

ORIGINAL RESEARCH

Estimation of the manufacturing industry sub-sectors' capacity utilization rates using support vector machines

Halil Ibrahim Erdal^{*1}, Alp Baray², Şakir Esnaf²

¹ *Turkish International Cooperation and Coordination Agency (TİKA), Atatürk Bulvarı, Ulus Ankara, Turkey*

² *Istanbul University, Industrial Engineering Department, İstanbul, Turkey*

Received: April 25, 2013

Accepted: October 16, 2013

Online Published: October 25, 2013

DOI: 10.5430/air.v3n1p1

URL: <http://dx.doi.org/10.5430/air.v3n1p1>

Abstract

Capacity utilization rate is one of the most important indicators of the efficiency of the manufacturing industry, and therefore of the return of the investments made. Estimation of these rates accurately renders it possible to make important economic decisions such as taking sectorial investment decisions, defining the optimal distribution of sectorial credits, determining non-competitive sectors, making development plans and developing unemployment policies. In this study, we estimated the capacity utilization rates of 21 sub-sectors of the Turkish manufacturing industry using support vector machines and compared the results with the results obtained from the methods of artificial neural networks and vector auto-regression. This study is the first in the literature in that it was carried out using this method.

Key Words: Analysis of variance, Artificial neural networks, Capacity utilization rates, Support vector machines, Vector auto-regression

1 Introduction and literature review

With the rise of international trade and financial transactions, the competition in today's markets has become global and this situation has had a massive influence on decision making units. In order to be successful in the global commercial environment, decision making units have begun to search for optimal selection and analysis methods, whereas the issue of estimating financial indicators has gained importance.

To be able to measure potential future risks and benefits, the accurate estimation of economic and financial indicators is of great importance for governments that execute monetary and financial policies, banks, intermediary institutions, firms, investors and savers. For a decision made on the basis of a wrong estimation might not only drive the financial system and the national economy into ruination starting from small savers, but also cause decision units to lose their

competitive advantages.

The disciplines of statistics and econometrics have assumed a pioneering role in not only testing economic theories empirically but also estimating economic and financial indicators. Numerous methods have been developed in order to estimate economic and financial indicators such as regression, auto-regression and vector auto-regression along with quantitative methods like artificial neural networks and fuzzy logic. All these methods are used to solve the problem of estimating the future.

Vapnik^[1] offered a new and powerful approach to the problem of estimation with his support vector machines (SVM). The basis of SVM dates back to the studies carried out by Vapnik from the late 1960s to the early 1990s,^[1-6] however, it was first used in 1995^[7] with the aim of solving a classification problem. SVMs are used in the analysis of classification and regression problems and the range of application is rapidly widening in parallel with the advancements in com-

^{*}**Correspondence:** Dr. Halil Ibrahim Erdal; Email: halilerdal@yahoo.com; Address: Turkish Cooperation and Coordination Agency (TİKA), Atatürk Bulvarı No:15 Ulus Ankara, Turkey.

puter technologies.

The literature contains numerous studies on the estimation of manufacturing industries' capacity utilization rates (CUR). While classical econometric methods are used in most of them, a very little number of them include artificial intelligence techniques. Berndt and Morrison^[8] estimated CUR using an econometric model based on the relationship of long-run average cost curve and capacity curve.

Corrado and Matthey^[9] projected an econometric model between CUR values and consumer price index, and reported a positive correlation between them. Kim^[10] suggested an internal model that estimates capacity utilization rates and the indicators of them. Kim, in this study, determined the negative influence of high raw material and investment costs on CURs and the positive influence of high energy prices. Burnside et al.^[11] used an econometric model based on electricity consumption and outstanding capital in order to make CUR estimation through a time series analysis. Hornstein,^[12] on the other hand, estimated CUR using econometric models based on capital and labour values. Azeez^[13] investigated the CUR trend of the Indian industrial manufacturing and the factors influencing CUR, using translog variable cost function. In this study, Azeez suggested that CUR demonstrates scarce resources as well as the level of demand, and that it is a criterion of industrial performance. Shapiro^[14] and Koenig^[15] examined the relationship between CUR and industrial manufacturing in the USA, and reported a strong correlation between them.

Bansak et al.^[16] formed an econometric model in order to estimate CUR for 111 sectors in the USA for the period of 1974–2000. In their model, they employed the data of the change in industrial production, the ratio of investments to capital and machine park average age. Similarly, Corrado et al.^[17] evaluated the CUR and industrial production data of the USA for the period of 1977–1992 and discussed the methods used in the estimation of these data. Beaumont,^[18] on the other hand, developed an econometric model based on labor statistics for the estimation of CUR, and tested this model on the US industry CUR values. In addition, whereas Lee^[19] compared primary and secondary parametric CUR estimations in the literature, Ragan,^[20] De Leeuw,^[21] Shapiro,^[22] Lalonde^[23] and Morin & Stevens^[24] compared the estimation models of different US public institutions. Ray et al.^[25] formed an econometric model between labour and CUR. Using the US Bureau of Labour Statistics data for 1970–2001; they made an industrial CUR estimation through data envelopment analysis. Färe et al.^[26] and Sahoo and Tone^[27] similarly employed data envelopment analysis in CUR estimation. Lieberman^[28] estimated the capacity increases and CURs of 40 chemical-production industries through Manne, Newsboy and Whitt-Luss models; and compared the results through multiple regression analysis. Kalyuzhnova and Vagliasindi^[29] estimated the CURs of Kazakhstan firms for a 7-year period including the Russian financial crisis using the panel data analysis method.

