Vector Machines (SVM) to predict the performance of machine learning algorithms based on features of the learning problems. This work is related to the Meta-Regression approach, which has been successfully applied to predict learning performance, supporting algorithm selection. Experiments were performed in a case study in which SVMs with different kernel functions were used to predict the performance of Multi-Layer Perceptron (MLP) networks. The SVMs obtained better results in the evaluated task, when compared to different algorithms that have been applied as meta-regressors in previous work.

1 Introduction

Algorithm selection is an important aspect to the success of the Machine Learning applications [1]. This task is traditionally supported by empirically evaluating the candidate algorithms using the available data, which can demand on expensive computational resources [2]. Algorithm selection can also be guided by expert rules, however such expert knowledge is not always easy to acquire, specially for new algorithms [3].

Considering the above context, different authors in literature have investigated the automatic acquisition of knowledge to predict the performance of learning algorithms, and to support algorithm selection, in a kind of Meta-Learning [1–12]. The knowledge in Meta-Learning is acquired from a set of meta-examples, in which each meta-example stores features of a learning problem solved in the past and information about the performance obtained by the candidate algorithms (the base-learners) on the problem. The knowledge in this case can be represented as a learning model (i.e., the meta-learner) that relates features of the problems and the performance of the learning algorithms.

A specific approach to Meta-Learning consists of using a regression algorithm to predict the value of a chosen performance measure (e.g., classification error) of the candidate algorithms based on the features of the problems. This approach is referred in the literature as Meta-Regression [10]. The meta-learner

is a regression model that may be used to select the best candidate algorithm considering the highest predicted performance measure. Different regression algorithms have been applied in this context, such as decision trees, linear regression and instance-based algorithms [3, 5, 6].

In the current work, we investigated the use of Support Vector Machines (SVMs) [13] to Meta-Regression. In the literature, SVMs have been successfully applied to many different problems, achieving very competitive results when compared to other machine learning algorithms [14]. In the context of Meta-Regression, there is no previous work that evaluated the use of SVMs to predict the performance of learning algorithms.

Experiments were performed in a case study which consists of predicting the performance of Multi-Layer Perceptron (MLP) networks [15]. A set of 50 meta-examples was produced from the application of the MLP on 50 different learning problems. In the meta-level, we applied SVMs with different kernel functions (among polynomial and RBF kernels) to predict the value of the Normalized Mean of Squared Errors (NMSE) of the MLPs. The predictions of the NMSE were based on a set of 10 pre-defined features of the learning problems (e.g., number of training examples and correlation between attributes). The performance of the SVMs was evaluated in a leave-one-out experiment performed on the 50 meta-examples.

As a basis of comparison, we also performed experiments in the meta-level using three different benchmark regression algorithms which were already used in previous work to Meta-Regression: the M5 algorithm (decision trees), the linear regression, and the 1-nearest neighbor algorithm. The performed experiments revealed that the SVMs with simple polynomial kernels obtained better performance in the meta-regression task when compared to the benchmark algorithms. The experiments also revealed that good performance can also be achieved by using SVM with RBF kernels, depending, in this case, on an adequate choice of the kernel's parameters.

Section 2 brings a brief introduction about Meta-Learning, including the Meta-Regression approach. Section 3 presents details of the proposed work, as well as the case study. Section 4 brings the experiments and obtained results. Finally, section 5 presents some final conclusions and the future work.

2 Meta-Learning

According to [16], there are different interpretations of the term *Meta-Learning*. In our work, we focused on the definition of Meta-Learning as the automatic process of acquiring knowledge that relates the performance of learning algorithms to the features of the learning problems [1]. In this context, each *meta-example* is related to a learning problem and stores: (1) the features describing the problem, called *meta-features*; and (2) information about the performance of one or more algorithms when applied to the problem. The *meta-learner* is a learning system that receives as input a set of such meta-examples and then acquires knowledge used to predict the algorithms performance for new problems being solved.

The meta-features are, in general, statistics describing the training dataset of the problem, such as number of training examples, number of attributes, correlation between attributes, class entropy, among others [7, 8]. In a strict formulation of Meta-Learning, each meta-example stores, as performance information, a class label which indicates the best algorithm for the problem, among a set of candidates [4]. In this case, the class label for each meta-example is defined by performing a cross-validation experiment using the available dataset. The meta-learner is simply a classifier which predicts the best algorithm based on the meta-features of the problem.

