

ORIGINAL ARTICLE

Drought forecasting using neural network and stochastic models

¹A. Fathabadi, ²H. Gholami, ³A. Salajeghe, ⁴H. Azanivand, ⁵H. Khosravi

¹MS student, watershed engineering, university of Tehran, Iran.

²MS student Of De-Desertification, Faculty of Natural Resources, university of Tehran, Iran.

³Assistant Professor, Faculty of Natural Resources, University of Tehran, Iran

⁴Assistant Professor, Faculty of Natural Resources, University of Tehran, Iran

⁵PhD Student of De-desertification, Faculty of Natural Resources, University of Tehran, Iran. hakhosravi@ut.ac.ir

A. Fathabadi H. Gholami A. Salajeghe H. Azanivand H. Khosravi, Drought forecasting using neural network and stochastic models: *Adv. in Nat. Appl. Sci.*, 3(2): 137-146, 2009.

ABSTRACT

Drought forecasting is necessary to drought management and mitigation. In this study using SPI index, drought condition at five stations in north western of Iran was investigated, and using Artificial Neural Network, time series and K-Nearest Neighbor (KNN) methods SPI values forecasted. At the first, SPI 3, 6, 9 and 12-month time scales, were calculated. Then, using artificial neural network (PML), KNN and time series models, SPI values were forecasted. Results showed that the time series method is superior to the ANN method and KNN in the modeling and forecasting the SPI values and SPI 9, 12 - month time scales were forecasted better than SPI 3, 6- month time scales.

Key words: Drought, SPI, K-nearest neighbor and time series

Introduction

Drought is natural hazard that affecting large regions and causing significant damages both in human lives and natural environment. This phenomenon is defined as temporary and recurring meteorological event, which caused from the lack of precipitation over long term period of time. In order to drought monitoring, hydrologist and meteorologist used drought indices. Drought indices are variables to describe the magnitude, duration, severity, and spatial extent of drought (Mori *et al*, 2006; Mori *et al*, 2007). Most of common indices are calculated using Precipitation data.

The most common Meteorological indices include Palmer Drought Severity Index (PDSI) (Palmer, 1965), Deciles (Gibbs and Maher, 1967), and the Standardized Precipitation Index (SPI) (McKee *et al.*, 1993). Among them the standardized precipitation index SPI is now widely used because of the following advantages.

- 1- SPI is calculated by rainfall alone so that drought assessment is possible even if other metro-hydrological measurements are not available.
- 2- The SPI is not adversely affected by topography.
- 3- The SPI has the ability to quantify the precipitation deficit for multiple time scales.
- 4- SPI index is standardized and can compare dry and wet periods on different locations.
- 5- The SPI detects moisture deficit more rapidly than PDSI

Forecasting methods for meteorological phenomena can be classified as dynamical and statistical methods. Dynamical methods are based on physical rules. The application of these methods needs a lot of data and complicated computational methods. Other methods for forecasting meteorological variable are statistical methods that are based on relation between input and output data and do not consider physical aspect of phenomena (Yosfi *et al*, 2007). Traditionally, statistical models have been used for hydrologic and climatologically drought

forecasting based on time series methods. These models suffer from the assumption linearity and stationery, While most hydrological and Meteorological process are non-stationery and nonlinear. Recently a growing interest of Artificial Intelligent (Artificial Neural network and Fuzzy Logic) in field that relation between input and output are non linear. These techniques are less subjected to the constraints of physical description and able to map the logical input and output relation on basis of observed data set only.

Another method that in the past decade used for forecasting and simulation hydrological and meteorological variables is nonparametric K-Nearest Neighbor approach. This method has successful application in hydrology and meteorology (Lall and Sharma. 1996; Lall *et al.*. 1996; Rajagopalan and Lall., 1999; Buishand; Brandsma., 2001).

Various methods such as Markov chain, theory of run, time series model and neural network are used in literature to characterization the duration and severity of droughts. Some applications are: Mishra *et al* (2007) using the alternative renewable process and run theory, investigated the distribution of drought interval time, mean drought interval time, joint probability density function and transition probabilities of drought events in the Kansabati River basin in India. Paulo and Pereira (2007), Paulo *et al* (2005) and Steinemann (2003) used Markov chain to modeling (characteristic) SPI and Palmer drought index. Mishra and Desai (2005) used linear time series models ARIMA and (SARIMA) to forecast standardized precipitation index (SPI) series in the Kansabati river basin India. Mishra and Desai (2006) compared linear stochastic models (ARIMA/SARIMA), recursive multistep neural network (RMSNN) and direct multi-step neural network (DMSNN) for standardized precipitation index (SPI) series as drought index in the Kansabati River Basin, in India. Moridet *et al* (2007) used Artificial Neural Network to forecasting quantitative values of drought indices SPI and EDI in the Tehran province in Iran. Kim and Valdés (2008) forecasted Palmer Drought Severity Index in the Conchos River Basin at various lead times using conjunction model based on dyadic wavelet transforms and neural networks. Vasiliades and Loukas (2008) used linear deterministic (e.g. regression models) and stochastic time series models (ARIMA, SARIMA models, etc.), and nonlinear models (e.g. artificial neural network models) to forecast droughts using standardized precipitation index (SPI) time series at multiple time scales in the Pinios River Basin in central Greece.