Sun^[30] employed input-output model analysis for the estimation of services sector CUR. Shaikh and Moudud^[31] set up an econometric model for CUR estimation and tested this model through IMF data for 9 OECD countries.

Dergiades and Tsoulfidis^[32] estimated the CURs of 14 European countries through vector auto-regression method. Tsoulfidis and Dergiades^[33] also estimated the US industrial CUR values through the same method for the period between 1948:1 and 2004:2. Dergiades and Tsoulfidis,^[34] in their study, selected capital and production values as independent variables, and estimated the US and Canadian CUR values obtained through questionnaire using vector auto-regression. Morrison^[35] employed dynamic optimization method in estimating CUR. Al-Ghandoor and Samhour^[36] estimated the relationship between CUR and electricity consumption in Jordan using the methods of linear regression and fuzzy artificial neural networks. While Gökçekuş^[37] estimated the Turkish plastic industry CUR using the generalized Leontief cost function, Gajanan and Malhotra^[38] employed the same method in order to estimate Indian industrial CUR values for the period of 1976–1996.

2 Methods

2.1 Artificial Neural Networks

Artificial neural networks (ANN) have been developed through inspiration from the data processing function of the biological neural system of human brain. Doctor Warren McCulloch and mathematician Walter Pitts developed the first ANN model in 1943. McCulloch and Pitts,^[39] being inspired by the human brain's calculation skill, modeled a simple neural network using electric circuits. After Frank Rosenblatt^[40] developed Perceptron, studies on ANN gathered speed. "Perceptron", which is a single-layered trainable network model that contains a single output, was created as a result of studies aimed at modeling brain's functions. There are numerous sources that provide detailed theoretical information regarding ANN, which are used widely today in economics and finance.

2.2 Vector Auto-regression

Vector auto-regression (VAR) method was developed by Sims in 1980. Each variable in the VAR model is a dynamic model of the equations that contain values dependent on the past movements of that variable and all other variables in the model.^[41] Relations between variables can be explained using single- or multi-variable models. In single-equation models, the relationship between the dependent and independent variables is explained as a one-way cause and effect relationship. It is not always easy to decide which variable is the cause and which variable is the effect. Therefore, the problem is solved using multi-equation models that explain the cause and effect relationship between the variables reciprocally.^[42]

While regression analysis is performed for relations that can be explained with a single equation, models with multiple equations require the solution of simultaneous equation systems. In simultaneous models, all equations are analyzed at the same time and thus the system's coefficients are estimated. Since the variables in these models are separated into two as internal and external variables, researchers may encounter problems in defining which one is the internal and which one is the external variable, and also in meeting the definition requirement that is necessary for the solution of the system.^[43]

2.3 Support Vector Machines

SVM are used widely in financial applications. Time series analysis is among the areas in which SVM are employed widely in practice. In classification, SVM use structural risk minimization (SRM) that is based on Vapnik-Chervonenkis (VC) dimension, which addresses learning from a statistical perspective and ensures the minimization of the upper limit error rate. In short, VC dimension defines the capacity of the functions cluster; if VC dimension is low, it is inferred that the expected probability of error is also low.^[44]

SRM is an induction method used in machine learning. The generalized model is chosen from a finite data set in machine learning; however, this most of the time creates the overfitting problem. SRM method overcomes this problem by balancing the complexity of the model but it sacrifices from training success while doing this. Despite this, SRM indirectly brings about a better generalization.

In addition, SVM also use empirical risk minimization that is aimed at reducing the error rate among the examples. SVM try to divide two classes through a hyperplane in a way to maximize the distance between two examples that are closest to threshold while minimizing the empirical classification error. What should not be overlooked here is that empirical risk minimization (ERM) is necessary for a good solution even though it does not guarantee a single solution. Since SVM use ERM and SRM together, it can escape from the risk of falling into the local minimum, unlike learning machines such as artificial neural networks.

SVM uses hyperplane to be able to make binary classification. SVM divides two training class' with a hyperplane, and while defining this hyperplane, it solves an optimization problem in order to maximize the hyperplane margin. The example vectors, which are the closest to the margins of the hyperplane that maximizes the distance between the two closest vectors to the borders of training clusters, are called support vectors. They can be linearly separated if the inequalities in the equation(1) are satisfied for the weight vector (W) and the threshold value (b) suitable for the training cluster in the equation.

$$(x_1, y_1), \dots, (x_n, y_n), \quad x_i \in R^d, \quad y_i \in \{-1, +1\} \quad (1)$$

$$w \cdot x_i + b \geq 1 \quad y_i = 1; \quad w \cdot x_i + b \leq -1 \quad y_i = -1; \quad (2)$$

The inequality in the equation(2) is generalized and the equation(3) is obtained; from which then the hyperplane equation(4) is derived.

$$y_i(w \cdot x_i + b) \geq 1, \quad i = 1, \dots, l \quad (3)$$

$$w \cdot x + b = 0 \quad (4)$$

In order to maximize the hyperplane margin (d); $\|w\|$ in equation (5), which is the norm of w is minimized.

$$d = \frac{2}{\|w\|} \quad (5)$$

Equation(5) is converted to equation(7) Lagrangian function under the equation(2) constraints.

$$L(w, b, \alpha) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^l \alpha_i [y_i(w \cdot x + b) - 1] \quad (6)$$

While Lagrangian multipliers are applied to equation-constrained non-linear models, Kuhn-Tucker Conditions are applied to inequation-constrained problems, and equation(6) is converted to equation(7).

