

Stacked-based *in-silico* ensemble model for toxicity prediction of the chemical data

Priyanka Shit*, Sidhanta Kumar Balabantaray, Smita Ranee Barik
Department of Computer Science and Engineering
Gandhi Institute for Education and Technology, Baniatangi, Bhubaneswar

KAMALINI SINGH, Department of Computer Science and Engineering, Rajdhani Engineering College, Bhubaneswar

Abstract-The stacked-based ensemble method is based on a meta-learning approach that has been effectively used in various fields of application. In the area of computational chemistry, the prediction of toxicity is a very challenging task. The required problem is solved by different methodologies which have been used in the literature. The ensemble classifier is very efficient to use for any classification problem. In this study, we have used a stacked-based ensemble learning method to classify toxic and non-toxic data. For feature reduction of the chemical data, we have used a hybrid filter-wrapper approach. This approach has given an optimized feature set for our classification problems. For comparing the result, we have taken various pre-existing machine learning methods like artificial neural networks, support vector machines, random forest, etc., and we got improved results with a stacked-based ensemble approach.

Keywords: Toxicity, Feature Reduction, Ensemble method, Machine Learning

Introduction

Chemicals are today an essential part of human life and comfort and consequently, a large number of industrial chemicals are manufactured and used every day. However, many of these have been identified as potentially toxic to the humans. The regulatory agencies have been emphasizing their safety assessment prior to their manufacture and use (Casalegno et al., 2005). A 40-h toxicity test, expressed in terms of the median growth inhibition concentration (IGC₅₀) value to the freshwater ciliate *Tetrahymena pyriformis* is considered appropriate for toxicological testing and safety evaluation of the chemicals (Cronin et al., 2002). *T. pyriformis* is an ubiquitous ciliated protozoan belonging to free living, freshwater genus and is sensitive to growth conditions. It has several advantageous characteristics for toxicity studies, such as they play key role in the transfer of energy and matter within the microbial loop. They occupy one of the first trophic levels in aquatic ecosystems and thus, are early warning of a toxic danger (Lukacinova et al., 2007). In *T. pyriformis*, the toxicological end-points have usually been restricted to a measure of growth inhibition rather than of cell viability (Schultz, 1997). The main relevance of this endpoint is to characterize aquatic toxicity. According to Duchowicz and Ocsachoque (2009), the IGC₅₀ encapsulates the most aqueous

toxicity information. Toxicity evaluation of the high production volume compounds, such as industrial chemicals and pesticides is a top global regulatory priority. Accordingly, several experimental databases on IGC₅₀ of different groups of chemicals have been developed by various research groups (Aptula et al., 2005; Cronin and Schultz, 2001; Cronin et al., 2002, 2004; Devillers, 2004; Deweese and Schultz, 2001; Melagraki et al., 2005; Netzeva and Schultz, 2005; Netzeva et al., 2003; Ren, 2003; Ren and Schultz, 2002; Roy et al., 2005; Schultz, 1999; Schultz et al., 2005; Schuurmann et al., 2003; Serra et al., 2001). Because the experimental determination of toxicological properties is a costly and time consuming process, it is essential to develop predictive mathematical relationships to theoretically quantify toxicity (Melagraki et al., 2006). Quantitative structure-activity relationships (QSARs) are increasingly being used as a tool to assess regulatory agencies into toxicological assessment of chemicals substances (Cronin et al., 2002). Desirable qualities of QSAR include the model being transparent, easily portable, and developed with interpretable descriptors. Several quantitative structure-toxicity relationship (QSTR) models, based on multiple linear regression (MLR), partial least squares regression (PLSR), k-nearest neighbor (k-NN), artificial neural networks (ANN), support vector machines (SVM), decision tree (DT) approaches have been proposed by various research groups

(Chenget al., 2011; Ivanciuc, 2004; Jalali-Heravi and Kyani, 2008; Melagraki et al., 2006; Schuurmann et al., 2003; Toropov et al., 2010; Zhu et al., 2008) for predicting *T. pyriformis* toxicity of chemicals. However, the main problem in any QSAR analysis is the evaluation and control of the predictive ability of the developed model (Toropov et al., 2010). Moreover, these local models could provide acceptable predictive accuracy for a very limited chemical domain; they were not applicable to assess a large diverse set of chemical structures (Chenget al., 2011). Also, all these approaches considered different types of molecular descriptors as estimators, and selection and computation of relevant descriptors to extract information from compound structures is the major limitation of this research field.