Although the strict Meta-Learning approach (as described above) has been applied by different authors (such as [2, 4, 9, 17, 18]), certain information loss may be introduced in the definition of the class labels associated to meta-examples. For instance, the performance of two algorithms may be very similar, and this information will be lost by merely recording the best algorithm as class label [5].

In order to overcome the above difficulty, the Meta-Regression approach [10] tries to directly predict the numerical value of accuracy (or alternatively the error rate) of each candidate algorithm. In this case, the meta-examples store as performance information the numerical values of accuracy obtained in previous problems. The meta-learner, in turn, is a regression model that may be used either to select the best candidate algorithm based on the highest predicted accuracy or to provide a ranking of algorithms based on the order of predicted accuracies.

In [3], the authors evaluated different algorithms as meta-regressors, including linear regression models, piecewise linear models, decision trees and instance-based regression. In [6], the authors used linear regression models to predict the accuracy of 8 classification algorithms, and the experiments revealed good results. In [5], the authors performed comparative experiments with both the strict Meta-Learning and Meta-Regression approaches, and observed that the latter one performed better when used to support algorithm selection.

3 Meta-Regression with SVMs

As seen in section 2, Meta-Regression is a flexible approach to predicting the performance of learning algorithms, supporting a more informative solution to algorithm selection.

In the current work, we investigate the use of Support Vector Machines (SVMs) in the context of Meta-Regression. SVM is a state-of-the-art algorithm in the Machine Learning field, successfully applied to different classification and regression problems [13, 14]. Previous work in Meta-Regression has applied different regression methods to produce meta-learners (as seen in section 2), yielding relative success. However, to the best of our knowledge, there is no evaluation of the use of SVMs in this task.

Figure 1 brings the architecture of the Meta-Regression which summarizes our proposal. Each meta-example is composed by the meta-features of a learning task and the performance information derived from the empirical evaluation

of the learning algorithm on the task. The set of generated meta-examples is given as input to the SVM algorithm, which will produce a regression model responsible for predicting the algorithm’s performance for new problems based on its meta-features.

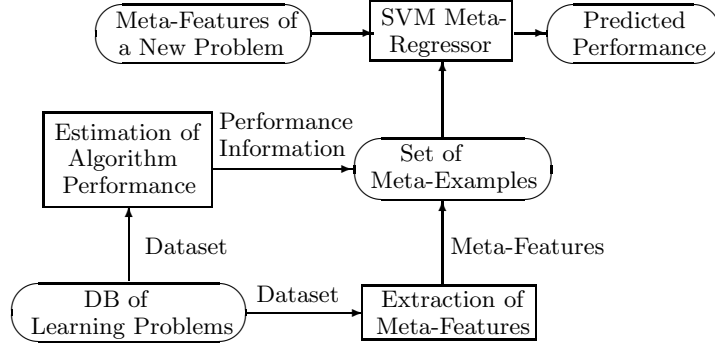


Fig. 1. Architecture of Meta-Regression

A case study was performed in our work in a meta-learning task which corresponds to predict the performance of Multi-Layer Perceptron (MLP) networks [15] in regression problems. In this case study, we generated a set of 50 meta-examples from the application of the MLP in 50 regression problems (see section 3.1). Each meta-example is associated to a single problem and stores: (1) the value of 10 descriptive meta-features (see section 3.2); and (2) the test error obtained by the MLP when evaluated in the regression problem (see section 3.3). The set of meta-examples is given as input to a regression algorithm which will be able to predict the error of the MLP for new problems only based on the meta-attributes of the problems.

In the following sections, we provide details about the construction of the set of meta-examples, as well as the details on the SVMs used for Meta-Regression.

3.1 Datasets

In order to generate meta-examples, we collected 50 datasets corresponding to 50 different regression problems, available in the WEKA project ¹. On average, the collected datasets presented 4,392 examples and 13.92 attributes.

We observed in these datasets a large variability in both the number of training examples and the number of attributes. This variation may be convenient

¹ These datasets are specifically the sets provided in the files *numeric* and *regression* available to download in <http://www.cs.waikato.ac.nz/ml/weka/>

to Meta-Learning studies, since it is expected that the algorithms present significantly different patterns of performance, depending on the problem being solved.

The attributes in the original datasets were normalized to the $[-1; +1]$ interval, aiming a better treatment by the MLP network. The order of the examples in each dataset were randomly permuted, in order to avoid eventual biases derived from the original dataset acquisition.

