

Spike Sorting Method Using Exponential Autoregressive Modeling of Action Potentials

Sajjad Farashi

Abstract—Neurons in the nervous system communicate with each other by producing electrical signals called spikes. To investigate the physiological function of nervous system it is essential to study the activity of neurons by detecting and sorting spikes in the recorded signal. In this paper a method is proposed for considering the spike sorting problem which is based on the nonlinear modeling of spikes using exponential autoregressive model. The genetic algorithm is utilized for model parameter estimation. In this regard some selected model coefficients are used as features for sorting purposes. For optimal selection of model coefficients, self-organizing feature map is used. The results show that modeling of spikes with nonlinear autoregressive model outperforms its linear counterpart. Also the extracted features based on the coefficients of exponential autoregressive model are better than wavelet based extracted features and get more compact and well-separated clusters. In the case of spikes different in small-scale structures where principal component analysis fails to get separated clouds in the feature space, the proposed method can obtain well-separated cluster which removes the necessity of applying complex classifiers.

Keywords—Exponential autoregressive model, Neural data, spike sorting, time series modeling.

I. INTRODUCTION

SEQUENTIAL analysis of single-site recordings includes important knowledge about the physiological function of the nervous system. In this purpose placing a single microelectrode in the vicinity of neuron membrane is a common task. The recorded time series by microelectrode contains spikes (neuronal discharges) of the neurons in microelectrode field of view. In such recording, as the microelectrode tip is surrounded by several neurons so the recorded signal contains the activity of more than one neuron therefore the separation of spikes and allocating each to one distinct neuron is a common problem in neural signal analysis. Such task is referred as spike sorting.

The shape of extracellularly recorded spikes depends on the electrical characteristics of the microelectrode, the relative position of microelectrode with respect to the target neurons, and on the electrical properties of neuronal membrane. Therefore in a stationary recording site, the shape of the produced action potentials by each distinct neuron is highly stereotyped. This means that temporal characteristics of spike waveforms can be considered as a useful tool for spike sorting [1], [2]. The sorting procedure classifies extracted spikes in some clusters according to their origination. So each cluster should represent spikes related to one distinct neuron. The

most important step in clustering is the feature extraction where more distinctive features create more compact and well-separated clusters in feature space. This allows simpler classifiers to be applied.

In practice it is difficult to guess the optimal features for classification beforehand. A great deal of efforts in spike sorting is focused on optimal feature selection. Simple features like peak-to-peak amplitude or spike duration aren't optimal and discriminative, especially in the low signal to noise ratios (SNR). On the other hand such features in the cases where spikes have similar amplitude and different shape aren't discriminative [3], [4]. Transformation of spikes into a new space by mathematical tools can reveal some hidden characteristics of time domain representation which can help better discrimination between spikes. In this regard wavelet methods for feature extraction have gained considerable attention. As multiresolution wavelet transform gives a time-frequency representation of each spike, so temporal and frequency content of the spike can be used for finding distinguishable features [5]-[8]. In wavelet domain, any coefficient is the projection of spike on the translated and dilated version of mother wavelet which exhibits a fraction of energy content of the spike in a certain frequency band. Greater similarity between mother wavelet and spike waveform causes sparser transform. This makes mother wavelet selection a critical issue.

Projection of spikes on some first principal components computed from the covariance matrix of spike dataset reduces the dimensionality of spikes to some limited number of coefficients called scores. Each score is the weight of the projection of spike in the direction of one principal component. The scores can be used as features for clustering. Although PCA is a powerful tool in feature extraction but it fails to distinguish spikes different in small-scale structures [7].

There are number of studies focused on the statistical solutions for waveform clustering. An example is the clustering based on the mixture of Gaussian model. This method is based on the assumption of Gaussian distribution of background noise and spike waveforms of each neuron. Such statistical assumption makes it possible to apply Gaussian mixture decomposition and use the model parameters as features [9]. Another choice for spike sorting is methods utilizing neural network. The main weakness of feature extraction methods based on neural networks is their necessity for learning procedure and optimal structure that relies on a-priori knowledge about the data which usually aren't accessible in neural data processing [10]. Some techniques

S Farashi is with the Faculty of Medicine, Shahid Beheshti University of Medical Sciences, Tehran, Iran (e-mail: farashi@sbmu.ac.ir).