In recent years ensemble learning (EL) methods (Snelder et al., 2009) have emerged as unbiased tools for modeling the complex relationships between set of independent and dependent variables and have been applied successfully in various research areas (Yang et al., 2010). In general, these methods are designed to overcome problems with weak predictors (Hancock et al., 2005) and have the advantage of alleviating the small sample size problem by averaging and incorporating over multiple classification models to reduce the potential for over-fitting the training data (Dietterich, 2000). Decision tree forest (DTF) and decision tree boost (DTB) implementing bagging and boosting techniques, respectively are relatively new methods for improving the accuracy of a predictive function (Yang et al., 2010). These techniques are inherently non-parametric statistical methods and make no assumption regarding the underlying distribution of the values of predictor variables and can handle numerical data that are highly skewed or multimodal in nature (Mahjoobi and Etemad-Shahidi, 2008). To our knowledge, ensemble learning methods have not yet been applied to the toxicity prediction modeling.

Selection of appropriate molecular descriptors into toxicity prediction is yet another important issue. A large number and variety of such descriptors have been used in several earlier studies, generally derived through highly complicated semi-empirical and empirical methods based on quantum mechanical calculations (Eroglu et al., 2007; Wanget al., 2010; Inetal., 2012). Hence, it would be desirable to develop toxicologically relevant QSTRs using simple properties that can be derived directly from a chemical's structure. Moreover, in view of the regulatory toxicology requirements, models discriminating compounds merely between toxic and non-toxic classes are not enough and it is very much desirable to have more efficient screening tools capable of classifying compounds in several toxicity classes, such as highly toxic, toxic, harmful, and non-harmful; as well as capable of predicting the toxicity end-points in a quantitative manner.

In this study, the basic objectives were to construct the ensemble learning based models (DTB and DTF) for predicting the toxicity of the diverse chemical compounds using simple molecular descriptors. Accordingly, classification and regression models were constructed to predict the toxicity classes and the toxicity end-point ($-\log \text{IC}_{50}$) of the diverse chemicals using a set of selected molecular properties/descriptors as estimators. The predictive and generalization abilities of the DTB and DTF classification and regression models constructed here were evaluated using several statistical criteria. Parameters and performance of these models were tested using external datasets. Moreover, the predictive ability of the DTB and DTF regression models was compared with kernel partial least squares regression (KPLSR), a basic modeling approach.

Materials and methods

Dataset

For developing the ensemble learning based QSTR models for toxicity prediction of chemicals in *T. pyriformis*, data from multiple

sources were considered (Cheng et al., 2011; Tetko et al., 2008; Xue et al., 2006). The chemically heterogeneous dataset is comprised of 1450

chemicals representing different groups. The reported 50% growth-inhibitory concentration values (IGC_{50}) are based on the standard *T. pyriformis* test protocol (Schultz and Netzeva, 2004). For these compounds, values of 40-h exposure based median growth inhibition concentration (IGC_{50}) for *T. pyriformis* are reported as $-\log IGC_{50}$ ($pIGC_{50}, mmolL^{-1}$), where the logarithm is mistakenly contracted the dataset to a computationally efficient range (Serra et al., 2001). Toxicity level generally increases with increasing value of $-\log IGC_{50} (mmolL^{-1})$, and compounds with positive values are generally considered to be toxic or weakly toxic. Complete dataset of 1450 compounds with end-point values are represented in Table S1 (Supporting Information). Since, validation aims to stimulate the predictivity of a model towards new, unknown chemicals, external datasets were collected from literature for model validation (Zhang et al., 2010). Accordingly, the external datasets (Table S2 of Supporting Information) contained comparative toxicity (IGC_{50}) data of chemicals to *T. pyriformis*, and median effective concentration (EC_{50}) values of chemicals in *Vibrio fischeri* (marine bacterium) and *Scenedesmus obliquus* (algae). The EC_{50} values in view of their high correlation reported with IGC_{50} values in different species (Zhang et al., 2010), were also considered for external validation of the constructed models.