$$L(\alpha) = \sum_{i=1}^l \alpha_i - \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l \alpha_i \alpha_j y_i y_j (x_i x_j) \quad (7)$$

In cases the problem cannot be solved linearly, ξ_i slack variables are added to the equation(3) to get equation(8).

$$y_i(w \cdot x_i + b) \geq 1 - \xi_i, \quad i = 1, \dots, l \\ \xi_i \geq 0 \quad i = 1, \dots, l \quad (8)$$

If the value ξ_i is lower than 1, it means that it is within the example hyperplane margins, and this condition is not regarded as an error operationally. However, if the value ξ_i is higher than 1, the upper limit of the number of training errors is calculated with $\sum_i \xi_i$. In this case, the quadratic equation in the equation(9) under the constraints in the equation(10) is obtained.^[45]

$$P = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^l \xi_i \quad (9)$$

$$\sum_{i=1}^l \alpha_i y_i \quad \text{ve} \quad 0 \leq \alpha_i \leq C \quad i = 1, \dots, l \quad (10)$$

In non-linear cases, on the other hand, the training cluster is mapped to a feature space that is higher dimensioned than an input space like in equation (11).

$$\phi: R^d \rightarrow R^N$$

$$\phi(x_i) = \phi_1(x_i), \phi_1(x_i), \dots, \phi_N(x_i), \quad i = 1, \dots, l \quad (11)$$

After mapping, equation(7) gets the form in equation(12). To make it easier for the operation weight, the kernel function of the inner product of $\phi(x) \cdot \phi(x_i)$ is expressed as (K) , as in equation(13).

$$L(\alpha) = \sum_{i=1}^l \alpha_i - \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l y_i y_j \alpha_i \alpha_j \phi(x_i) \phi(x_j) \quad (12)$$

$$K(x, x_i) = \phi(x) \cdot \phi(x_i) \quad (13)$$

Following this, equation(12) is converted to equation(14).

$$L(\alpha) = \sum_{i=1}^l \alpha_i - \frac{1}{2} \sum_{i=1}^l \sum_{j=1}^l y_i y_j \alpha_i \alpha_j K(x_i, x_j) \quad (14)$$

3 Application and EMPIRICAL results

3.1 Application and data

We obtained the data regarding the variables used in this study from the official websites of the Central Bank of the Republic of Turkey (CBRT) and Turkish Statistical Institute (TSI). As dependent variables; under the scope of International Standard Industrial Classification of All Economic Activities (Rev.3), CURs of 21 sub-sectors were used (ISIC 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26, 27, 28, 29, 31, 32, 33, 34, 35, 36). On the other hand, the independent variables used in CUR estimation are the exchange rate of dollar, interest rate on deposits, import and export values, industrial production index and consumer price index. The data are monthly and composed of 210 observations covering the period between March 1991 and August 2008.

In the study, firstly, we examined the data set composed of 210 months (1993M03-2008M08), then performed stationarity tests for economic series. Among the variables used in the study, the manufacturing industry CUR and sub-sector CUR were found stationary at their levels; however, since they were not stationary at the levels of other series, we took their natural logarithms and removed seasonality. After this operation, since these variables were still not stationary, their first differences were taken and the growth rates of interest, export, import and price index were obtained, and thus the variables became stationary.

The manufacturing industry CUR and sub-sector CUR were found stationary at their levels. We took natural logarithms

of all other variables of the study and removed seasonality. These variables were not stationary at their levels; therefore, by taking their first differences, the growth rates of interest, export, import and price index were obtained. The stationarity analysis was tested through Augmented Dickey Fuller (ADF) test for the stationary or trend-stationary model. According to the Schwartz information criterion, ADF test was performed when the maximum delay was 12. Time series are either stationary or non-stationary. Before statistically analyzing a time series, it is firstly necessary to investigate whether the process that created that series is constant over time or not. Stationarity of the time series analyzed is of importance, since the “t” and “F” tests and the “R” value might lead to fallacious results in the analysis. There are numerous studies that demonstrate that regression analyses conducted using non-stationary time series produce problematic results in terms of the statistical features. If the mean and variance of a stochastic process do not exhibit a systematic change over time, the series is called stationary. Time series are not generally stationary.

After making the variables stationary, we made CUR estimation through SVM, ANN and VAR. In the estimation, firstly, we trained using the CUR data set of 186 months (1991M03-2006M07) and then the data set of 24 months (2006M07-2008M08) was tested. Then, after explaining the technical details of the methods used in the study, we refined the estimation results and compared adjusted correlation of determination, mean absolute error and root mean square error within the framework of the evaluation criteria. In measuring the general success of the application, we used one-way variance analysis (ANOVA).

In this study, we used Weka 3.5.7 in SVM and ANN applications. We made many tests to select the parameters that would yield the best result. After these tests, the SVM parameters that give the best estimation are the following: The complexity parameter was selected as “C=10”. We normalized the data. As the kernel function, we selected the polynomial kernel and the exponent value as “e=10”.

Similarly, many tests were performed in the selection of the parameters that would give the best result. The ANN parameters that gave the best estimation are the following: number of hidden layers “h=5”, momentum “m= 0.3” and learning rate “l=0.3”. In the applications where the data were normalized, we set the number of iterations to 5000 and the error limit to 1/10000. We used Eviews 5.0 in VAR application. In the VAR algorithm, which was used in CUR estimation, we used the series as two lag according to the Akaike information criterion.