Most drought indices monitor current drought conditions while in order to drought management and mitigation drought forecasting necessary. In this study using SPI index drought conditions at five stations in Iran were investigated, and using Artificial Neural Network, time series and KNN methods SPI values were forecasted.

Material and methods

SPI

The Standardized Precipitation Index (SPI) was developed by McKee *et al.* (1993, 1995). The SPI quantifies the precipitation deficit for multiple time scales. Different time scales reflect lags in the response of different water resources to precipitation anomalies.

The SPI is computed by fitting a probability density function to the frequency distribution of precipitation summed over the time scale of interest. This is performed separately for each month (or any other time scale of the raw precipitation time series) and for each location in space. In this study the SPI is computed as follows: First, gamma probability density function was fitted to the monthly series for 3, 6, 9, 12, and 24 timescale. The cumulative distribution is then transformed using equal probability to a normal distribution with a mean of zero and standard deviation of one. The values of the standard normal variables are actually the SPI values. Drought classifications for SPI values are shown in table (1).

Table 1: Classification of the SPI values and drought category

SPI values	Category
2.00 and above	above Extremely wet
1.50 to 1.99	Very wet
1.00 to 1.49	Moderately wet
-0.99 to 0.99	Near normal
-1.00 to -1.49	Moderately dry
-1.50 to -1.99	Severely dry
-2.00 and less	Extremely dry

K-NN methods

The method of nearest- neighbor is used here to forecast SPI values. There are two versions of the k nearest-neighbor method in literature. The first version (*k*-NN1), refers to nearest-neighbor resembling, in which for, one of the *k*-NN is randomly selected according to some predefined weight-function to simulate time series. The second version (*k*-NN2), referred here as the nearest-neighbor averaging, forecasted time series as a weighted average of all *k*-NN using the same weight-function as in *k*-NN (Brandsma and Konnen, 2006). For implement the K-NN algorithm have the following steps:

1- First define the “successor and composition of feature vector D_t of dimension d

$$\begin{aligned} D_t &: (x_{t-1}) \\ D_t &: (x_{t-1}, x_{t-2}) \\ D_t &: (x_{t-\tau_1}, x_{t-\tau_2}, \dots, x_t - M_1 \tau_1; x_{t-\tau_2}, x_{t-2\tau_2}, \dots, x_t - M_2 \tau_2) \end{aligned} \quad (1)$$

Where, (e.g., 1 month) and (e.g., 12 months) are lag intervals, and M_1 , M_2 are the number of such lags considered in the model. Here we have chosen to use the vector of SPI value on the previous steps (month) as the feature vector.

2- Denote the current feature vector as D_i and determine its k nearest neighbors among the historical state vectors D_m using the weighted Euclidean distance.

$$r_{im} = \left(\sqrt{\sum_{j=1}^d w_j (v_{ij} - v_{mj})^2} \right)^{1/2} \quad (2)$$

Where, v_{ij} is the j^{th} component of D_i and the w_j are weights. Here we choose the weights w_j equal one.
 3- Denote the ordered set of nearest-neighbor indices by J_i, k . An element j (i) of this set records the time t associated with the j th closest D_t to D_i . Denote $x_{j(i)}$ as the successor to $D_j(i)$.
 4- The weights of the k neighbors were based on their rank distance to the value of the target weight function that is calculated as:

$$k(j(i)) = \frac{1/j}{\sum_{j=1}^k 1/j} \quad (3)$$

Where, j is the rank of the hindcast years in ascending order. The weight function assigns weights to each of the k -NN. The neighbor with the shortest distance was assigned the highest weight, whereas the neighbor with the longest distance was assigned the smallest weight.

5- m -step-ahead forecast is obtained by using the corresponding generalized regression estimator.

$$g_{GNN}(x_{i,m}) = \sum_{j=1}^k k(j(i)) x_{j(i),m} \quad (4)$$

Where $x_{i,m}$ and $x_{j(i),m}$ denote the m^{th} successor to i , and $j(i)$, respectively (Lall and Sharma, 1996).

Time series model

Time series model used in data that have serial correlation. Conventional AR, MA, ARMA and ARIMA models had extended application in hydrology. In AR models actual data value in previous steps is used to modeling and forecasting. In MA models error in previous steps used to modeling and forecasting. ARIMA models are special model that used in series are not stationery and by differencing were stationery.

$$\begin{aligned} y_t &= \theta_0 + \phi_1 y_{t-1} + \phi_2 y_{t-2} + \dots + \phi_p y_{t-p} \\ &+ \varepsilon_t - \theta_1 \varepsilon_{t-1} - \theta_2 \varepsilon_{t-2} - \dots - \theta_q \varepsilon_{t-q} \end{aligned} \quad (5)$$

Where, y_t and ε_t are the actual values and random error at time period t , respectively; ($i=1; 2; \dots; p$) and ($j=0; 1; 2; \dots; q$) are model parameters. P and q are integers and often referred to as orders of the model. Random errors are assumed to be independently and identically distributed with a mean of zero and a constant variance of. If $q = 0$, then equation (1) becomes an AR model of order p and when $p = 0$, the model reduces to an MA model of order q .