Molecular descriptors and feature selection

In toxicological studies, the molecular descriptors represent structural and physicochemical properties of compounds. The descriptors were calculated for each molecule using Toxmatch (Ideaconsult Ltd.). Molecular descriptors (physical, constitutional, geometrical, and topological) were computed by 2D structures of the molecules, which were taken in the form of SMILES (simplified molecular input line entry system). A set of 60 different molecular properties of each of the compound were selected initially. Since, all the molecular properties may not be relevant to the modeling; elimination of less significant descriptors can improve the accuracy of prediction, and facilitate the interpretation of the model through focusing on the most relevant variables. Initial features were selected by model-fitting approach. EL modeling was performed. For optimal values of the model parameters, the EL models were trained by using the complete set of features computing the respective scoring functions to rank the contribution of features in the current set. The lowest ranked features were then removed (Xue et al., 2006). The EL systems were retained by using the remaining set of features, and the corresponding prediction accuracies (misclassification rate, and mean squared error of prediction) were computed by means of 10-fold cross validation. These selected descriptors are logarithmic form of octanol-water partition coefficient ($\log P$), molecular weight (MW), molecular surface area (MSA), charge partial surface area 22 (CPSA-22), connectivity index order one (C100), eccentric connectivity index (ECI), number of atoms in largest chain (NALC), and number of atoms in largest pi-system (NALPS). Finally a set of six descriptors for classification and five descriptors for regression modeling were considered in this study. The basic statistics of these selected descriptors for different datasets are given in Table 1 and Table S3 (Supporting information).

Data processing

Since the aim of present study is to build a robust model capable

of making accurate and reliable predictions of toxicity of new compound, the model derived from a training set should be validated/tested using new chemical moieties for checking its predictive ability. The validation strategies check the reliabilities of models for their possible application on a new dataset, and confidence in the prediction can thus be judged. For predictive modeling the end point (IGC_{50}) values were expressed as $-\log IGC_{50} (pIGC_{50}, mmolL^{-1})$. Categorization of compounds as *T. pyriformis* toxic (TPT) and non-toxic to *T. pyriformis* (Non-TPT) was based on the criteria of Xue et al. (2006). According to the criteria

Table1
 Basic statistics of the selected
 molecular descriptors into toxicity prediction of chemical compounds.

Descriptors	Model	Min	Max	Mean	SD	CV
CPSA.22	C	0.09	1.00	0.43	0.14	31.86
CIOO	R,C,C*	2.56	33.45	12.45	4.48	35.98
ECI	C,C*	2.00	573.00	96.28	57.91	60.15
LogP	R, C,C*	-2.41	6.86	2.08	1.26	60.43
MSA	R	0.00	157.69	38.33	23.79	62.06
MW	R, C,C*	32.03	623.74	153.53	51.47	33.52
NALC	C*	0.00	50.00	11.20	9.54	85.23
NALPS	R,C	0.00	20.00	6.33	4.09	64.54
pIGC ₅₀	R	-2.67	3.34	0.33	1.03	315.02

SD—standard deviation, CV—coefficient of variation; R—regression; C—two-category classification; C*—four-category classification modeling.

(Xue et al., 2006), compounds with pIGC₅₀ value of $\leq -0.5 \text{ mmol L}^{-1}$ are categorized as non-TPT and those with pIGC₅₀ value of $\geq -0.5 \text{ mmol L}^{-1}$ as TPT. In the selected dataset, total 310 compounds were taken as non-TPT and the remaining 1140 compounds were considered as TPT. However, for regulatory purpose, classification of chemicals merely in two categories (TPT, non-TPT) will not be sufficient, and needs a more precise categorization of chemicals. Therefore, a four category classification system of chemicals was considered. Four toxicity classes of the compounds were generated according to the intervals: pIGC₅₀ $\geq 1.5 \text{ mmol L}^{-1}$ for class 1 (highly toxic); $1.5 \text{ mmol L}^{-1} > \text{pIGC}_{50} \geq 0.5 \text{ mmol L}^{-1}$ for class 2 (toxic); $0.5 \text{ mmol L}^{-1} > \text{pIGC}_{50} \geq -0.5 \text{ mmol L}^{-1}$ for class 3 (harmful); $\text{pIGC}_{50} < -0.5 \text{ mmol L}^{-1}$ for class 4 (non-harmful). In this toxicity categorization scheme, the cut-off limit for non-TPT compounds was similar ($\text{pIGC}_{50} \leq -0.5 \text{ mmol L}^{-1}$) to proposed by Xue et al. (2006). These criteria rendered 191 compounds in class I, 458 in class II, 491 in class III, and 310 in class IV. In this study, for classification and regression modeling, data were split into training (80%) and test (20%) subsets using the random distribution approach. Such test sets (when defined prior to analysis) come close to external validation set, which are commonly accepted as the gold standard to assess real predictivity (Benigni et al., 2007).