3.2 Performance measures and variance analysis

In this study, we used the criteria of adjusted correlation of determination (R^2_{adj}), mean absolute error (MAE) and root mean square error (RMSE) in measuring the success of the estimations conducted through SVM by comparing them

with ANN and VAR. R^2_{adj} takes a value between 0 and 1. The value 1 represents the perfect correlation. RMSE is used to determine the rate of error between calculation values and model forecasts. MAE questions the absolute error between calculation values and model forecasts. As RMSE and MAE values approach zero, the forecasting capacity of the model increases.

In measuring the general success of the application, on the other hand, we used non-iterative single factor variance analysis (ANOVA). ANOVA is a statistical method that is used to define the differences between the means of various populations. It was designed to define the differences between populations that represent different behaviors. It is a combined test; the equality of the means of varying numbers of populations is tested simultaneously and together. ANOVA tests whether the means of a certain number of populations are equal or not by looking at the two estimators of

the population variance. We conducted ANOVA application using Microsoft Excel and took the “p” value as 0.05.

3.3 Empirical results

In the manufacturing industry CUR estimation, in all 21 sub-sectors, SVM yielded the best result for the training stage R^2_{adj} values. In the test stage, on the other hand, we obtained two of the best results for R^2_{adj} values from ANN, two of them from VAR and the remaining 17 of them from SVM. We used ANOVA test to measure the general success of the results. As required by the method, null hypothesis results were tested for all sectors’ training stage R^2_{adj} values.

Null Hypothesis: There is no difference between the columns (analysis methods; SVM, ANN, VAR), see Table 1 and Table 2.

Table 1. Table for all sectors training and test stages R^2_{adj} values

R^2_{adj}	TRAINING			TEST		
	SVM	ANN	VAR	DVM	YSA	VAR
CUR						
CUR 15	0.8390	0.6162	0.5684	0.7530	0.5413	0.4831
CUR 16	0.7929	0.6854	0.6058	0.4820	0.4385	0.4173
CUR 17	0.8673	0.8181	0.7889	0.5974	0.4371	0.2791
CUR 18	0.7571	0.5629	0.5367	0.7485	0.1324	0.4734
CUR 19	0.7834	0.5591	0.4594	0.1600	0.2475	0.3718
CUR 20	0.8362	0.6760	0.6551	0.6536	0.4032	0.4615
CUR 21	0.8143	0.6416	0.5556	0.3919	0.0520	0.0637
CUR 22	0.8143	0.4782	0.4252	0.3919	0.0817	0.0669
CUR 23	0.8431	0.6086	0.5877	0.7597	0.4762	0.3167
CUR 24	0.7825	0.6126	0.6304	0.2010	0.0533	0.1805
CUR 25	0.8743	0.7396	0.7155	0.4938	0.2285	0.3549
CUR 26	0.8366	0.7633	0.7115	0.7714	0.6478	0.7113
CUR 27	0.7391	0.5454	0.5065	0.5625	0.0400	0.4019
CUR 28	0.8560	0.7429	0.7356	0.1533	0.2268	0.0580
CUR 29	0.8428	0.7993	0.7407	0.6799	0.5278	0.3722
CUR 31	0.8326	0.7038	0.6946	0.5102	0.2378	0.2076
CUR 32	0.8682	0.7920	0.7873	0.6482	0.5714	0.3846
CUR 33	0.8187	0.5653	0.5843	0.5819	0.4171	0.4797
CUR 34	0.8835	0.7969	0.7665	0.3461	0.3120	0.2186
CUR 35	0.7964	0.7356	0.7519	0.4455	0.5604	0.5492
CUR 36	0.8335	0.7498	0.7454	0.5498	0.4155	0.6100

Table 2. Table for the results of all sectors training stage R^2_{adj} values variance analysis

SUMMARY	Number	Total	Mean	Variance		
Column 1	21	17.31172	0.824368	0.001467		
Column 2	21	14.1923	0.675824	0.009885		
Column 3	21	13.55274	0.645369	0.01214		
ANOVA						
Source of Variance	SS	df	MS	F	P-value	F criterion
Columns	0.385234	2	0.192617	76.63774	2.08E-14	3.231733
Error	0.100534	40	0.002513			
Total	0.855092	62				

Table 3. Table for test results of the Variance analysis for all sectors test stage R^2_{adj} values

SUMMARY	Number	Total	Mean	Variance		
Column 1	21	10.88152	0.518167	0.036803		
Column 2	21	7.048035	0.335621	0.036777		
Column 3	21	7.461843	0.355326	0.031433		
ANOVA						
Source of Variance	SS	df	MS	F	P-value	F criterion
Columns	0.421603	2	0.210802	14.92777	1.44E-05	3.231733
Error	0.564857	40	0.014121			
Total	2.521838	62				

Since the P value for all sectors training stage R^2_{adj} values is $2.08E-14 < 0.05$, the null hypothesis is rejected. Column1 (SVM) total: 17.31172, mean: 0.824368 > column 2 (ANN) total: 14.1923, mean: 0.675824 > column 3 (VAR) total: 13.55274 mean: 0.645369. We stated earlier that the bigger the R^2_{adj} values, the better the results. Therefore, for the R^2_{adj} criterion, in the training stage of estimating all sectors' CUR values, the best method was found to be SVM, the second is ANN and the third is VAR. We performed ANOVA test for all sectors test stage R^2_{adj} values. Null hypothesis was tested (see Table 3).