Box and Jenkins (1976) developed a practical approach to building ARIMA models. The Box–Jenkins methodology includes three iterative steps of model identification, parameter estimation and diagnostic checking. The identification stage involves transforming the data (if necessary) to improve the normality and determine the

differencing required to producing stationary and also determine the order of both the seasonal and non-seasonal AR and MA operators for a given series. By plotting original series (monthly series), seasonality, trends in the mean and variance may be revealed. The nonparametric test (Mann Kendall) can be applied to decide whether trend exists in the yearly data. Autocorrelation function (ACF) and partial autocorrelation function (PACF) should be used to gather information about the seasonal and non-seasonal AR and MA operators for the monthly series. During the estimation stage the model parameters are calculated using the method of moments, least square methods, or maximum likelihood methods. Finally, diagnostic checks of the model are performed to reveal possible model inadequacies and to assist in selecting the best model.

Artificial neural network

ANN is a simulation method that inspired by biological nerve system. The most widely used type of ANN in hydrological modeling is feed-forward multi-layer Preceptor (MLP). It is composed of three layers including input, output and one or more hidden layers which are used to connect input layers to output layers.

In artificial neuron the scalar p and a are input and output from neuron respectively. The scalar input p is transmitted through a connection that multiplies its strength by the scalar weight w , to form the product wp , again a scalar. Here the weighted input wp is the only argument of the transfer function f , which produces the scalar output a . The neuron on the right has a scalar bias, b , the bias as simply being added to the product wp as shown by the summing junction or as shifting the function f to the left by an amount b . The bias is much like a weight, except that it has a constant input of one. The transfer function net input n , again a scalar, is the sum of the weighted input wp and the bias b . This sum is the argument of the transfer function (Demuth and Beale, 1998).

Case study

In this study observed monthly rainfall data from five meteorological stations include Tabriz, Orumiyeh, Sanandaj, Khoy and Zanjan in north-western Iran have been used. The length of data at these stations is from January 1965 to December 2003. The mean annual precipitations of these stations are 288.83, 339.94, 466.27, 299.42, and 306.62 respectively.

Application

The first SPI values with the lead times at 3, 6, 9, and 12 months were calculated. Then using time series method SPI values were predicted. Time series model was developed using Box-Jenkins methodology that includes three steps of model identification, parameter estimation and diagnostic checking.

Before using data in ANN, it is necessary to apply some pre-processing applications to data. Pre-processing can viewed as any application to data that causes the performance of models increased. Because of transfer function in hidden layer used in this study was sigmoid function and slope of this function about zero and one is height and gradually slopes are decreased. Because of have not threshold number zero and one in input data equation (6) used to data normalized data between 0.1 and 0.9. Then training (75%), validation (10%) and test (15%) data set were selected.

$$y = .8 * \frac{X_i - X_{\min}}{X_{\max} - X_{\min}} + .1 \quad (6)$$

Where, Y and X_i represent the original variable and the standardized value respectively, while $X_{\max}$ and $X_{\min}$ are the maximum and the minimum values.

After preprocessing the best structure of network was determined. In ANN the number of input and output variables, the number of hidden layers and the number of neurons in each hidden layer, type of transfer function and training process are structure of network. In ANN gain is construct a relation between input spaces and output space using observed data. In these methods input parameters should be selected such that network can determined any hidden characteristic in input data. Input vector was different combination of antecedence (previous) SPI values and the appropriate input vector has been selected based on trail and error criteria. The input data set are presented in table 2.

Table 2: Various input vector

Input	Name model
S_{t-1}	Model 1
S_{t-1}, S_{t-2}	Model 2
$S_{t-1}, S_{t-2}, S_{t-3}$	Model 3
$S_{t-1}, S_{t-2}, S_{t-3}, S_{t-4}$	Model 4
$S_{t-1}, S_{t-2}, S_{t-3}, S_{t-4}, S_{t-5}$	Model 5

The number of neurons in input and output layer equals to number input and output vector, most important problem in ANN is determined the best number of hidden layer and the number of neurons in hidden layers which for this case Hornik *et al* (1989) showed that network with one hidden layer and sufficiently large number of neurons can approximate any measurable functional relationship between the input and the output variable to any desired accuracy. In this study network with one hidden layer that has sigmoid function in hidden layer and linear function in output layer and number neurons were determined by trail and error procedure were tried. And levenberg –maquaret learning algorithm is used for its speed and effectiveness.

The correlation coefficients (R), Root Mean Square Error (RMSE), and average Absolute Relative Error (AARE) statistics are used as the comparing the application of various models. The RMSE and AARE statistics are denoted as below

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (SPI_t^o - SPI_t^c)^2} \quad (6)$$

$$AARE = \frac{1}{N} \sum_{I=1}^n \left| \frac{(SPI_t^o - SPI_t^c)}{Q_t^o} * 100 \right| \quad (7)$$

Where, Q_t^o is the observed SPI at time t , and Q_t^c is the predicted SPI at time t .