Diversity analysis

The diversity of a dataset involves defining a diverse subset of representative compounds and is important for global model development (Zhao et al., 2006). The diversity and similarity of the dataset were explored through radar chart analysis (Singh et al., 2013a) and Tanimoto similarity analysis (John et al., 1998). A radar chart displays multivariate data in the form of a two-dimensional chart with several quantitative variables represented on axis starting from the same point. Tanimoto similarity index (TSI), is an appropriate distance metric for topology-based chemical similarity studies. Tanimoto distance method calculates Tanimoto similarity between fingerprint of a chemical and a consensus fingerprint, which is 1024 bit fingerprint (Toxmatch, Ideaconsult Ltd.). The fingerprint generation is based on the fingerprint implementation of the open source cheminformatics library (Steinbeck et al., 2003). For given molecule all possible paths for a predefined length are generated, the path is submitted to a hash function which uses it as a seed to a pseudorandom generator, the hash function outputs a set of bits, which is added to the fingerprint. For a molecule, TSI is calculated as;

$$TSI_{AB} = \frac{1}{4} \frac{Z_{AB} + Z_{AA} + Z_{BB} - Z_{AB}}{Z_{AB}} \quad (1)$$

where Z is the similarity matrix, A and B are the two molecules

being compared. The TSI ranges from 0 (no similarity) to 1 (pair-wise similarity). Smaller TSI means compounds have good diversity (Cheng et al., 2011). A good cut-off for biologically similar molecules is 0.7 or 0.8.

Predictive modeling

Here, we constructed the EL approach based DTB and DTF models for classification and regression problems to predict the toxicity class of compounds as TPT versus non-TPT, and among four categories (highly toxic, toxic, harmful, and non-harmful); as well as the toxicity (pIGC₅₀) of the structurally diverse chemicals using a set of non-quantum mechanical molecular descriptors as the estimators. All computations were performed using the EXCEL 97 and the modeling algorithms were implemented in MATLAB (Mathworks, Inc, Natwick, MA). Brief theoretical description of the EL based modeling approaches used here is provided below.

Ensemble learning approaches

An ensemble contains a number of base learners (Ishwaran and Kogalur, 2010) and their generalization ability is usually much stronger. Bagging and stochastic gradient boosting algorithms are considered here for constructing the classification and regression decision tree (DT) models (DTF and DTB). Bagging uses different perturbed data and feature sets for training base classifiers, whereas in boosting, diversity is obtained by increasing the weight of misclassified samples in an iterative manner (Yanget al., 2010). These methods used decision trees as base classifiers because DTs are sensitive to small changes on the training set (Dietterich, 2000).

Decision Tree Boost. DTB combines the strengths of regression tree and stochastic gradient boosting algorithm. Boosting improves the accuracy of a predictive function by applying it repeatedly in a series and combining the output of each function with weighting, so that the total error of prediction is minimized (Friedman, 2002). Gradient boosting is a sequential stage-wise forward iterative algorithm to find an additive predictor (Fig. 1a). The DTB algorithm creates a tree ensemble and it uses randomization during the tree creations. The goal is to minimize the loss function in the training set, $\{x, y\}$. After each iteration, F represents the sum of all trees built so far:

$$F_m(x) = F_{m-1}(x) + \alpha \cdot \text{Tree}_m(x) \quad (2)$$

where m is the number of trees in the model. Regardless of the loss-function, the trees fitting the gradient on pseudo residuals are regression trees trained to minimize mean squared error (MSE). The DTB model for classification is essentially the same as for regression except $\logit(\text{probability})$ values are fitted rather than raw target values. The DTB uses the Huber M-regression loss function which makes it highly resistant to outliers (Huber, 1964).

Decision tree forest. In DTF, a large number of independent trees are grown in parallel, and they do not interact until after all of them have been built (Fig. 1b). Bootstrap re-sampling method (Efron, 1979) and aggregating, the basis of bagging, are incorporated in DTF. Different training subsets are drawn at random with replacement from the training dataset. Separate models are produced and used to predict the entire data from a fore-said sub-sets. Then various estimated models are aggregated by using the mean for regression problems or majority voting for classification problems. Bagging can reduce variance when combined with the base learner generation with a good performance (Wanget al., 2011). The DTFs gain strength from bagging technique

using the out-of-bag data rows for model validation. This provides an independent test set without requiring a separate dataset or holding back rows from the tree construction. The stochastic element in DTF algorithm makes it highly resistant to over-fitting.