have considered sorting of overlapped spikes caused by simultaneously firing of several neurons in the microelectrode neighborhood. The blind source separation methods like independent component analysis (ICA) can be used for removing cross talk between different neurons then an artificial neural network in a supervised or unsupervised fashion can be used for classifying separated source spikes [11]

In the present work a method for optimal feature extraction is proposed which is based on the spike modeling by exponential autoregressive model (EXPAR) combined with self-organizing feature map (SOFM). The method needs no a-priori information about the characteristics of spikes. The paper is organized as the following. In Section II the method for modeling of spikes by an EXPAR model is explained. Also in this section SOFM is used for space dimension reduction. In Section III the performance of EXPAR model and proposed feature extraction method are compared with some traditional methods. The results show that the nonlinear EXPAR model outperforms its linear autoregressive counterpart in modeling action potential. The paper is concluded in Section IV.

II. MATERIAL AND METHODS

A. Modeling Action Potential

The main idea of spike modeling is representation of a time series via some model coefficients. It is useful for signal compression, classification and reconstruction. In this paper mathematical modeling of action potential is considered to express each spike based on its model coefficients. One of the well-recognized models for time series modeling is autoregressive (AR) model. The AR model specifies that the output variable depends linearly on its own previous values. In the case of action potentials, as neurons behave inherently in a nonlinear fashion so it seems that the nonlinear model is more capable to represent action potential [12] so an exponential autoregressive (EXPAR) model as a non-linear tool is used for action potential modeling [13]. An EXPAR model is mathematically expressed as (1):

$$y_t = \sum_{i=1}^p (\varphi_i + \pi_i \exp(-\gamma y_{t-1}^2)) y_{t-p} \quad (1)$$

where p is the model order, γ is a nonlinear coefficient, and finally φ_i and π_i are linear model coefficients. It should be noted that setting all π_i coefficients to zero, EXPAR model is

$$\begin{bmatrix} y(N) \\ y(N-1) \\ \vdots \\ y(4) \\ y(3) \end{bmatrix} = \begin{bmatrix} y(N-1) & y(N-1)e^{-\gamma y^{(N-1)}} \\ y(N-2) & y(N-2)e^{-\gamma y^{(N-2)}} \\ \vdots & \vdots \\ y(3) & y(3)e^{-\gamma y^{(3)}} \\ y(2) & y(2)e^{-\gamma y^{(2)}} \end{bmatrix} \begin{bmatrix} y(N-2) & y(N-2)e^{-\gamma y^{(N-1)}} \\ y(N-3) & y(N-3)e^{-\gamma y^{(N-2)}} \\ \vdots & \vdots \\ y(2) & y(2)e^{-\gamma y^{(3)}} \\ y(1) & y(1)e^{-\gamma y^{(2)}} \end{bmatrix} \begin{bmatrix} \varphi_1 \\ \pi_1 \\ \varphi_2 \\ \pi_2 \end{bmatrix} + e_t \quad (6)$$

In (6), N is the length of data to be modeled and e_t is the residual of the estimation. The concise form of above can be expressed as $\tilde{Y} = B\tilde{\theta} + \tilde{E}$, where $\tilde{\theta}$ is the vector of parameters that should be estimated based on the known values of data

reduced to a simple AR model. For the large values of signal samples, EXPAR model acts like an AR model with coefficients φ_i and for the small values of samples it acts like an AR model with coefficients $\varphi_i + \pi_i$. For the large initial values of signal, the characteristic equation of model is expressed as (2) and for small initial values it is represented by (3).

$$z^p - \varphi_1 z^{p-1} - \dots - \varphi_p = 0 \quad (2)$$

$$z^p - (\varphi_1 + \pi_1)z^{p-1} - \dots - (\varphi_p + \pi_p) = 0 \quad (3)$$

For limit cycle considerations, it is necessary that all roots of (2) be placed inside the unit circle and some roots of (3) be outside the unit circle. It guaranties that all predicted samples by the model in the long term prediction tend to limit cycle. For excluding the unstable singular points, the condition expressed by (4) must be satisfied.

$$\frac{1 - \sum_{j=1}^p \varphi_j}{\sum_{j=1}^p \pi_j} > 1 \text{ or } < 0 \quad (4)$$