3.2 Meta-Features

In the developed work, a total number of 10 meta-features was used to describe the datasets of regression problems:

1. LogE - Log of the number of training examples;
2. LogEA - Log of the ration between number of training examples and number of attributes;
3. MinCT, MaxCT, MeanCT and StdCT - Minimum, maximum, mean and standard deviation of the absolute values of correlation between predictor attributes and the target attribute;
4. MinCAtr, MaxCAtr, MedCAtr and DesvCAtr - Minimum, maximum, mean and standard deviation of the absolute values of correlation between pairs of predictor attributes.

The meta-feature LogE is an indicator of the amount of data available for training, and LogEA, in turn, indicates the dimensionality of the dataset. The meta-features MinCT, MaxCT, MedCT and DesvCT indicate the amount of relevant information available to predict the target attribute. The meta-features MinCAtr, MaxCAtr, MedCAtr and DesvCAtr, in turn, indicate the amount of redundant information in the dataset. This set of meta-features was chosen by considering features adopted in previous work. As this set is probably non optimal, in future work, we will consider new features.

3.3 Performance Information

The final step to generate the meta-examples is to evaluate the performance of the MLP in the 50 collected regression tasks. From this evaluation, we can produce the performance information stored in the meta-examples which will correspond to the target attribute in the Meta-Regression task. In order to measure the performance of the MLP in each problem, the following methodology of evaluation was applied.

Each dataset was divided in the training, validation and test subsets, in the proportion of 50%, 25% and 25%, respectively. As usual, the training subset was used in the adjustment of the MLP's weights, the validation subset is used to estimate the generalization performance of the MLP during training, and the test subset is used to evaluated the performance of the trained MLP. In our work, we applied the standard Backpropagation (BP) algorithm [15] to train the

MLP network². The optimal number of hidden nodes³ was empirically defined, by testing the values 1, 2, 4, 8, 16 and 32. For each number of hidden nodes, the MLP is trained 10 times with random initial weights. The stopping criteria of the training process followed the benchmarking rules provided in [20].

The optimal number of hidden nodes was finally chosen as the value in which the trained MLP obtained the lowest average NMSE (Normalized Mean of Squared Errors) on the validation subset over the 10 training runs. The NMSE measure is defined as:

$$NMSE = \frac{\sum_{i=1}^n (y_i - \tilde{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (1)$$

In the equation, n is the number of examples in the validation subset, y_i and $\tilde{y}_i$ are respectively the target and the predicted value for example i , and $\bar{y}$ is the average of the target attribute.

The performance information which is stored in the meta-example is the average NMSE obtained by the trained MLP (with optimal number of hidden nodes) on the test subset. According to [21], a property of NMSE which is interesting for Meta-Learning is that the values have no scale and are comparable across different datasets (which it would not be feasible if we had used a non-normalized error measure). Values of NMSE lower to 1 indicate that the MLP provided better predictions than the mean value. Values higher than 1 indicate that the MLP was not useful in the regression problem.

We highlight here that there are different algorithms to train MLPs, as well as other methodologies of training, that could have been used to improve performance. However, our aim in this work is not to achieve the best possible performance with MLPs but to *predict* the learning performance. Other learning algorithms to train MLPs (such as the Levenberg-Marquardt algorithm [22]) can be applied in the future as new case studies, possibly with more effective strategies to train the networks.

3.4 Meta-Regressor

In our case study, the Meta-Regression task is to predict the NMSE measure of the MLP based on the features of the problems given as input. In our experiments, this task is dealt with SVM regressors.

An important aspect in the development of SVM is the chosen kernel function. In our work, we evaluated the use of two different types of kernel functions in the SVMs. First, we evaluated homogeneous polynomial kernels represented as:

$$K(x, x') = (x \cdot x')^p \quad (2)$$

² The BP algorithm was implemented by using the NNET Matlab toolbox [19]. Learning rates were defined by default.

³ The MLP was defined with only one hidden layer.

In the above equation, x and x' are instances in an attribute space. In our experiments, we used the values $p = 1$ (linear kernel) and $p = 2$ (quadratic kernel), which are the more common options for polynomial kernels. We also deployed in our work a Radial Basis Function (RBF) kernel in the form:

$$K(x, x') = e^{-\gamma \cdot \|x - x'\|^2} \quad (3)$$

The RBF kernel is a non-linear function that, in comparison to the polynomial kernels, it is expected to handle more complex relationships between predictor and target attributes. In our experiments, we evaluated the values 0.01, 0.05 and 0.1 for the parameter γ .