Kernel partial least squares regression (KPLS)

KPLS is a nonlinear extension of linear PLS in which training samples are transformed into a feature space via a nonlinear mapping through the kernel trick, and the PLS algorithm is then implemented in the

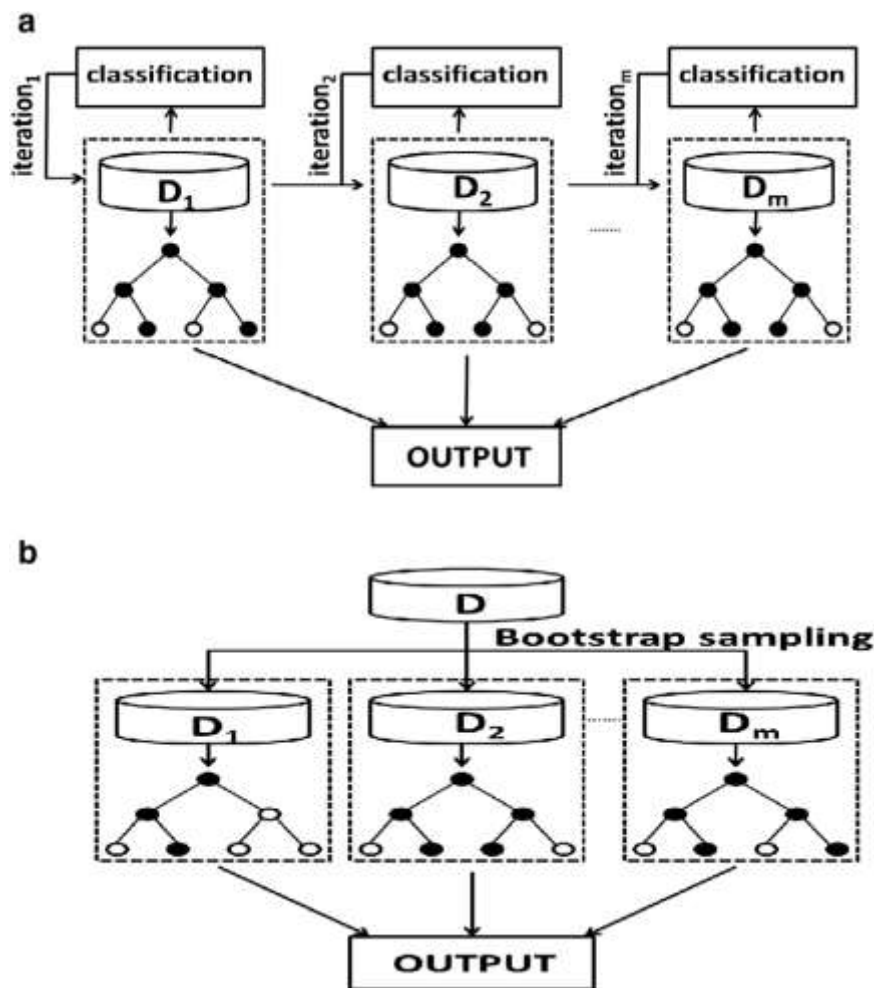


Fig.1. Conceptual diagram of the (a) DTB and (b) DTF models.

feature space. Nonlinear data structure in the original space is most likely to be linear after high-dimensional nonlinear mapping (Singh et al., 2013b). Therefore, KPLS can efficiently compute latent variables (LVs) in the feature space by means of integral operators and nonlinear kernel function. Here, we used the radial basis function (RBF) kernel. For the RBF kernel, the most important parameter is the width (σ) of the RBF which controls the amplitude of the kernel function. Here, the optimum value of σ and LVs were determined through CV procedure.

Model validation

The optimal architectures and model parameters of the classification and regression models constructed here were determined following both the internal and external validation procedures. For internal validation, a V-fold cross validation (CV) method was adopted. The V-fold CV is the most common procedure recommended to check the generalization ability of the model (Benigni et al., 2007). The advantage of this method is that it performs reliable and unbiased testing on data set. For external validation, a separate validation (test) sub-set of the data was used which was kept out during the training process (Singh et al., 2012, 2013a). In case of the predictive models, validation step using external data set provides information about the predictive ability of the trained model for the unknown data (Singh et al., 2012). Benigni et al. (2008) pointed out that the prediction reliability should be checked by means of an external test set with new chemicals not used in model building. Optimal models were selected on the basis of the misclassification rate (classification), MSE and the value of squared correlation

coefficient (R^2) between the measured and model predicted

response (regression) in the training and validation data (Singh et al., 2011).