Because of the nonlinear nature of model, it is difficult to estimate the coefficients. Here Genetic Algorithm (GA) accompanied by simple least square method is utilized for optimal model parameter estimation. The procedure of model parameter estimation is as following:

Step 1. Predefined range for nonlinear coefficient (γ) is specified: $\gamma \in [a \ b]$ and $\gamma > 0$

Step 2. An initial population for probable values of γ is created. Here a binary genetic algorithm is considered so the created γ values are in the binary string format c with l bit resolution. Each member of this initial population is called chromosome. The string c can be converted to a decimal value by (5).

$$\gamma = a + \frac{c(b-a)}{2^l - 1} \quad (5)$$

Step 3. For each value of γ the problem of parameter estimation is reduced to a linear case and the least square method is used for estimating other parameters (φ_i, π_i). For a second order model other parameters are specified as the following:

samples in matrix B and vector $\tilde{Y}$, and $\tilde{E}$ is the residual vector. The solution for $\tilde{\theta}$ is given by (7).

$$\tilde{\theta} = (B^T B)^{-1} B^T \tilde{Y} \quad (7)$$

Step 4. Knowing the value of γ and estimating other parameters by least square method (step 3), the signal is completely modeled. By Reconstructing data based on the estimated parameters and computing residual mean square value as $\sigma_e^2 = \sum_{p+1}^N e_t^2$, the fitness of estimated parameters is obtained. Here an exponential function of the residual mean square value is used for fitness computing as: $\text{fitness} = \exp(-\sigma_e^2)$.

Step 5. Choosing another chromosome (γ) and repeat steps 3 and 4, the fitness value for all possible solutions (estimated parameters) are computed.

Step 6. The Genetic operator (crossover and mutation) are used to generate new members which called offspring. The new generations are produced as follows:

Cross over: Sort chromosomes based on their fitness values. The most efficient chromosome is located in the top of the list. Select two chromosomes with higher fitness values for cross over. For increasing the chance of chromosomes with smaller fitness values to be included in offspring generation, selection of the chromosomes is done with a probability. For this aim the span of $[2^{n-1}, 2^n]$ is allocated to the n -th chromosome in the sorted population then a random number is generated and the chromosome which the generated number is belong to its span is selected. In this manner the chromosome with higher fitness value has the greater chance to be selected for cross over. For two selected chromosomes which both have l bit resolution, another random number k is generated in the interval $(1:l-1)$ which k is the location where cross over occurs. The bits from k to $l-1$ of the first chromosome are replaced by the corresponding bits in the second chromosome, and vice versa. By cross over, two new offspring are generated. By discarding some chromosomes with lower fitness value from initial populations, the new offspring are added to the end of population list.

Mutation: In each step of offspring generation, the small numbers of chromosomes are selected randomly from the population list and one randomly selected bit of them is reversed for mutation. It should be noted that the probability of mutation is decreased by time (iteration).

Step 7. The fitnesses of generated population members which contain new possible values of γ are computed and the population is sorted based on fitness values again. Follow steps 3 and 4 to obtain the (ϕ_1, π_1) coefficients of each γ coefficient.

Step 8. Steps 3 to 7 are iterated for a large number of predefined iterations. Finally from the population the γ coefficient with largest fitness value and related (ϕ_1, π_1) coefficients are selected as the optimal coefficients for model. The estimated coefficients for a second order model are arranged in a vector like $[\phi_1 \ \pi_1 \ \phi_2 \ \pi_2 \ \gamma]$.

B. Critical Points for Spike Segmentation

The most important parameter of EXPAR model is its order. The main approach for model order selection involves selecting a model order that minimizes one or more

information criteria evaluated over a range of model orders. Commonly used information criteria include Akaike Information Criterion (AIC) and the Bayesian Information Criterion (BIC) [14]. Briefly, such criteria are engaged with entropy rate or prediction error of the model and the number of freely estimated parameters or coefficients in the model. Increasing the model order increases the number of parameters. By minimizing both terms, the optimal order for the model is chosen. In the model order selection, data sample size is a critical issue. Also in modeling physiological signals it is important to consider underlying dynamic behavior which is difficult to understand in many physiological cases. By increasing complexity and variation levels in a signal it is essential to increase the model order for reducing prediction or forecasting error. Model with small order faces with increased forecasting error, impairs the frequency resolution of the signal and merge near peaks. Generally in practice it is better to examine more than one criterion to find the best model order. In this paper instead of getting involved in complex mathematical criterion for model order selection, action potential signals are detached to shorter segments in some critical points and in each segment with lower complexity, the parameter estimation is carried out [15]. For finding critical points in action potential time series, the morphological shape of them is considered.