Since the P value for all sectors test stage R^2_{adj} values is $1.44E-05 < 0.05$, the null hypothesis is rejected, which suggests that there exists difference between the analysis methods. Column1 (SVM) total: 10.88152, mean: 0.518167 > column 3 (VAR) total: 7.461843, mean: 0.355326 > column 2 (ANN) total: 7.048035 mean: 0.335621. For the R^2_{adj} criterion, in the test stage of estimating all sectors' CUR values, the best method was found to be SVM, the second is VAR and the third is ANN (see Table 4).

In the CUR estimation, we obtained the best results from SVM for the training stage mean absolute error values in all 21 sub-sectors. In the test stage, on the other hand, ANN gave one of the best results while SVM gave the remaining 20. We used ANOVA test to measure the general success of the results. For the training stage MAE values in all sectors, the results of the null hypothesis were tested.

Null Hypothesis: There is no difference between the columns (analysis methods; SVM, ANN, VAR).

Since the P value for all sectors training stage MAE values is $5.11E-12 < 0.05$, the null hypothesis is rejected, which means that there is difference between the analysis methods. Column1 (SVM) total: 53.35828, mean: 2.540871 < column 3 (VAR) total: 100.0679, mean: 4.76514 < column 2 (ANN) total: 103.021 mean: 4.905761. We stated earlier that the closer the MAE values to zero, the better the results. Therefore, for the MAE criterion, in the training stage of estimating all sectors' CUR values, the best method was found to be SVM, the second is VAR and the third is ANN (see Table 5).

Since the P value for all sectors test stage MAE values is $9.32E-06 < 0.05$, the null hypothesis is rejected, which means that there is difference between the analysis methods. Column1 (SVM) total: 36.3637, mean: 1.731605 < column 2 (ANN) total: 54.56104, mean: 2.598145 < column 3 (VAR) total: 62.59379, mean: 2.980657. For the MAE criterion, in the test stage of estimating all sectors' CUR values, the best method was found to be SVM, the second is ANN and the third is VAR (see Table 6 and Table 7).

In the manufacturing industry CUR estimation, SVM gave the best results for the training stage root mean square error values in 20 of the sub-sectors while VAR gave the best result in one of them. In the test stage, ANN gave two of the best results, VAR gave four and SVM gave fifteen of them. We used ANOVA test in measuring the general success of the results. For all sectors training stage MAE values, we tested the results for the null hypothesis.

Null Hypothesis: There is no difference between the

columns (analysis methods; SVM, ANN, VAR).

Since the P value for all sectors training stage RMSE values is $6.54E-07 < 0.05$, the null hypothesis is rejected, which means that there is difference between the analysis methods. Column1 (SVM) total: 99.10506, mean: 4.719289 < column 3 (VAR) total: 131.2963, mean: 6.252205 < column

2 (ANN) total: 136.4984 mean: 6.499924. We stated earlier that the closer the RMSE values to zero, the better the results. Therefore, for the RMSE criterion, in the training stage of estimating all sectors' CUR values, the best method was found to be SVM, the second is VAR and the third is ANN (see Table 8).

Table 4. Table for all sectors test stage mean absolute error values

MAE	TRAINING			TEST		
	SVM	ANN	VAR	SVM	ANN	VAR
CUR						
CUR 15	1.24	2.62	2.70	0.90	1.47	1.40
CUR 16	2.30	3.85	4.05	2.37	3.18	3.25
CUR 17	1.13	2.19	1.93	1.19	1.35	1.47
CUR 18	1.86	3.58	3.56	0.63	1.26	1.14
CUR 19	2.96	5.35	5.31	1.44	2.15	1.53
CUR 20	2.61	4.95	5.06	1.27	1.94	2.06
CUR 21	2.22	4.15	4.29	1.24	1.80	2.66
CUR 22	2.97	7.09	6.73	1.99	2.88	2.39
CUR 23	2.34	5.83	5.11	2.14	4.92	4.55
CUR 24	1.86	3.41	3.20	1.22	2.11	2.30
CUR 25	0.63	1.71	3.49	1.00	1.99	3.77
CUR 26	1.65	3.19	3.34	1.55	2.43	2.29
CUR 27	1.93	3.48	3.45	1.01	1.84	1.66
CUR 28	1.99	4.00	3.59	1.30	1.51	2.23
CUR 29	2.53	4.21	4.63	1.54	2.36	2.45
CUR 31	2.02	4.21	3.64	1.03	1.42	1.69
CUR 32	3.52	6.83	6.84	1.65	2.28	3.09
CUR 33	5.06	10.71	11.25	2.76	2.20	5.92
CUR 34	4.17	7.54	8.24	5.86	7.22	7.00
CUR 35	4.91	7.65	7.34	2.95	6.37	4.17
CUR 36	3.11	5.67	2.05	1.58	2.78	5.64

Table 5. Table for the results of the variance analysis on all sectors training stage mean absolute error values

SUMMARY	Number	Total	Mean	Variance		
Column 1	21	53.35828	2.540871	1.300459		
Column 2	21	103.021	4.905761	4.5694		
Column 3	21	100.0679	4.76514	5.055289		
ANOVA						
Source of Variance	SS	df	MS	F	P-value	F criterion
Columns	73.91899	2	36.9595	53.38501	5.11E-12	3.231727
Error	27.69279	40	0.69232			
Total	292.422	62				

Since the P value for all sectors test stage RMSE values is $0.016084 < 0.05$, the null hypothesis is rejected, which means that there is difference between the analysis methods. Column1 (SVM) total: 60.5832, mean: 2.884914 < column 2 (ANN) total: 69.6986, mean: 3.318981 < column 3 (VAR)

total: 77.48362 mean: 3.689696. Therefore, for the RMSE criterion, in the test stage of estimating all sectors' CUR values, the best method was found to be SVM, the second is ANN and the third is VAR (see Table 9).