Results

At the first, SPI values were calculated with the lead times of 3, 6, 9, and 12 months. The SPI values for lead times of 3 and 9 months at Khoy station are showed in figures 2 and 3 respectively. It is observed that at all studied stations with increase of timescale, dry and wet period severities have decreased, but the run of dry and wet periods has increased. It is caused because in SPI index calculation, cumulative precipitation quantities of previous months are calculated. So, in short-term timescales total precipitation quantities are less than long-term timescales and every unexpected change in one month precipitation quantities has significant effect on SPI values.

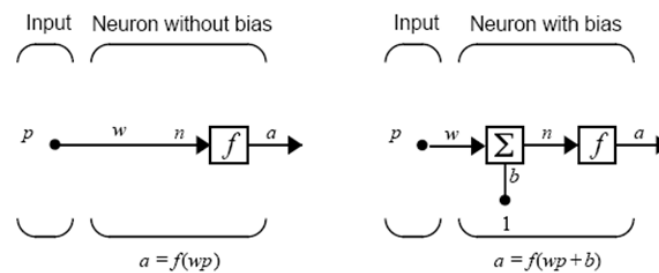


Fig. 1: Schematic view of neuron

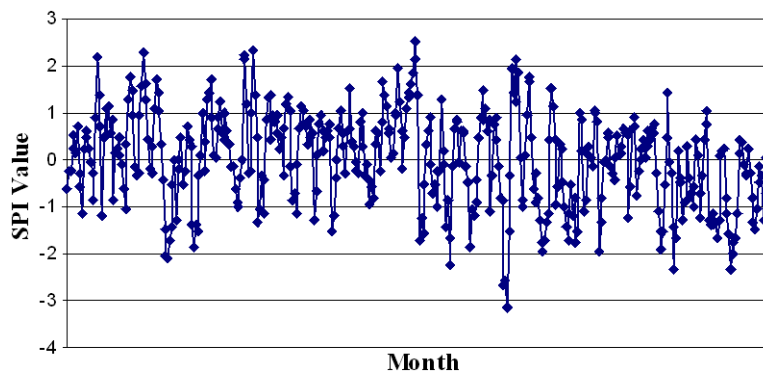


Fig. 2: SPI₃ at Khoy station

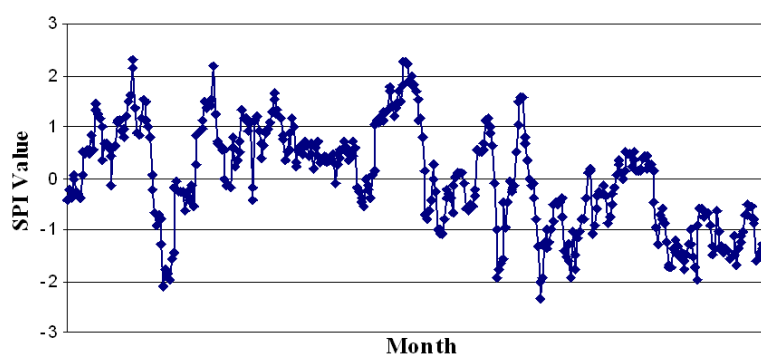


Fig. 3: SPI₉ at Khoy Station

After the calculation of SPI values, SPI values were forecasted by various models. For time series models, ARIMA model (1, 0, 0) (0, 0, 1), that differential order of them were proportional to utilized months number in SPI values calculation was selected as the best model and SPI values were forecasted by this model. The results of forecasted SPI values by time series model are shown in table 3. Table 3 indicates that in all of five stations SPI values with lead time of 9 and 12 months were forecasted better than SPI values with lead time of 3 and 6 months. For Artificial Neural Network, a network with one hidden layer that neurons numbers of hidden layers were 1-16 was tried. Then optimized numbers of hidden neurons were determined in trail and error procedure using statistical criteria of correlation coefficient, RMSE and AARE. Results of the best neural network model are shown in table 4. Table 4 demonstrate that the best model performance happened when neurons number are maximum seven. Also, table 4 shows that input type has less effect on performance of different models. So, models with different inputs almost have similar performance. To forecasting drought, feature vector and neighbor numbers are two main subjects in K-Nearest Neighbor method. For feature vector three input including SPI values in 1, 2, and 12 previous time steps (input1), SPI values in 1, 2 previous time steps (input2) and SPI values in 1 and 12 previous time steps were selected. Results for KNN are presented in table 5. Table 5 shows that type of feature vector has much effect on model performance and the best results related to the state that feature vector of SPI quantities has been in 1 and 2 previous time steps (input 2). According to Lal and Shama (1996) neighbor numbers are appointed equal to Square of total data (18 neighbors) and in the next stage neighbor numbers changes (5, 9, 25, 30 and 42). The results are shown in table 5. In general, models with 30 and 42 neighbors had better performance. In figure 4 observed and predicted SPI values using Nearest Neighbor model at Tabriz station have been shown. Figures 4 and 5 present that with increase of utilized month numbers in SPI values calculation, models performance has increased.