The action potential generation can be divided into five phases including the rising phase, the peak phase, the falling phase, the undershoot phase, and the refractory period which are managed by ionic gates. In rising phase the membrane potential becomes more positive and depolarization occurs. This phase concludes with action potential peak where membrane potential reaches the maximum. After peak location the falling phase starts which the membrane potential falls down. Then an undershoot event which called hyperpolarization is possible. In such phase the membrane potential is more negative than resting potential. Finally, the time during a subsequent action potential is impossible or difficult to fire is called the refractory period, which may overlap with the other phases.

As usually in action potential detection procedure, extracted spikes are aligned with their dominant peak location, so the peak is selected as one critical point. The instance before dominant peak with maximum slope and the instance after dominant peak with maximum absolute slope value are selected as other critical points. Such points are highly depends on the action potential waveform which is important in the process of sorting.

The probable valley location produced by hyperpolarization is another critical point. Due to the activity of other neurons the lower amplitude part of an action potential (initial and tail segments of spike) are contaminated much more by interferences so to reduce the effect of such interferences the initial and tail segments of each spike are discarded in modeling procedure. Discarded segments have duration equal to 0.1 times of action potential length. Selected critical points and produced segments between them for a sample action potential are shown in Fig. 1. Segmentation causes the

variability in each segment be reduced so the lower order model can be used for modeling each segment. Note that the slow variation in each segment can be modeled by exponential behavior of EXPAR model.

It is possible that interferences caused by the activity of neighboring neurons destroy the hyperpolarization valley. In such case associated critical point of the hyperpolarization valley will be chosen in the mid distance between critical point with maximum absolute slope and starting point of the tail segment. As the amplitude, duration and shape of the action potential are determined largely by the properties of the excitable membrane and not the amplitude or duration of the stimulus so it is expected that the model coefficients for each segment of action potentials of one neuron are as close as possible to each other. EXPAR model with order p contains $2p+1$ free coefficients. Here the initial and tail segments of action potential are discarded for modeling purposes and other 5 segments are considered for modeling. In this manner each action potential can be represented by $5 \times (2p+1)$ coefficients in a vector representation. After representing all spikes by the vector of coefficients, the limited numbers of estimated coefficients are selected as final features. For this aim SOFM is utilized.

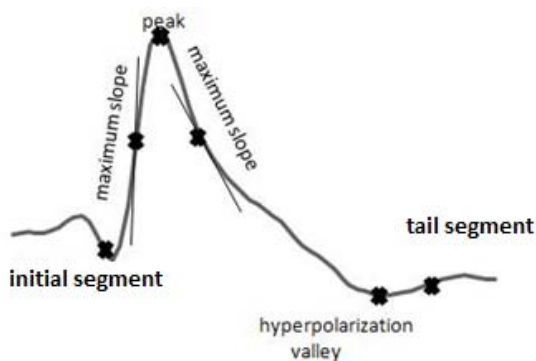


Fig. 1 Sample action potential and selected critical points. Critical points divide each spike to seven segments. The initial and tail (last) segments are discarded in modeling procedure and other segments are modeled by EXPAR model

C. Self-Organizing Feature Map (SOFM)

SOFM is an unsupervised learning artificial neural network (ANN) which is used to produce a low-dimensional representation of the input space of the training samples, called map. SOFM may be considered as a nonlinear generalization of PCA [11]. The map contains a set of nodes or neurons with specified weight which are arranged in a 2D rectangular or hexagonal grid. Associated weight of each neuron has the same length as training data. SOFM consist of two different phases including training and mapping. In training step a high dimensional input training set is used for adjusting the weight and location of map nodes. In a competitive manner for any training data, node with closest weight to the training data is selected as winner node or best matching unit (BMU). The weight and location of the BMU and its neighbor neurons are altered toward input data as (8).

$$w(i+1) = w(i) + \alpha(i)\theta(u,v,i)(D(t) - w(i)) \quad (8)$$

where in (8) w is the map node, i is training step, u is the index of the BMU, v is the index of neighboring neurons in the map and D is the input vector or training data. α is a monotonically decreasing learning coefficient and θ is neighborhood function and its value depends on distance between neuron u and v . Usually a Gaussian function is selected as neighborhood function. The update value decreases with time and with distance from the BMU so the BMU is adapted the most and other SOFM nodes farer away are adapted the least. For each input data the update procedure is carried out and eventually the patterns or clusters in training dataset make the nodes a lower representative of dataset[16].

