

SUPPORT VECTOR MACHINES AND LEARNING ABOUT TIME

Stefan Rüping and Katharina Morik

University of Dortmund
CS Department, AI Unit
Dortmund, Germany

ABSTRACT

The analysis of temporal data is an important issue in current research, because most real-world data either explicitly or implicitly contains some information about time. The key to successfully solving temporal learning tasks is to analyze the assumptions and prior knowledge that can be made about the temporal process of the learning problem and find a representation of the data and a learning algorithm that makes effective use of this knowledge. This paper will present a concise overview of the application of Support Vector Machines to different temporal learning tasks and the corresponding temporal representations.

1. INTRODUCTION

There is a multitude of learning tasks related to temporal phenomena and, correspondingly, there are many possible representations for temporal data. Learning tasks and representations are closely related: the No Free Lunch Theorem [1] implies that finding an adequately biased representation can make a hard learning problem easy (and, vice versa, that finding this representation itself is hard).

Statistical time series analysis has developed two big classes of representations, namely those in the time domain and those in the frequency domain [2]. Analysis in the time domain is based on the correlation between the current and previous observations, while the frequency domain tries to decompose the time series into cyclic components at different frequencies. For learning tasks, time series analysis has the following objectives: description (to describe a time series by certain statistics), explanation (to understand the process behind a time series), prediction (to predict future values of the time series) and control (control the process behind the time series to generate certain future values).

For Machine Learning, an overall theory of temporal analysis is much less developed. Learning tasks are usually taken from specific real-world problems and represen-

tations are often constructed ad hoc. Morik [3] differentiates between two different aspects of time, the linear precedence of events and immediate dominance of temporal categories. These terms originate from natural language theory [4]. Immediate dominance refers to the construction of higher-level categories of the time-dependent elements, exemplary learning tasks based on the concept of immediate dominance are the discovery of frequent episodes [5] or first order logic learning based on Allen's interval relations [6]. The aspect of linear precedence refers to the linear temporal order of the single events and is most prominent in the framework of time series analysis.

The focus of this paper lies on the time series representation and the types of learning tasks that can be solved with support vector machines (SVMs, [7]). Support vector machines have been applied to very different kinds of learning problems, for example to time series prediction by Mukherjee et al. in [8] and by Müller et al. in [9]. A regression problem for time series has been solved with SVMs in [10], where certain coefficients of chemical components have been predicted from chromatography time series. Chang et al. [11] have presented an approach for time series segmentation with SVMs.

All of these applications are based on a single time series representation, the phase space representation, i. e. creating d -dimensional examples by moving a window of length d over the time series. The theoretical foundation for this is the theorem of Takens [12, 13] which states that for dynamical systems of a certain type, the phase-space reconstruction and the unobserved internal structure of the system are topologically identical, given the embedding dimension is large enough. At first glance, this seems like the perfect tool: We know that we can find out all we need to know about a time series simply by looking at large enough time windows and in the SVM we have a learning algorithm that is well known to perform well on high-dimensional data. But this conclusion is fallacious. First of all, Taken's theorem does only hold for dynamical systems which can be described by differential equations of a certain form. Generally, one cannot decide whether this is the case for a given real-world data set. Second, the theory of structural risk minimization [7],

The financial support of the Deutsche Forschungsgemeinschaft (SFB 475, "Reduction of Complexity for Multivariate Data Structures") is gratefully acknowledged.

on which the SVM is based, is only formulated for independent, identically distributed data. Clearly, the independence assumption is violated for time series data. Although versions of the central theorems of structural risk minimization do also hold for dependent data of weak dependence structure [14] and in practice, SVMs have been shown to perform quite well on time series data, one should be careful to transfer results of the SVM to this type of data. Finally, even if the premises of Taken's theorem and the structural risk minimization principle do hold, a different representation of the data may lead to a much easier generalization (it is the "large enough dimension"-part of Taken's theorem that can cause much trouble).

In the next section we will show the relation between the SVM and statistical time series modeling, in particular autoregressive models and the fourier transform. After that, in section 3, we will discuss alternative representations that were used in different applications with real-world data. Finally, section 4 will present novel temporal learning tasks which can be solved using SVMs.