Table 3: Results of SPI forecasting using time series model

Stations		Model	AARE	ARMSE	R
Zanjan	SPI ₃	ARIMA(1,0,0)(1,0,1) ₃	94.19	0.73	0.75
	SPI ₆	ARIMA(1,0,0)(1,0,1) ₆	69.76	0.53	0.89
	SPI ₉	ARIMA(2,0,0)(0,0,1) ₉	269.12	0.37	0.95
	SPI ₁₂	ARIMA(1,0,0)(0,0,1) ₁₂	43.31	0.34	0.96
Orumiyeh	SPI ₃	ARIMA(1,0,0)(0,0,1) ₃	120.07	0.66	0.75
	SPI ₆	ARIMA(1,0,0)(1,0,1) ₆	97.04	0.44	0.92
	SPI ₉	ARIMA(1,0,0)(0,0,1) ₉	62.47	0.3	0.97
	SPI ₁₂	ARIMA(1,0,0)(1,0,1) ₁₂	28.07	0.26	0.98
Tabriz	SPI ₃	ARIMA(1,0,0)(0,0,1) ₃	128.17	0.54	0.7
	SPI ₆	ARIMA(1,0,0)(0,0,1) ₆	95.54	0.37	0.88
	SPI ₉	ARIMA(1,0,0)(0,0,1) ₉	71.39	0.25	0.95
	SPI ₁₂	ARIMA(2,0,0)(0,0,1) ₁₂	25.69	0.22	0.97
Sanandaj	SPI ₃	ARIMA(1,0,0)(1,0,1) ₃	179.01	0.75	0.74
	SPI ₆	ARIMA(2,0,0)(0,0,1) ₆	107.55	0.54	0.9
	SPI ₉	ARIMA(1,0,0)(0,0,1) ₉	61.84	0.39	0.95
	SPI ₁₂	ARIMA(1,0,0)(0,0,1) ₁₂	44.95	0.28	0.98
Khoy	SPI ₃	ARIMA(2,0,0)(0,0,1) ₃	194.1	0.61	0.74
	SPI ₆	ARIMA(1,0,0)(0,0,1) ₆	73.06	0.43	0.85
	SPI ₉	ARIMA(1,0,0)(0,0,1) ₉	78.97	0.4	0.86
	SPI ₁₂	ARIMA(2,0,0)(0,0,1) ₁₂	30.24	0.28	0.92

Figures 4 and 5 suggest that the time series method is superior to the ANN method and KNN in the modeling and forecasting the SPI values also Neural Network has performed better than KNN. In general, performances of three models were good and three models had forecasted SPI₉ and SPI₁₂ values better than SPI₃ and SPI₆ values.

In drought forecasting, in addition to SPI values forecasting, forecasting classes of every drought and wet season severities are important. So, in this study the performance of each model in forecasting of SPI severity classes has been studied. The results of Tabriz, Sanandaj and Orumiyeh stations are shown in table 6. In

Table 5: Results of SPI forecasting using KNN model

Stations		Neighbor Number	Input model	RMSE	AARE	R
Zanjan	SPI3	42	Input(1)	0.72	86.06	0.79
	SPI6	30	Input(1)	0.85	149.79	0.61
	SPI9	30	Input(2)	0.91	223.45	0.49
	SPI12	18	Input(2)	0.46	49.97	0.93
Orumiyeh	SPI3	42	Input(2)	0.69	145.76	0.72
	SPI6	42	Input(2)	0.87	133.41	0.56
	SPI9	30	Input(2)	0.94	67.21	0.46
	SPI12	30	Input(2)	0.97	56.18	0.41
Tabriz	SPI3	30	Input(2)	0.93	31.04	0.31
	SPI6	42	Input(2)	0.91	77.97	0.35
	SPI9	42	Input(2)	0.81	103.01	0.46
	SPI12	42	Input(2)	0.61	154.98	0.61
Sanandaj	SPI3	25	Input(3)	0.69	158.51	0.83
	SPI6	42	Input(1)	0.86	122.55	0.65
	SPI9	9	Input(3)	0.92	100.59	0.51
	SPI12	30	Input(2)	0.93	31.04	0.31
Khoy	SPI3	25	Input(2)	0.59	209.70	0.75
	SPI6	25	Input(3)	0.78	119.67	0.61
	SPI9	30	Input(2)	0.78	133.88	0.49
	SPI12	42	Input(2)	0.90	35.81	0.33

Table 6: Number of class differences between observed and forecasted SPI values.