Note that nodes replicate the main templates or patterns in training dataset because in training, similar input vectors tend to excite adjacent nodes in the map therefore similar patterns are mapped close together and dissimilar ones apart.

D. Method

The steps of proposed method in feature extraction for spike sorting procedure are as following:

1. Segment each spike in critical points.
2. Model each segment except the initial and tail segments. These segments due to their low amplitude are likely to be contaminated with interferences caused by neighborhood neurons. The selected order for modeling of each segment is 2 so EXPAR model gives 5 coefficients for each segment as $[\varphi_1 \pi_1 \varphi_2 \pi_2 \gamma]$. Due to the presence of five segments for each spike, finally a vector of 25 parameters will be produced as a new representation of each spike due to 5 selected segments for each spike. In this step the matrix of spike dataset is mapped to a matrix of model coefficients.
3. SOFM is used to obtain the main patterns of coefficient vectors in the matrix of coefficients. A 6 by 6 grid of nodes in a rectangular network is used as SOFM map. Associated weights for each node have the same size as training vector. In training step, training dataset with size of 500 spikes is fed to EXPAR model and a 500 by 25 training matrix is generated. SOFM reduces training matrix to a 36 by 25 matrix which replicate the main templates in training matrix (matrix contains EXPAR coefficients).
4. The final step is referred as feature extraction. As the exact number of classes aren't specified so instead of using produced SOFM map for clustering the weights of map nodes are compared to find two parameters with the maximum separation ability. For achieving better performance these two parameters are searched in two different segments.
5. For clustering a simple linear discriminant analysis (LDA) is used.

III. RESULTS & DISCUSSION

A. Model Performance

The proposed method for spike sorting consists of two distinct phases. In the first phase each spike in dataset is modeled by EXPAR model and hence spike representation is reduced to the vector of model coefficients for all spike segments. In the second phase the main templates of vectors are searched using SOFM. In Table I the performance of EXPAR model is compared with some other available modeling methods which error refers to the model prediction error and is computed based on the sum of squared difference between the main signal and the reconstructed one based on model. The difference is normalized by the energy of the main signal (spike) as (9). The values in Table I are based on the average values of prediction errors in modeling of 100 spike waveforms and the values inside the parenthesis indicate the standard deviation of prediction error.

$$error = \frac{\sum (reconstructed(i) - signal(i))^2}{\sum signal^2(i)} \quad (9)$$

The polynomial model fits each segment of spike with a polynomial in a least square sense. In Table I Global polynomial is a polynomial curve with a predefined degree that is fitted to the spike without segmentation. As can be seen from Table I the polynomial method applied to segmented spike outperforms EXPAR and AR models but when it is applied to the spikes without segmentation (Global polynomial) it gets the worst result in comparison with others. Critical points cause the spikes to be segmented to shorter and smoother parts which can be represented by a lower order model. For Global polynomial, increasing the model order decreases the prediction error. Also comparison between non-linear EXPAR model with its linear counterpart (AR) indicates that the non-linear model enables to model the behavior of a non-linear system like neuron with the lower error.

B. Feature Extraction Performance

From Table I the best performance of EXPAR model is obtained by the order two so all spikes in dataset are represented by coefficients of a second order EXPAR model. For each segment of spike (five segments), excluding the initial and last segments, five coefficients are estimated which produces a 25-length feature vector for each spike. The transformed spikes to model parameter space are depicted in Fig. 2 (a). By SOFM the main templates of feature vectors are extracted as the weight of SOFM map nodes (Fig. 2 (b)). In training phase of SOFM the node weights are adopted to the main templates in the coefficient vectors. In Fig. 3, the 6 by 6 SOFM map is displayed which is obtained by coefficient vectors of training spike dataset consisting of three classes. Note to the existence of three distinct patterns in the map which are shown by closed curves. The weights in each closed curves are similar approximately. Finally two coefficients, among all SOFM weight vectors which discriminate nodes from each other are selected. These selected features are two

parameters of the model which can be selected manually or by normality test [8]. For mapping spike dataset to a two dimensional space the selected model parameters are considered as final features. For our training dataset such parameters are specified by arrows in Fig. 2 (b). Finally LDA classifier is used for clustering purposes. The performance of the proposed feature extraction method is compared with some other methods.