2. TIME SERIES MODELS

Using the phase space model with a linear prediction function leads to the class of autoregressive (AR) models [2]. Obviously, AR models can be learned by a SVM with linear kernel, so it does not surprise that the SVM does not perform very different on data generated from an AR model than other methods for AR model estimation, like the Yule-Walker equations. However, it can be seen that the SVM is more robust against outliers in the data. The following table compares the mean absolute error of a AR model learned using the Yule-Walker equations against a SVM model. In the first case, the data was generated from an AR model, in the second case 10% of outliers were added (validation set results averaged over 4 runs with different models).

Data	Tool	d=2	d=3	d=4	d=5
AR model	AR	0.640	0.624	0.624	0.633
	SVM	0.648	0.634	0.632	0.639
AR+outliers	AR	0.942	0.922	0.919	0.911
	SVM	0.885	0.826	0.832	0.827

Fig. 1. Performance of AR and SVM model.

For time series analysis in the frequency domain, the fourier transformation can be used to transform the examples for the SVM. There also exist kernel function, which make this transformation explicitly, e. g. the fourier kernels proposed by Vapnik in [7], Ch. 11. or the time-frequency kernel of Davy et al. [15].

3. TIME SERIES REPRESENTATIONS FOR REAL-WORLD DATA

Time series models are a well understood research area. However, in practice we see over and over again that much data manipulation has to be done in order to get good results. In this section we will show some examples of representation tricks that can be used when analysing time series with Support Vector Machines. The main advantage of SVMs is that the relevant equations that describe the generalization errors of SVMs [7] do not depend on the dimensionality of the data but only on the margin of the separating hyperplane, which makes the SVM especially suited for high-dimensional data. While this reasoning is not strictly valid - the margin depends on the geometry of the data and hence also on the dimension - empirical evidence show that this property of SVMs does hold in practice (see e. g. [16]). This allows us to enhance the representation with certain attributes which improve the temporal analysis of the data.

Often, time series can be decomposed into a long-term linear trend, a cyclical component and a rest, where the trend and the cyclical component are fitted before the actual analysis of the rest is being performed. For the trend, when analysing a time series on the basis of the phase space representation and a linear kernel (i. e. using an AR model), we can simply add the time t as another attribute to the data. By this, the SVM will learn a function $w_1 x_1 + \dots + w_d x_d + w_{d+1} t$, i. e. it will automatically decide, which effects in the data to attribute to the trend $w_{d+1} t$ and which to the rest model $w_1 x_1 + \dots + w_d x_d$.

The case of cyclical components is more complicated, because they cannot be as easily identified and filtered from the data as a simple trend. But often the most difficult problem in practice is that lots of statistical procedures for modeling periodic functions cannot be applied, because what looks like a periodic component actually is not one. In Figure 2 you can see the weekly sales of a certain item in a retail store. One can easily see the effects of christmas sales in the 51st and preceeding weeks and also the lower effects of easter sales in the 12th week. One can easily imagine that these effects will occur periodic every year at the same time. But they don't! In some years, christmas is in the 51st week of the year, but in some it is already in the 50th week. On a daily time scale, christmas will repeat itself every 365 days in most years, but in 366 days in a leap year. The data of the eastern can vary about five weeks. Therefore, in [17] 20 additional binary attributes were used to mark the presence of holidays, special sale promotions, and other significant events in that particular week. In the example of Figure 2, this reduced mean absolute error by more than 20% from 116.15 to 91.063 (test error over one year, averaged over 4 different stores).

Another case of sales time series is shown in Figure 3.

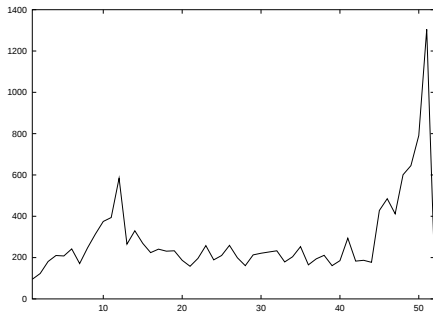


Fig. 2. Retail Store Sales per Week.

This time, the subject are newspapers and the sales are reported on a daily basis. We can clearly see that there is a cycle of 6 days in the data (the newspaper does not appear on sundays) and that sales are especially high for the weekend edition on Saturdays. Hence, instead of using the phase space representation on the original time series, we can also reason that we have indeed 6 different time series of sales (one for each weekday) and use the phase space representation there. And third, we can also use a mixed representation where we use the last couple of days plus the last couple of weekdays to predict the time series. Testing these three representations, we can see that modeling only one time series achieves a medium absolute test error of 8.102, the approach with 6 time series has an error of 5.942 and the mixed approach reaches an error of 5.654.