	DC	SPI 3			SPI 6			SPI 9			SPI 12		
		ANN	AR	KNN	ANN	AR	KNN	ANN	AR	KNN	ANN	AR	KNN
Tabriz	0	65	65	65	63	62	63	69	72	68	65	69	66
	1	17	16	16	20	22	18	14	12	15	16	15	17
	2	2	3	3	1	0	3	2	1	1	4	1	2
	3	1	1	1	1	1	1	0	0	1	0	0	0
	4	0	0	0	0	0	0	0	0	0	0	0	0
Orumiyeh	0	54	53	56	57	64	50	44	55	46	57	65	56
	1	25	25	21	23	17	27	37	29	34	29	20	24
	2	5	6	7	5	5	8	5	2	6	0	1	6
	3	2	2	2	1	0	1	0	0	0	0	0	0
	4	0	0	0	0	0	0	0	0	0	0	0	0
Sanandaj	0	51	51	48	53	57	51	62	63	62	69	75	57
	1	23	25	26	25	26	26	22	21	22	15	9	28
	2	8	7	8	6	2	6	2	2	2	1	2	0
	3	4	3	4	2	1	3	0	0	0	1	0	1
	4	0	0	0	0	0	0	0	0	0	0	0	0

table 6, number 2 in the DC column (DC stands for differences in classes) means that there are two classes of differences between observed and forecasted values. For example, for SPI₃ in Orumiyeh station from 86 cases in 54 cases SPI severity classes have forecasted precisely, in 25 cases with one class of difference, in five cases with two classes of differences and in two cases with three classes of differences. At all five studied stations maximum difference between observed and forecasted SPI severities classes were three classes. It can be indicated from table 6 time series method is superior to the KNN and ANN in forecasting drought classes.

In figure 5 comparison of observed and forecasted SPI values by time series model and neural network at Tabriz station are shown. Figure 4 and 5 indicate that three models have the same performances in forecasting SPI values with lead time of 3 months, but time series model from SPI₆ to SPI₁₂ has performed better than two other models. It is observed with increase of lead times SPI, graphs are smoother and forecasting is better.

Discussion

In this study SPI values of five stations in northwestern of Iran using time series model, Nearest Neighbor and Artificial Neural Network were forecasted. At all five stations the application of three models in SPI values forecasting were good. Time series model has the best performance. Nearest Neighbor model was the worst model. In every three models SPI values with lead time of 9 and 12 months were forecasted better than SPI values with lead time of 3 and 6 months. To modeling, in time series serial correlation between data is used and to determination of the best model systematically several stage are manipulated. This feature rarely is observed in other models such as Artificial Neural Network that to determination of the best network structure true and error procedure is used. On the other hand in time series model uncertainty in SPI values forecasting is raised by probability but in artificial intelligent model this uncertainty is raised by uncertainty in parameters or maybe process is nonlinear. In Nearest Neighbor model, determination of feature vector and neighbor numbers used in SPI values forecasting are very important. For this aim there is not specific algorithm and must used true and error procedure. If an algorithm for determination of optimum neighbor numbers is found capability of model will be significantly increased. In time series models because of using serial correlation, long-term memory