TABLE I
COMPARISON THE ERROR OF EXPAR MODEL AGAINST SOME OTHER MODELS

Modeling method	Modeling error (%)		
	2	3	4
EXPAR	0.078(±0.01)	0.11(±0.01)	0.42(±0.05)
AR	0.224(±0.05)	0.38(±0.05)	0.5(±0.01)
Polynomial	0.028(±0.007)	0.023(±0.003)	0.02(±0.007)
Global polynomial	1.24(±0.1)	0.8437(±0.09)	0.8045(±0.09)

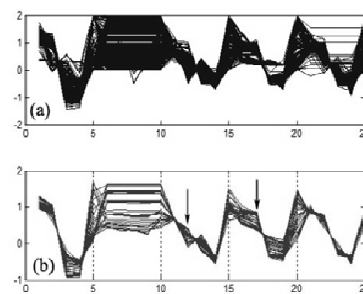


Fig. 2 (a) Feature vectors for spike dataset. Each spike is segmented in its critical points and for each segment the EXPAR model coefficients are estimated. For a second order model there are five parameters estimated for each segment. (b) Feature vector Templates produced by SOFM. Each template is the weight of one node of a 6 by 6 SOFM map. The selected weight members which discriminate the templates are selected by normality test or manually and depicted by arrows

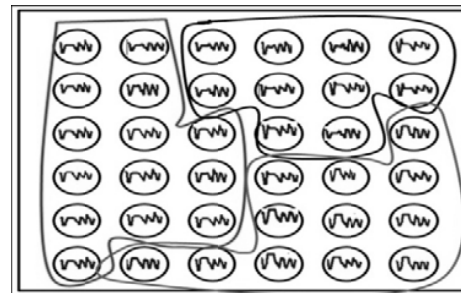


Fig. 3 SOFM map. The map is a 6 by 6 grid of nodes. Each spike is represented by a vector of EXPAR model coefficients. The SOFM decreases the dimensionality of dataset by finding the most probable patterns as the node weights. The similar patterns are assigned to neighbor nodes which are indicated by closed border

For comparison, two other feature extraction methods are implemented. The first is wavelet transform and the second is principal component analysis (PCA). In the wavelet transform based method, multiresolution wavelet transform (MWT) of all action potentials is computed in five decomposition levels.

The coefficients are concatenated into a 1D array starting with the approximation coefficients at the largest level followed by the detail coefficients at all levels in descending order. The normality test for each coefficient among all spikes is carried out and the coefficients with maximum difference from normal probability distribution function (pdf) are searched [8].

Due to the morphology of spike a distinct coefficient is gained for the first and second half of spikes [7]. The two selected coefficients for each spike construct wavelet feature space. Finally LDA classifier is used for clustering extracted features. The wavelet feature selection is done for different mother wavelets. The mother wavelets are selected based on their similarities with templates of spikes in the dataset.

The second method for comparison is PCA which converts the observation of correlated variables into a set of values which linearly uncorrelated. This method finds the directions of variability in the dataset based on the variance as a second order statistics. The directions of variability are sorted where the first direction or principal component (PC) contains the largest variability. After finding PCs the projection of all spikes on the first two PCs are computed and considered as the features for each spikes. Such features also called scores. In this manner spike dataset is transformed to a two dimensional feature space. Again LDA classifier is used for clustering.

For training and testing the proposed feature extraction followed by LDA classifier two spike dataset consist of 500 spikes with known class labels are used. The first dataset consists of three classes which the templates are different waveforms [8] and the second dataset consists of two classes which their templates are different in small-scale structures. Each dataset is decomposed into two equal size data as train and test. The train and test data member Selection is done from dataset in a random manner without embedment. Such train and test data generation is carried out for several times and for each trial the classifier is trained by train data and tested by test data. The final accuracy of classifier is the average value of accuracy for all trials. In Table II the result for classifier accuracy is reported. Even though the PCA feature extraction outperforms our proposed method for data set 1 but for data set 2 our method outperforms PCA. This is due to the weakness of PCA in feature extraction for cases where dataset consist of members different only in details. In such cases the cloud of classes overlaps in feature space.