Table 6. Table for the results of the variance analysis on all sectors test stage mean absolute error values

SUMMARY	Number	Total	Mean	Variance		
Column 1	21	36.3637	1.731605	1.260758		
Column 2	21	54.56104	2.598145	2.629634		
Column 3	21	62.59379	2.980657	2.669133		
ANOVA						
Source of Variance	SS	df	MS	F	P-value	F criterion
Columns	17.20136	2	8.600681	15.69062	9.32E-06	3.231727
Error	21.92567	40	0.548142			
Total	148.3919	62				

Table 7. Table for all sectors training stage “root mean square error” values

RMSE	TRAINING			TEST		
CUR	SVM	ANN	VAR	SVM	ANN	VAR
CUR 15	2.26	3.36	3.38	1.49	1.76	1.87
CUR 16	4.29	5.32	5.45	3.91	4.12	4.10
CUR 17	1.98	2.67	2.43	1.79	1.79	1.76
CUR 18	3.67	4.71	4.71	1.00	1.51	1.33
CUR 19	4.96	7.20	6.86	2.09	2.66	1.90
CUR 20	4.86	6.53	6.68	2.04	2.45	2.42
CUR 21	4.00	5.53	5.67	1.88	2.36	3.03
CUR 22	5.87	9.57	8.97	2.52	3.40	2.95
CUR 23	4.43	7.62	6.64	3.37	5.91	5.45
CUR 24	3.38	4.49	4.16	1.88	2.63	2.87
CUR 25	1.15	3.43	4.83	1.23	1.99	4.96
CUR 26	3.49	4.08	4.43	2.54	2.98	2.76
CUR 27	3.75	4.71	4.76	1.73	2.32	2.06
CUR 28	3.66	4.94	4.56	2.64	1.99	2.49
CUR 29	4.87	5.55	6.07	2.40	3.06	3.44
CUR 31	3.71	5.33	4.74	1.90	1.95	2.02
CUR 32	7.00	9.01	9.07	3.61	3.55	4.33
CUR 33	9.33	14.00	15.25	4.68	3.29	6.27
CUR 34	7.60	10.16	10.64	9.64	10.44	9.16
CUR 35	8.56	9.97	9.48	5.68	7.00	5.10
CUR 36	6.10	7.44	2.39	2.44	3.54	7.48

Table 8. Table for the results of the variance analysis on all sectors training stage root mean square error values

SUMMARY	Number	Total	Mean	Variance		
Column 1	21	99.10506	4.719289	4.352932		
Column 2	21	136.4984	6.499924	7.696452		
Column 3	21	131.2963	6.252205	9.216705		
ANOVA						
Source of Variance	SS	df	MS	F	P-value	F criterion
Columns	39.073	2	19.5365	20.76145	6.54E-07	3.231727
Error	37.63995	40	0.940999			
Total	464.3948	62				

Table 9. Table for the results of the variance analysis on all sectors test stage root mean square error values

SUMMARY	Number	Total	Mean	Variance		
Column 1	21	60.5832	2.884914	3.717488		
Column 2	21	69.6986	3.318981	4.531744		
Column 3	21	77.48362	3.689696	4.290825		
ANOVA						
Source of Variance	SS	df	MS	F	P-value	F criterion
Columns	6.814623	2	3.407312	4.587281	0.016084	3.231727
Error	29.71095	40	0.742774			
Total	257.6158	62				

4 Conclusions

Researchers have been putting intensive emphasis on estimating the future, and in this framework, estimating manufacturing industry capacity utilization rates. In this study, we addressed the subject of modeling based on estimation through support vector machines, and as the application, we estimated the monthly sectorial CUR of 21 main manufacturing industries in Turkey for the period between March 1991 and August 2008 (210 months). We then compared the results with those of the methods of artificial neural networks and vector auto-regression.

Firstly in the study, we examined the data set consisting of 210 months (1993M03-2008M08), and then we tested the economic series for stationarity in order to eliminate the risk of spurious regression. In order to remove stationarity in the data set; we firstly cleared trend and seasonal, effects from the data and then took the logarithm and first difference of the data. To increase reliability, we used R^2_{adj} criterion instead of R^2 criterion. The variables used in the study were found stationary at the manufacturing industry sub-sector CUR level. Since the other economic series were not stationary, we took their natural logarithms and removed seasonality. However, after this operation, these variables were still not stationary at their levels. For this reason, we took their first differences and obtained the growth rates of interest, export, import and price index; and thus ensured sta-

tionarity for all variables.

After rendering the variables stationary, we estimated CUR through SVM, ANN and VAR. In the estimation, we firstly trained with the data set consisting of 186 months (1991M03-2006M07) and then tested with the data set consisting of 24 months (2006M07-2008M08).

Briefly, in the study, we estimated 21 sectorial CUR using SVM, ANN and VAR methods and evaluated the results for each sector separately with respect to the evaluation criteria of adjusted correlation of determination, mean absolute error and root mean square error under the headings of training and test stages. In the literature review on estimation through SVM, we observed that authors discuss which method is better by demonstrating which method has become successful in how many examples. In this study, along with comparing the successes of the methods in this way, we assessed the general success of the methods through ANOVA for the first time.