The developed classification and regression models were further validated using the Y-randomization test. Y-randomization has been frequently used to determine the possibility of chance of correlation during descriptor selection procedure (Xue et al., 2006; Masand et al., 2010). In Y-randomization, the dependent variable (category and pIC_{50}) vectors were randomly shuffled and new models were built using the original independent variables (Mahajan et al., 2013). The procedure was repeated a number of times and CV statistics were computed. The performance (misclassification rate and R^2) of the new models were compared with the original models. If the new models have higher misclassification rate (classification) and lower R^2 values (regression) for several trials, then the given model is thought to be robust. Thus, Y-randomization is useful to avoid any chance-comer correlation between dependent variables vector and independent variables (Mahajan et al., 2013).

Prediction verification

The statistical characterization of classification model is based on the 'confusion' matrix. The performance of the classification models were assessed in terms of misclassification rate (MR), sensitivity, specificity, accuracy of prediction, and Matthew's correlation coefficient (MCC) (Chen et al., 2011; Singh et al., 2011). In addition, the area under curve (AUC) for the receiver operating characteristic (ROC) was calculated. If the AUC of ROC curve is 1, a perfect classifier can be found, and the

AUC value of 0.5 suggests the classifier has no discriminative power at all (Chen et al., 2011).

Performance of the regression models used here was evaluated using different statistical criteria parameters: the root mean square error (RMSE), mean absolute error (MAE), and the squared correlation coefficient (R^2) between the measured and predicted values of the response (Singh et al., 2009a, 2010). Each performance criterion described above conveys specific information regarding the predictive performance efficiency of a specific model. Goodness of fit of the selected model was also checked through the analysis of the residuals.

Results

Basic statistics of the selected molecular descriptors (Table 1) suggest that the standard deviation and the coefficient of variation (CV) indicated high variability within the selected descriptors. The CV values of the descriptors oscillate between 31.86% (CPSA.22) and 85.23% (NALC). It may be noted that the constitutional descriptors exhibited highest variability (64.54%–85.23%) followed by geometrical (31.86%–62.06%), physico-chemical (33.52%–60.43%) and topological (35.98%–60.15%) descriptors. A wide variability in molecular properties of the considered chemicals reveals the importance of the selected descriptors for proposed modeling studies here.

Pattern of the frequency distribution of the experimental toxicity values (pIC_{50}) of the chemicals in training and test sets used for the classification and regression modeling was assessed through constructing the histogram (Fig. 2a). Frequencies of the experimental toxicity values in various sub-ranges (bins) are represented as vertical bars. The histogram shows an nearly normal distribution of the experimental toxicity values for the selected set of chemicals. The radar chart analysis (Fig. 2b) shows that the compounds used in our dataset covered a sufficiently large chemical space. Tanimoto similarity analysis yielded small values of TSI for the toxic (0.011), non-toxic (0.013), and completed data (0.011) (Fig. 2c), indicating structural diversity of chemicals. In Tanimoto similarity assessment, the cut-off values for similar/dissimilar molecules are generally accepted between 0.85 and 0.70 (Manley et al., 2010).

Classification modeling

Classification modeling was performed to categorize the chemicals among two categories (TPT and non-TPT) as well as four categories (highly toxic, toxic, harmful, and non-harmful) of the chemicals. Accordingly, EL-approach based models (DTB and DTF) were constructed for two and four-category classifications using the selected molecular descriptors (Table 1). Optimal architecture and the model parameters were determined through the internal and external validation procedures. Internal validation was performed using the 10-fold CV, whereas, for external validation, a subset of data (20%) was used. The CV results both for two and four-category classifications are given in Table 2. For the TPT versus non-TPT classification, the average (ten runs) values of MR rendered by DTB and DTF models are between 8.90% and 8.97%, whereas, in the four-category classification, the models yielded the average MR values of 25.72% and 26.69%. The results indicate that the toxicity prediction accuracies of both the models are comparable in two and four-category classifications. The results have also shown no obvious over-fitting of data.