In Table II the accuracy is defined as the number of true label assignments for test dataset divided by test data size. In Fig. 4 the result of classification for a test dataset1 which consists of three different templates is depicted. The clusters of PCA features and selected EXPAR model parameters are well-separated but the clusters of wavelet coefficients overlap which this increases the LDA classifier error. Also the results for wavelet-based feature extraction methods indicate the sensitivity of such methods to the mother wavelet selection because mother wavelet directly determines how the energy of each spike dispersed in wavelet coefficients and this affects the quality of feature extraction.

The main weakness of PCA in classification arises in the cases where the spikes in dataset are different in small-scale

structures. In such cases PCA gets overlapped cluster. In Fig. 5 the result of feature extraction is displayed for dataset2 which contains two classes of approximately similar spikes. Visually inspection of such result shows that our proposed method get well-separated clouds in feature space in comparison with PCA and wavelet based feature selection methods. In Table II the result of classification for data set2 is reported which indicates the higher ability of proposed method in comparison with PCA and wavelet-based method.

IV. CONCLUSION

Spike sorting is an essential step in studying nervous system mechanism. In this work a sorting method based on the modeling of action potential is proposed. As neurons behave inherently in a nonlinear fashion so their representation should be considered as a nonlinear time series [12]. For such reason an exponential autoregressive (EXPAR) model as a nonlinear model has been utilized. For decreasing the model order, each spike is segmented in some critical points. The critical points are selected based on the physiological properties of action potential. Each segmented part of spike is modeled by a second order EXPAR model and finally the spikes are transformed to feature space by their model coefficients. The results show that segmentation of spikes in critical points reduces the modeling error. Although results show that using a nonlinear AR model outperforms its linear counterpart which it is the result of inherent nonlinear behavior of neuron. For optimal feature selection SOFM is utilized which in its map the weight of neighbor nodes show the main variability in feature space. Among map node elements, two of them which maximize the distance between weight nodes from normal distribution are selected as final features for each spike which obtain a 2D feature space. Results show that such selected features get well-separated clusters in comparison with some traditional feature selection methods. In such situation a simple linear classifier can be utilized for sorting purposes.

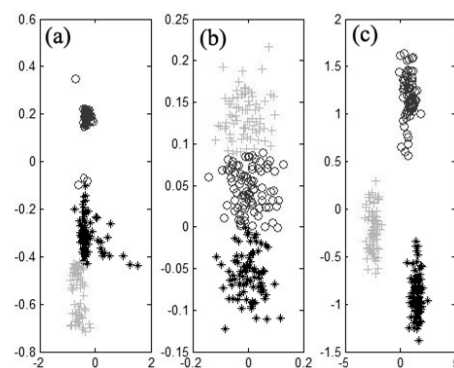


Fig. 4 The result of classification for test dataset 1 for different feature extraction methods (a) Feature selected from EXPAR model coefficients of spikes using SOFM map (b) Features extracted from multiresolution wavelet coefficients selected based on normality test (c) Features extracted from the projection of each spike on first and second principal components. Extracted features are fed to a linear discriminant classifier (LDA)

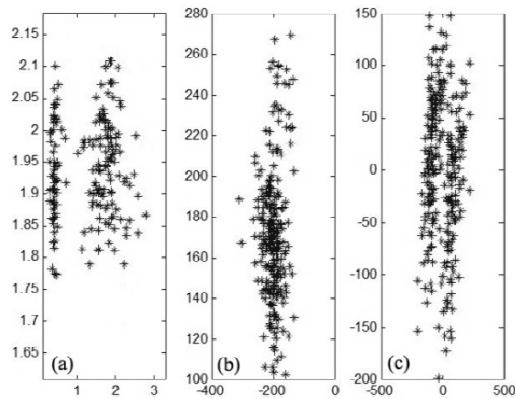


Fig. 5 Result of feature extraction in dataset 2 which contains two classes which are different in small-scale structures (a) Proposed method feature selection (b) Selected wavelet features by normality test (c) Selected features by PCA