In conclusion, in the estimation of CUR values of 21 sub-sectors, we found that SVM yield better results than ANN and VAR in all sectorial and overall evaluations at training and test stages with respect to three different evaluation criteria. Therefore, we conclude that support vector machines can be used successfully in estimation-oriented modeling in general, and in estimating sectorial CUR in particular.

References

- [1] VAPNIK, V.N. *The Nature of Statistical Learning Theory*, Springer. 1995; 0-387-98780-0.
- [2] VAPNIK, V.N. *Estimation of Dependences Based on Empirical Data*, Springer, Berlin. 1982.
- [3] VAPNIK, V.N. and CHERVONENKIS, A. *Theory of pattern recognition*. Nauka, Moscow. 1974.
- [4] VAPNIK, V.N. and CHERVONENKIS, A. On the uniform convergence of relative frequencies of events to their probabilities. *Theory of Probability and its Applications*. 1972; 16(2): 264–281. <http://dx.doi.org/10.1137/1116025>
- [5] VAPNIK, V.N. and CHERVONENKIS A. A note on one class of perceptrons. *Automation and Remote Control*. 1964; 25: 838–845.
- [6] VAPNIK, V.N. and LERNER, A. Pattern recognition using generalized portrait method, *Automation and Remote Control*. 1963; 24: 774–780.
- [7] VAPNIK, V.N., GOLOWICH, S. and SMOLA, A. Support vector method for function approximation, regression estimation, and signal processing, *Advances in Neural Information Processing Systems 9*, In: Mozer M.C., Jordan M.I., and Petsche T. (Eds.) MA, MIT Press, Cambridge. pp. 281–287. 1997.
- [8] BERNDT, E. R. and MORRISON, C. J. Capacity utilization measures: underlying economic theory and an alternative approach. *American Economic Review*. 1981; 71(2): 48–52.
- [9] CORRADO, C. and MATTEY, J. Capacity utilization. *Journal of Economic Perspectives*. 1997; 11(1): 151–167. <http://dx.doi.org/10.1257/jep.11.1.151>
- [10] KIM, H.Y. Economic capacity utilization and its determinants: Theory and evidence. *Review of Industrial Organization*. 1999; 15(4): 321–339. <http://dx.doi.org/10.1023/A:1007747813591>
- [11] BURNSIDE, C., EICHENBAUM, M. and REBELO S. Capital utilization and returns to scale. NBER Working Paper Series, no. 5125. 1995.
- [12] HORNSTEIN, A. Towards a theory of capacity utilization: Shift-work and the workweek of capital. *Federal Reserve Bank of Richmond Quarterly*. 2002; 86(2): 65–86.
- [13] AZEEZ, E. A. Economic reforms and industrial performance. An analysis of capacity utilisation in Indian manufacturing. Centre for Development Studies, Working Paper 334, Trivandrum, Kerala, India. 2002.
- [14] SHAPIRO, M.D. Cyclical productivity and the workweek of capital. *American Economic Review, Papers and Proceedings*. 1993; 83(2), 229–233.
- [15] KOENIG, E.F. Capacity utilization as a real-time predictor of manufacturing output. *Economic and Financial Policy Review*. 1996; Q III, 16–23.
- [16] BANSACK, C., MORIN, N. and STARR, M. Technology, capital spending, and capacity utilization. *Economic Inquiry*. 2007; 45(3): 631–645. <http://dx.doi.org/10.1111/j.1465-7295.2007.00019.x>
- [17] CORRADO, C., GILBERT, C., and RADDOK, R. Industrial production and capacity utilization: historical revision and recent developments. *Federal Reserve Bulletin*, Board of Governors of the Federal Reserve System (U.S.). 1997; issue Feb, 67–92.
- [18] BEAUMONT P.M. A production function approach to estimating capacity utilization. Research Paper 8607, West Virginia University, Regional Research Institute. 1985.
- [19] LEE, J.K. Comparative performance of short-run capacity utilization measures. *Economics Letters*. 1995; 48(3–4): 293–300. [http://dx.doi.org/10.1016/0165-1765\(94\)00640-N](http://dx.doi.org/10.1016/0165-1765(94)00640-N)
- [20] RAGAN, J. F. Measuring capacity utilization in manufacturing. *Federal Reserve Bank of New York Quarterly Review*. Winter 1976: 13–20.
- [21] DE LEEUW, F. Why capacity utilization rates differ. *Staff Studies 105*, Board of Governors of the Federal Reserve System (U.S.). 1979.
- [22] SHAPIRO, M.D. Assessing the federal reserve's measures of capacity and utilization. *Brookings Papers on Economic Activity*. 1989; 1: 181–242. <http://dx.doi.org/10.2307/2534498>
- [23] LALONDE, R. The United States capacity utilization rate: a new estimation approach. Working Paper 99-14. Bank of Canada. 1999.
- [24] MORIN, N. and STEVENS, J.J. Diverging measures of capacity utilization: an explanation. *Finance and Economics Discussion Series 2004-58*, Board of Governors of the Federal Reserve System (U.S.). 2004.
- [25] RAY, S. C., MUKHERJEE, K. and WU, Y. Direct and indirect measures of capacity utilization: A nonparametric analysis of U.S. manufacturing. *Manchester School*. 2005; 74(4): 526–548. <http://dx.doi.org/10.1111/j.1467-9957.2006.00507.x>
- [26] FÄRE, R., GROSSKOPF, S. and KOKKELENBERG E. C. Measuring plant capacity, utilization and technical change: A nonparametric approach. *International Economic Review*. 1989; 30(3): 655–666. <http://dx.doi.org/10.2307/2526781>
- [27] SAHOO, B.K. and TONE, K. Decomposing capacity utilization in data envelopment analysis: An application to banks in India, *European Journal of Operational Research*. 2009; 195(2): 575–594. <http://dx.doi.org/10.1016/j.ejor.2008.02.017>
- [28] LIEBERMAN, M.B. Capacity utilization: Theoretical models and empirical tests, *European Journal of Operational Research*. 1989; 40(2), 155–168. [http://dx.doi.org/10.1016/0377-2217\(89\)90327-5](http://dx.doi.org/10.1016/0377-2217(89)90327-5)
- [29] KALYUZHNOVA, Y., VAGLIASINDI, M. Capacity utilization of the Kazakhstani firms and the Russian financial crisis: A panel data analysis. *Economic Systems*. 2006; 30(3): 231–248. <http://dx.doi.org/10.1016/j.ecosys.2006.03.001>
- [30] SUN, Y.Y. Adjusting input–output models for capacity utilization in service industries. *Tourism Management*. 2007; 28(6): 1507–1517. <http://dx.doi.org/10.1016/j.tourman.2007.02.015>
- [31] SHAIKH, A. and MOUDUD, J. Measuring capacity utilization in OECD countries: A cointegration method, Working Paper No. 415, The Jerome Levy Economics Institute of Bard College. 2004.
- [32] DERGIADES, T. and TSOUFIDIS, L. A new method for the estimation of capacity utilization: Theory and empirical evidence from 14 EU countries. *Bulletin of Economic Research*. 2007; 59(4): 361–381. <http://dx.doi.org/10.1111/j.0307-3378.2007.00263.x>
- [33] TSOUFIDIS, L. and DERGIADES, T. Time series estimates of capacity utilization based on profits and investment. The case of US Economy. *International Journal of Economics*, 1(1) ISSN: 0973-6719. 2007.
- [34] DERGIADES, T. and TSOUFIDIS, L. Estimating capacity utilization using a SVAR model: An application to the US and Canadian economies. *Economics Bulletin*. 2007; 5(4): 1–12.
- [35] MORRISON, C.J. Productivity measurement with non-static expectations and varying capacity utilization: An integrated approach, *Journal of Econometrics*. 1986; 33(1–2): 51–74. [http://dx.doi.org/10.1016/0304-4076\(86\)90027-8](http://dx.doi.org/10.1016/0304-4076(86)90027-8)
- [36] AL-GHANDOOR, A. and SAMHOURI, M. Electricity consumption in the industrial sector of Jordan: application of multivariate linear regression and adaptive neuro-fuzzy techniques. *Jordan Journal of Mechanical and Industrial Engineering*. 2009; 3(1): 69–76.
- [37] GÖKÇEKÜŞ, O. Trade liberalization and capacity utilization: New evidence from the Turkish rubber industry. *Empirical Economics*. 1999; 23(4): 561–571
- [38] GAJANAN, S., MALHOTRA, D. Measures of capacity utilization and its determinants: a study of Indian manufacturing, *Applied Economics*. 2007; 39(6): 765–776. <http://dx.doi.org/10.1080/00036840500447732>
- [39] MCCULLOCH, W. S. and PITTS, W. A logical calculus of the ideas immanent in nervous. *Activity Bulletin of Mathematical Biophysics*. 1943; 5: 115–133. <http://dx.doi.org/10.1007/BF02478259>
- [40] ROSENBLATT, F. The perceptron: A probabilistic model for information storage and organization in the brain. *Cornell Aeronautical Laboratory, Psychological Review*. 1958; 65(6): 386–408. PMID:13602029 <http://dx.doi.org/10.1037/h0042519>