Obviously, there are other possibilities of kernel functions (e.g., sigmoid kernels), as well as parameter settings that we intend to experiment in future work. Besides, automated procedures can be used in the future to define the values of these parameters [23].

In this work, we deployed the implementation of SVMs provided in the WEKA environment [24], which applies the sequential minimal optimization algorithm proposed by [13] for training the SVM regression model.

4 Experiments and Results

In this section, we evaluated the performance of the Meta-Regression process for the generated set of meta-examples. In our experiments, the SVM meta-regressor was evaluated by using a leave-one-out procedure. At each step, 49 meta-examples are used as the training set, and the remaining meta-example is used to test the trained meta-learner. This step is repeated 50 times, using at each time a different test meta-example. The meta-learning performance is then evaluated based on the predictions generated by the meta-learner in the test samples.

Two different criteria were used to evaluate the meta-learner: (1) the Normalized Mean of Squared Errors (NMSE) computed for the predictions of the meta-learner in the test meta-examples; and (2) the Correlation (COR) between the predictions generated by the meta-learner and the true values of the target attribute stored in the test meta-examples.

As a basis of comparison, the above methodology of experiments was also applied using three different methods as meta-regressors:

1. M5 algorithm: which was proposed by Quinlan [25] to induce regression trees;
2. 1-Nearest Neighbor (1-NN) algorithm: a special case of instance-based learning algorithm [26];
3. Linear Regression (LR) model.

All these methods were already used in previous work as meta-regressors (see, for instance, [3]). We highlight that these algorithms are representatives of different families of regression methods, providing different inductive biases for the learning problem being solved in the case study.

Table 1. Results obtained by Meta-Regression in the leave-one-out experiment.

	NMSE	COR
SVM (linear)	0.486	0.731
SVM (quadratic)	0.457	0.758
SVM (RBF, $\gamma = 0.1$)	0.374	0.794
SVM (RBF, $\gamma = 0.05$)	0.484	0.724
SVM (RBF, $\gamma = 0.01$)	0.595	0.701
M5	0.605	0.631
1-NN	0.539	0.710
LR	0.595	0.658

As it can be seen in Table 1, the SVM regressor with polynomial kernels (linear and quadratic) obtained better results on both evaluation measures when compared to the benchmark methods M5, 1-NN and LR. The best result among the polynomial kernels was obtained by the quadratic kernel, which also yielded a very competitive result (the second best one) when we include the SVM with RBF kernel in the comparison.

Considering the runs with RBF kernel, we observed that the performance of the SVM was very sensitive to the value of parameter γ . For $\gamma = 0.1$, the SVM obtained the best result over all evaluated algorithms and parameter settings. For $\gamma = 0.05$, the SVM also performed well compared to the other algorithms (yielding the third best result). However, for $\gamma = 0.01$, the SVM obtained a performance which was only better than the method M5. This result indicates that an adequate choice of the parameter of RBF kernels (i.e., the model selection) is an aspect that should be more carefully addressed in the case of SVM meta-regressors. Model selection is in fact a relevant topic of research in SVMs, which can be handled, for instance, by considering the characteristics of the data [13, 21, 27] and by deploying search techniques [23, 28, 29]. Such strategies will be considered in future work to model selection in our case study.

5 Conclusion

In this work, we presented the use of SVMs to Meta-Regression, which aims to predict the performance of learning algorithms based on the features of the learning problems being solved.

Experiments were performed in a case study which consisted in predicting the numerical performance of MLP networks when applied to regression tasks. A set of 50 meta-examples was generated in this case study in order to perform a comparative analysis of different meta-regressors. The results of a leave-one-out experiment revealed the viability of using SVMs in the investigated case study. The performance of the SVM meta-regressor was in general better than the performance obtained by the benchmark algorithms applied as a basis of comparison.

The current work has limitations that will be dealt with in future work. Although we applied, in the case study, polynomial and RBF kernels which are widely used in the literature, different other functions can be evaluated. Furthermore, we only tested few values for the kernel parameters, which could be optimized by using more sophisticated approaches. We also include as future work, the evaluation of SVM meta-regressors in new case studies related to other machine learning algorithms.