TABLE II
COMPARISON OF CLUSTERING ACCURACY FOR DIFFERENT FEATURE
EXTRACTION METHOD

Method	Accuracy (%)	
	dataset1	dataset 2
EXPAR model feature selection	97.24(± 0.5)%	96.9(± 0.5)%
Wavelet feature selection	db2 96.04(± 1.5)%	81.6(± 3.5)%
	Sym2 95.4(± 1.2)%	84.6(± 5.1)%
	db10 78.4(± 4)%	74.4(± 3)%
PCA feature selection	99.6(± 0.1)%	94.3(± 0.05)%

The reported results are averaged values for different test data

REFERENCES

- [1] T. I. Aksenova, O. K. Chibirova, O. A. Dryga, I. V. Tetko, A. L. Benabid, and A. E. Villa, "An unsupervised automatic method for sorting neuronal spike waveforms in awake and freely moving animals," *Methods*, vol. 30, pp. 178-187, 2003.
- [2] Z. Nenadic and J. W. Burdick, "Spike detection using the continuous wavelet transform," *IEEE Trans Biomed Eng*, vol. 52, pp. 74-87, Jan 2005.
- [3] P. H. Thakur, H. Lu, S. S. Hsiao, and K. O. Johnson, "Automated optimal detection and classification of neural action potentials in extracellular recordings," *J. Neurosci. Meth*, vol. 162, pp. 364-376, 2007.
- [4] M. S. Lewicki, "A review of methods for spike sorting: the detection and classification of neural action potentials," *Network-Comp Neural*, vol. 9, pp. 53-78, 1998.
- [5] J. C. Letelier and P. P. Weber, "Spike sorting based on discrete wavelet transform coefficients," *J. Neurosci. Meth*, vol. 101, pp. 93-106, 2000.
- [6] E. Hulata, R. Segev, Y. Shapira, M. Benveniste, and E. Ben-Jacob, "Detection and sorting of neural spikes using wavelet packets," *Phys Rev Lett*, vol. 85, pp. 4637-4640, Nov 20 2000.
- [7] A. Pavlov, V. Makarov, I. Makarova, and F. Panetsos, "Sorting of neural spikes: When wavelet based methods outperform principal component analysis," *Natural Computing*, vol. 6, pp. 269-281, 2007.
- [8] R. Q. Quiroga, Z. Nadasdy, and Y. Ben-Shaul, "Unsupervised spike detection and sorting with wavelets and superparamagnetic clustering," *Neural Comput*, vol. 16, pp. 1661-1687, Aug 2004.
- [9] D. Peel and G. J. McLachlan, "Robust mixture modelling using the t distribution," *Stat Comput*, vol. 10, pp. 339-348, 2000.
- [10] P. M. Horton, A. U. Nicol, K. M. Kendrick, and J. F. Feng, "Spike sorting based upon machine learning algorithms (SOMA)," *J. Neurosci. Meth*, vol. 160, pp. 52-68, 2007.
- [11] T. Hermle, C. Schwarz, and M. Bogdan, "ICA and SOM for spike sorting of multielectrode recordings from CNS," *J Physiol*, vol. 98, pp. 349-356, 2004.
- [12] R. Prenger, M. Wu, S. V. David, and J. L. Gallant, "Nonlinear V1 Responses to Natural Scenes Revealed by Neural Network Analysis," *Neural Networks*, vol. 17, pp. 663-679, 2004.
- [13] R. Baragona, F. Battaglia, and D. Cucina, "A note on estimating autoregressive exponential models," *Quad Stat* vol. 4, pp. 71-88, 2002.
- [14] P. Stoica and P. Babu, "On the proper forms of BIC for model order selection," *IEEE Trans Signal Process*, vol. 60, pp. 4956-4961, 2012.
- [15] A. Ahmadian, S. Karimifard, H. Sadoughi, and M. Abdoli, "An efficient piecewise modeling of ECG signals based on Hermitian basis functions," in *Conf Proc IEEE Eng Med Biol Soc*, 2007, pp. 3180-3183.
- [16] T. Heskes, "Self-organizing maps, vector quantization, and mixture modeling," *IEEE Transactions on Neural Networks*, vol. 12, pp. 1299-1305, 2001.