- [41] KEATING, J.W. Identifying VAR models under rational expectations. *Journal of Monetary Economics*. 1990; 25(3): 453-476. [http://dx.doi.org/10.1016/0304-3932\(90\)90063-A](http://dx.doi.org/10.1016/0304-3932(90)90063-A)
- [42] KHALID, A.M., KAWAI, M. Was financial market contagion the source of economic crisis in Asia?: Evidence using a multivariate VAR model. *Journal of Asian Economics*. 2003; 14(1): 131-156. [http://dx.doi.org/10.1016/S1049-0078\(02\)00243-9](http://dx.doi.org/10.1016/S1049-0078(02)00243-9)
- [43] BLANGIEWICZ, M., CHAREMZA, W.W. East European economic reform: Some simulations on a structural VAR model. *Journal of Policy Modeling*. 2001; 23(2): 147-160. [http://dx.doi.org/10.1016/S0161-8938\(00\)00031-4](http://dx.doi.org/10.1016/S0161-8938(00)00031-4)
- [44] TRAFALIS, T. B. and İNCE, H. Support vector machine for regression and applications to financial forecasting, In: *IJCNN 2000: Proceedings of the IEEE-INNS-ENNS International Joint Conference on Neural Networks: Volume 6* edited by Shun-Ichi Amari, et al., p. 6348-6354. 2000.
- [45] MIKA, S., RÄTSCH, G., WESTON, J., SCHÖLKOPF, B., SMOLA, A. J. and MÜLLER, K.-R. Invariant feature extraction and classification in kernel spaces, *Advances in Neural Information Processing Systems 12*, S. A. Solla and T. K. Leen and K.-R. Müller (Eds.), MIT Press, Cambridge, MA, pp.526-532. 2000.