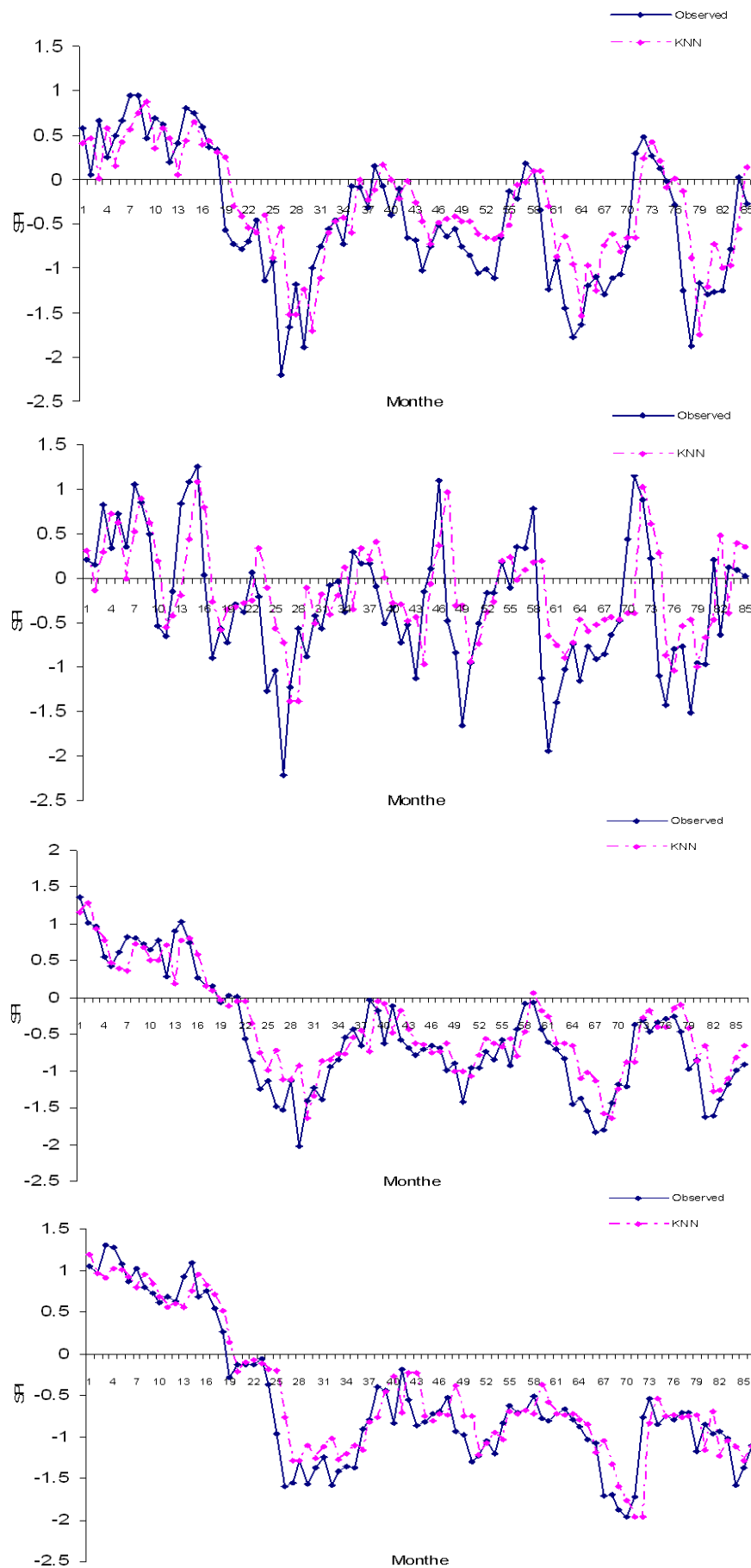


Fig. 4: Comparison of observed and forecasted SPI by KNN at the Tbriz station with lead times of 3 months(a), 6 months (b), 9 months (c) and 12 months (d).