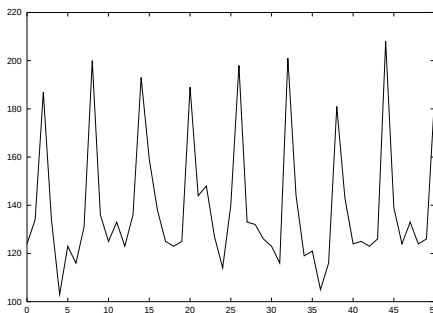


Fig. 3. Newspaper Sales per Day.

Sometimes, best results are achieved if one drops the idea of a time series at all. For example, for the task of recommending drug administration from recorded vital signs of intensive care patients - a high dimensional, noisy classification problem on multivariate time series - it was found in [18] that the best representation was to ignore time dependencies completely and make a non-temporal classification based on the last observation only. In the field of chromatography, Ritthoff et al. [10] solved the problem of predicting certain chemical coefficients based on the chromatographical analysis of a substance (which is a time series of

intensities of chemical components) by describing the time series by chosen analytical properties, e. g. the location of its maximum. That is, they did not use the time series observations at all, but only new, aggregated features. This technique alone reduced the error by 50%.

4. ADVANCED TEMPORAL LEARNING TASKS

The phase space representation bases on the assumption that the temporal dependence structure of the time series can be sufficiently captured in a short finite window of observations. This allows the examples generated from each window to be treated as if they were generated independently and thus, the order in which the examples are presented to the learning algorithm is not important. This assumption fails, if the process that generates the time series changes over time. This scenario is called concept drift. Usually, concept drift is treated by using only a certain number of the newest examples, where the actual number of examples used is chosen heuristically. For SVMs, Klinkenberg and Joachims [19] proposed an approach where this number is chosen based on efficient performance estimators for SVMs [20]. This approach was generalized to weighted examples and more flexible ways of choosing the examples in the training set in [21].

Another interesting problem is the detection of outliers in time series. The procedure of Bauer [22] constructs a certain ellipse (derived from an assumed autoregressive model) around the phase space representation of the time series all points outside of this ellipse are declared as outliers. This procedure can be generalized by using a SV estimation [23] of the support of the points in the phase space. This directly applies the definition of Davies and Gather [24] of the α -outlier-region as the region of points with the lowest probability, so that the probability of the whole region is α .

5. SUMMARY AND CONCLUSIONS

This paper presented a concise overview of time series analysis using Support Vector Machines. Of course, due to space constraints we could not cover all existing approaches in as much detail as they deserved.

In conclusion, we find that the key to successful time series analysis with SVMs lies in finding the right representation. The excellent generalization properties of the SVM, especially its good performance on high-dimensional data, make it easy to improve results by adding additional temporal features or constructing specialised kernel functions.

6. REFERENCES

- [1] D.H. Wolpert and W.G. Macready, "No free lunch theorems for optimisation," *IEEE Trans. on Evolutionary*

Computation, vol. 1, pp. 67 – 82, 1997.