References

1. Giraud-Carrier, C., Vilalta, R., Brazdil, P.: Introduction to the special issue on meta-learning. *Machine Learning* **54**(3) (2004) 187–193
2. Kalousis, A., Hilario, M.: Representational issues in meta-learning. In: *Proceedings of the 20th International Conference on Machine Learning*. (2003) 313–320
3. Gama, J., Brazdil, P.: Characterization of classification algorithms. In: *Progress in Artificial Intelligence, 7th Portuguese Conference on Artificial Intelligence, EPIA-95*. (1995) 189–200
4. Aha, D.: Generalizing from case studies: A case study. In: *Proceedings of the 9th International Workshop on Machine Learning*, Morgan Kaufmann (1992) 1–10
5. C. Koepf, C.T., Keller, J.: Meta-analysis: Data characterisation for classification and regression on a meta-level. In: *Proceedings of the International Symposium on Data Mining and Statistics*. (2000)
6. Bensusan, H., Alexandros, K.: Estimating the predictive accuracy of a classifier. In: *Proceedings of the 12th European Conference on Machine Learning*. (2001) 25–36
7. Brazdil, P., Soares, C., da Costa, J.: Ranking learning algorithms: Using IBL and meta-learning on accuracy and time results. *Machine Learning* **50**(3) (2003) 251–277
8. Kalousis, A., Gama, J., Hilario, M.: On data and algorithms - understanding inductive performance. *Machine Learning* **54**(3) (2004) 275–312
9. Prudêncio, R.B.C., Ludermir, T.B.: Meta-learning approaches to selecting time series models. *Neurocomputing* **61** (2004) 121–137
10. Koepf, C. In: *Meta-regression: performance prediction*. (2006) 89–106
11. Prudêncio, R.B.C., Ludermir, T.B.: Active learning to support the generation of meta-examples. In: *Proc. of the International Conference on Artificial Neural Networks*. (2007) 817–826
12. Prudêncio, R.B.C., Ludermir, T.B.: Active selection of training examples for meta-learning. In: *Proc. of the International Conference on Hybrid Intelligent Systems*. (2007)
13. Smola, A., Scholkopf, B.: A tutorial on support vector regression. *Statistics and Computing* **14**(3) (2004) 199–222
14. Wang, L.: *Support Vector Machines: Theory and Applications*. Springer (2005)
15. Rumelhart, D.E., Hinton, G.E., Williams, R.J.: Learning representations by back-propagating errors. *Nature* **323** (1986) 533–536
16. R. Vilalta, C. Giraud-Carrier, P.B., Soares, C.: Using meta-learning to support data-mining. *International Journal of Computer Science Application* **I**(31) (2004) 31–45
17. Prudêncio, R.B.C., Ludermir, T.B., de Carvalho, F.A.T.: A modal symbolic classifier to select time series models. *Pattern Recognition Letters* **25**(8) (2004) 911–921

18. Leite, R., Brazdil, P.: Predicting relative performance of classifiers from samples. In: Proceedings of the 22nd International Conference on Machine Learning. (2005)
19. Demuth, H., Beale, M.: Neural Network Toolbox: For use with MATLAB: User's Guide. The Mathworks (1993)
20. Prechelt, L.: A set of neural network benchmark problems and benchmarking rules. Technical Report 21/94, Fakultät für Information, Universität Karlsruhe, Karlsruhe, Germany (1994)
21. Soares, C., Brazdil, P., Kuba, P.: A meta-learning approach to select the kernel width in support vector regression. *Machine Learning* **54**(3) (2004) 195–209
22. Levenberg, K.: A method for the solution of certain non-linear problems in least squares. *Quarterly Journal of Applied Mathematics* **II**(2) (1944) 164–168
23. O. Chapelle, V. Vapnik, O.B., Mukherjee, S.: Choosing multiple parameters for support vector machines. *Machine Learning* **46**(1) (2002) 131–159
24. Witten, I.H., Frank, E., eds.: WEKA: machine learning algorithms in Java. University of Waikato, New Zealand (2003)
25. Quinlan, J.: Learning with continuous classes. In: Proceedings of the Australian Joint Conference on Artificial Intelligence. (1992) 343–348
26. Aha, D., Kibler, D.: Instance-based learning algorithms. *Machine Learning* **6**(3) (1991) 37–66
27. P. Kuba, P. Brazdil, C.S., Woznica, A.: Exploiting sampling and meta-learning for parameter setting support vector machines. In: Proceedings of the IBERAMIA 2002. (2002) 217–225
28. Cawley, G.: Model selection for support vector machines via adaptive step-size tabu search. In: Proceedings of the International Conference on Artificial Neural Networks and Genetic Algorithms. (2001) 434–437
29. S. Lessmann, R.S., Crone, S.: Genetic algorithms for support vector machine model selection. In: Proceedings of the International Joint Conference on Neural Networks 2006. (2006) 3063–3069