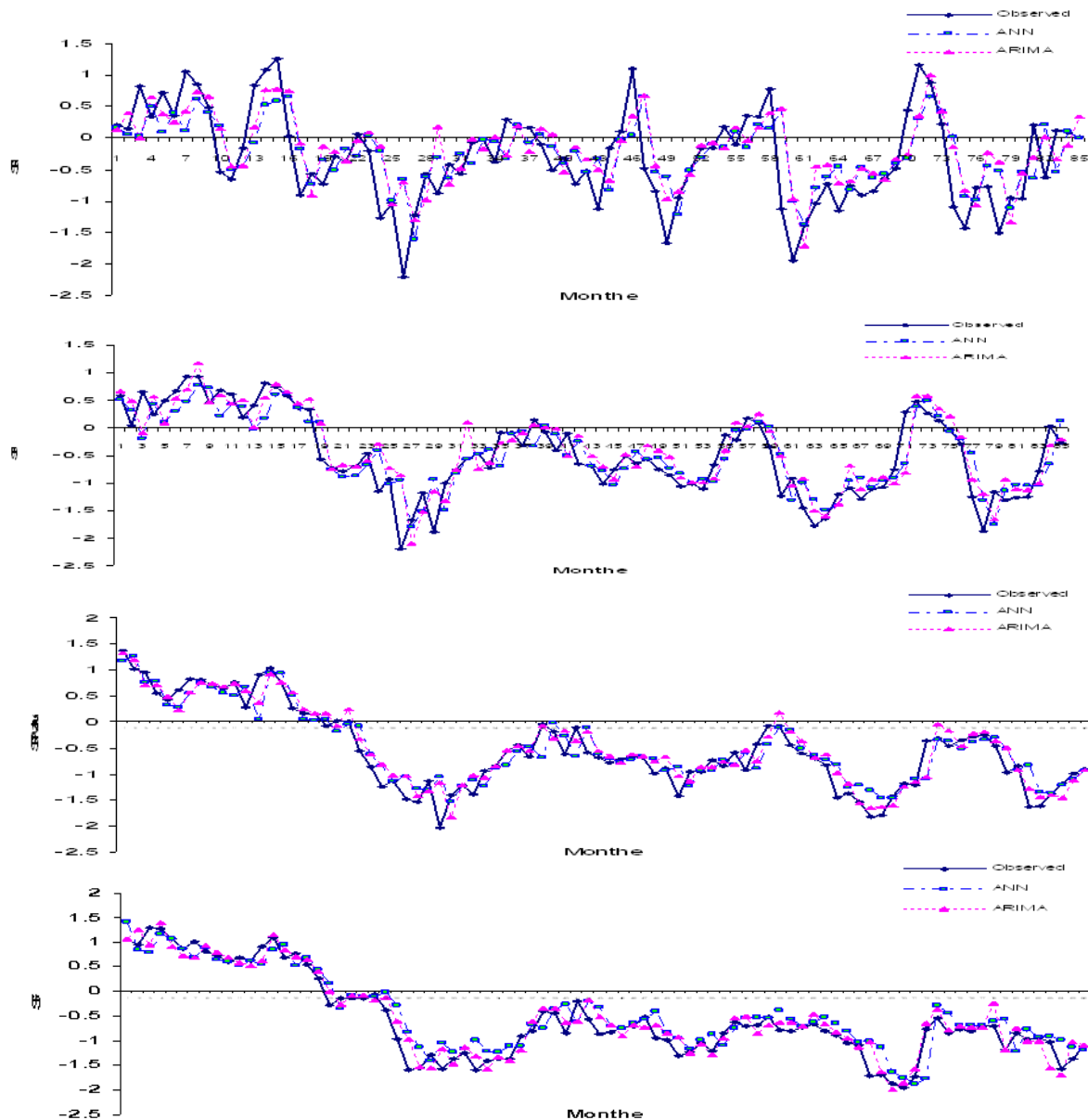


Fig. 5: Comparison of observed and forecasted SPI by ANN and ARIMA at the Tbriz station with lead times of 3 months(a), 6 months (b), 9 months (c) and 12 months (d)

between data is simulated better. This feature is rarely observed in two other models that are a type of nonlinear regression between inputs and outputs. Using from other neural networks such as recursive network that for modeling used from serial correlation and long-term memory in data- or other parameters such as rain, North Atlantic Oscillation and Southern Oscillation Index and other feature inputs cause better model performance to drought and creating nonlinear between data are suggested for future research.

Reference

- Box, G.E.P. and G.M. Jenkins, 1976. Time Series Analysis Forecasting and Control. San Francisco: Holden-Day.
- Brandsma, T. and G.P.K. Konnen, 2006. Application of nearest neighbor resampling for homogenizing temperature records on a daily to sub daily level. *Int. J. Climatol.*, 26: 75-89.
- Buishand, T.A., T. Brandsma, 2001. Multi-site simulation of daily precipitation and temperature in the rhine basin by nearest-neighbor resampling. *Water Resources Research*, 37: 2761-2776.
- Demuth, H. and M. Beale, 1998. Neural Network Toolbox For Use with MATLAB. MathWorks, Inc.
- Hornik, K., M. Stinchcombe and H. White, 1989. Multilayer feedforward networks are universal approximators. *Neural network*, 2(5): 359-366.
- Hayes, M.J., M.D. Svoboda, Wilhite and D.A. Vanyarkho, 1999. Monitoring the 1996 drought using the standardized precipitation index. *Bull Am Meterol Soc.*, 0: 429-438.

- Kim, T.W. and J.B. Valdés, 2003. Nonlinear model for drought forecasting based on a conjunction of wavelet transforms and neural networks. *Journal of Hydrologic Engineering*, Vol. 8, No. 6, November/December, pp: 319-328.
- Lall, U. and A. Sharma, 1996. A nearest neighbor bootstrap for resampling hydrologic time series. *Water Resour. Res.*, 32(3): 679-693.
- McKee, T.B., N.J. Doesken, J. Kleist, 1993. The relationship of drought frequency and duration to time scales. In: 8th Conference on Applied Climatology. American Meteorological Society, Boston, MA, pp: 179-184.
- Mishra, A.K. and V.R. Desai, 2005. Drought forecasting using stochastic models. *Stoch Environ Res Risk Assess* 19: 326-339, DOI 10.1007/s00477-005-0238-4.
- Mishra, A.K. and V.R. Desai, 2006. Drought forecasting using feed-forward recursive neural network. *ecological modelling*, 19(8): 127-138.
- Mishra, A., V.P. Singh and V.R. Desai, 2007. Drought characterization: a probabilistic approach. *Stoch Environ Res Risk Assess* DOI 10.1007/s00477-007-0194-2.
- Morid, S., V. Smakhtin and M. Moghaddasi, 2006. Comparison of seven meteorological indices for drought monitoring in Iran. *International Journal of Climatology*, 26: 971-985.
- Morid, S., V. Smakhtin and K. Bagherzadeh, 2007. Drought forecasting using artificial neural networks and time series of drought indices. *Int. J. Climatol.*, 27: 2103-2111.
- Paulo, A.A. and L.S. Pereira, 2006. Prediction of SPI Drought Class Transitions Using Markov Chains. *Water Resour Manage*, 21: 1813-1827. DOI 10.1007/s11269-006-9129-9.
- Paulo, A.A., E. Ferreira, C. Coelho and L.S. Pereira, 2005. Drought class transition analysis through Markov and Loglinear models, an approach to early warning. *Agricultural Water Management*, 77: 59-81.
- Rajagopalan, B. and U. Lall, 1999. A k-nearest-neighbor simulator for daily precipitation and other weather variables. *Water Resources Research*, 35: 3089-3101.
- Steinemann, A., M.J. Hayes and L. Cavalcanti, 2005. Drought Indicators and Triggers. *Drought and Water Crises: Science, Technology, and Management Issues*. Marcel Dekker, NY, 2005.
- Vasiliades, L. and A. Loukas, 2008. Meteorological drought forecasting models. *Geophysical Research*, Vol. 10, EGU2008-A-08458.
- Yosfi, N., S. Hejam and P. Iranneja, 2007. Estimation wet and drought condition probably using Markov chain and Normal distribution (case study: Ghazvin province). *Geographic Research Journal*. (in Persian).