- [2] Christopher Chatfield, *The Analysis of Time Series: An Introduction*, Chapman and Hall, 3rd edition, 1984.
- [3] Katharina Morik, “The representation race - preprocessing for handling time phenomena,” in *Proc. of the European Conference on Machine Learning (ECML 2000)*, Ramon López de Mántaras and Enric Plaza, Eds. 2000, vol. 1810 of *LNAI*, Springer Verlag Berlin.
- [4] Gerald Gazdar and Chris Mellish, *Natural Language Processing in PROLOG*, Addison Wesley, Workingham u.a., 1989.
- [5] H. Mannila, H. Toivonen, and A. Verkamo, “Discovering frequent episode in sequences,” in *Procs. of the 1st Int. Conf. on Knowledge Discovery in Databases and Data Mining*. 1995, AAAI Press.
- [6] J. F. Allen, “Towards a general theory of action and time,” *Artificial Intelligence*, vol. 23, pp. 123–154, 1984.
- [7] V. Vapnik, *Statistical Learning Theory*, Wiley, Chichester, GB, 1998.
- [8] Sayan Mukherjee, Edgar Osuna, and Frederico Girosi, “Nonlinear prediction of chaotic time series using support vector machines,” in *Proc. of IEEE NNSP 97*, Amelia Island, FL, Sep 1997.
- [9] K. Müller, A. Smola, G. Ratsch, B. Schölkopf, J. Kohlmorgen, and V. Vapnik, “Predicting time series with support vector machines,” in *Proc. of the International Conference on Artificial Neural Networks*. 1997, Lecture Notes in Computer Science, Springer.
- [10] Oliver Ritthoff, Ralf Klinkenberg, Simon Fischer, and Ingo Mierswa, “A hybrid approach to feature selection and generation using an evolutionary algorithm,” in *Proc. of the UKCI-02*, J. Bullinaria, Ed., Birmingham, UK, 2002, pp. 147–154, Univ. of Birmingham.
- [11] Ming-Wei Chang, Chih-Jen Lin, and Ruby C. Weng, “Analysis of nonstationary time series using support vector machines,” in *Pattern Recognition with Support Vector Machines — First International Workshop, SVM 2002*, Seong-Whan Lee and Alessandro Verri, Eds. Aug 2002, LNCS 2388, pp. 160–170, Springer.
- [12] F. Takens, “Detecting strange attractors in turbulence,” in *Dynamical systems and turbulence*, D. A. Rand and L. S. Young, Eds., vol. 898 of *Lecture Notes in Mathematics*, pp. 366 – 381. Springer, Berlin, 1980.
- [13] N. H. Packard, J. P. Crutchfield, J. D. Farmer, and R. S. Shaw, “Geometry from a time series,” *Physical Review Letters*, vol. 45, pp. 712–716, 1980.
- [14] Thomas Fender, *Empirische Risikominimierung für dynamische Datenstrukturen*, Ph.D. thesis, Department of Statistics, University of Dortmund, Germany, work in progress.
- [15] M. Davy, A. Gretton, A. Doucet, and P.W.J. Rayner, “Optimised support vector machines for nonstationary signal classification,” *IEEE transactions on Signal Processing*, vol. 9, no. 12, Dec. 2002.
- [16] Thorsten Joachims, “Text categorization with support vector machines: Learning with many relevant features,” in *Proceedings of the European Conference on Machine Learning*, Claire Nédellec and Céline Rouveirol, Eds., Berlin, 1998, pp. 137 – 142, Springer.
- [17] Stefan Rüping, “Zeitreihenprognose für Warenwirtschaftssysteme unter Berücksichtigung asymmetrischer Kostenfunktionen,” M.S. thesis, Universität Dortmund, 1999, in German.
- [18] Katharina Morik, Michael Imhoff, Peter Brockhausen, Thorsten Joachims, and Ursula Gather, “Knowledge discovery and knowledge validation in intensive care,” *Artificial Intelligence in Medicine*, 2000.
- [19] Ralf Klinkenberg and Thorsten Joachims, “Detecting concept drift with support vector machines,” in *Proc. of the Seventeenth International Conference on Machine Learning (ICML)*, P. Langley, Ed., San Francisco, CA, 2000, pp. 487–494, Morgan Kaufmann.
- [20] Thorsten Joachims, “Estimating the generalization performance of a SVM efficiently,” in *Proceedings of the International Conference on Machine Learning*, Pat Langley, Ed., San Francisco, 2000, pp. 431–438, Morgan Kaufman.
- [21] Ralf Klinkenberg and Stefan Rüping, “Concept drift and the importance of examples,” in *Text Mining – Theoretical Aspects and Applications*, Jürgen Franke, Gholamreza Nakhaeizadeh, and Ingrid Renz, Eds. Springer, Berlin, Germany, 2002, To appear.
- [22] M. Bauer, U. Gather, and M. Imhoff, “Analysis of high dimensional data from intensive care medicine,” in *Proceedings in Computational Statistics*, R. Payne, Ed., Berlin, 1999, Springer Verlag.
- [23] Bernhard Schölkopf, Robert C. Williamson, Alex J. Smola, and John Shawe-Taylor, “SV estimation of a distribution’s support,” in *Neural Information Processing Systems 12*, S.A. Solla, T.K. Leen, and K.-R. Müller, Eds. 2000, MIT Press.
- [24] Laurie Davies and Ursula Gather, “The identification of multiple outliers,” *J. Am. Statist. Ass.*, vol. 88, pp. 782–792, 1993.