The Y-randomization was performed both in two- and four-category classifications of chemicals using 10-fold CV procedure. In two-category classification, the average MR value of these scrambled DTB and DTF models were found to be 16.36% and 16.25%, respectively, whereas in four-category classification the respective models yielded MR

values of 36.58% and 36.16%, which are significantly higher than those of the original DTB and DTF classification systems (Tables 3a, 3b). This suggests that the original classification models are relevant and unlikely to arise as a result of chance of correlation.

Two-category classification

In the selected optimal DTB model, the total number of trees in series, maximum depth of tree, number of average group splits, and shrinkage factor values were 394, 10, 1028.1, and 0.01, respectively. In two-category classification, the DTB model fits the logit values (probability). The shrinkage factor improves the predictive accuracy of a DTB series (Friedman, 2001). On the other hand, the total number of tree in series, maximum depth of a tree, and the number of average group splits in optimal DTF model were 188, 17, and 84.2, respectively. The optimal DTB and DTF models were applied to the test and completed datasets. Contribution of the selected descriptors in classification models ranged between 16.40%–100% (DTB) and 17.85%–100% (DTF). The discriminating descriptors in each model were determined in view of their importance in corresponding model. The importance of the independent variables in each model was determined using the difference in MR calculated using actual data values of all predictors and the observed

puted through randomly rearrangement values of each predictor.

The performance parameters of DTB and DTF models for the training, test and completed data are represented in Table 3a. The MR values yielded by DTB and DTF models were 1.10% and 1.17% in complete data. The overall accuracy of the training and test sets were 100% and 94.48% for DTB, and 100% and 94.14% for DTF, whereas the values of MCC for the respective models in training and test data were 1.0.83 (DTB) and 1.0.83 (DTF). As shown in the Table 3a, the performance of both the models was reasonably good. The results showed that the sensitivity and specificity values for DTB and DTF classification were more than 96% in completed data.

Further, the performance of a classification model can be quantified by calculating the area under the ROC curve (AUC). Both the DTB and DTF models yielded AUC values of 1.0 in training and 0.94 in validation. Moreover, the average gain values in DTB and DTF models ranged between 1.136–1.929 and 1.209–2.150, respectively in validation data.

Four-category classification

The optimal DTB and DTF models with the total number of trees in series (400, 124), maximum depth of tree (10, 21), number of average group splits (5699.5, 241.9), and shrinkage factor (DTB, 0.01) were used for four-category toxicity classification of the chemicals. Contribution of the selected descriptors in four-category classification models ranged between 27.52%–100% (DTB) and 43.43%–100% (DTF). The mean sensitivity, specificity, accuracy and MCC values of different models for the training, test and completed data are summarized in Table 3b. Both the DTB and DTF models yielded the MR value of 3.72% in completed data. The mean accuracy and MCC values in the training and test sets for both the models were 100%, 90.69%, 1.0, and 0.76, respectively. The average gain values in DTB and DTF models ranged between 1.481–1.975 and 1.482–1.989, respectively in validation data. As evident, the performance of the DTB and DTF models constructed here were reasonably good.

Regression modeling

Regression modeling was performed to predict the toxicity (pIGC₅₀) of chemicals using the EL based modeling methods (DTB, DTF). Selected descriptors for regression modeling are given in Table 1. Optimal architecture and the model parameters were determined through the

10-fold CV. A criterion of minimum MSE value (training and validation) was used to determine the optimal model parameters. The average values of MSEs and R² in CV and training data for the proposed regression model are presented in Table 4a. It is evident that the values of MSEs and R² in ten different data folds ranged between 0.155–0.279, and 0.712–0.854 (DTB) and 0.173–0.277, and 0.717–0.835 (DTF),

respectively and their average values 0.220, 0.793 (DTB), and 0.227, 0.786 (DTF). These values are quite comparable to the results obtained when establishing the models in the training (DTB 0.042, 0.963; DTF 0.047, 0.960) and test (DTB 0.128, 0.874; DTF 0.129, 0.876) treatment. Similarly, for the given compounds, predictions are generally very

similar in the ten runs. The Y-randomization results for DTB and DTFre-
gression models derived using 10-fold CV procedure yielded R^2
values of 0.001 in both the cases, which revealed that the original regression

models are relevant and unlikely to arise as a result of chance of correlation.
These results indicate that both the models herein
investigated are robust and showed no over-fitting of data.

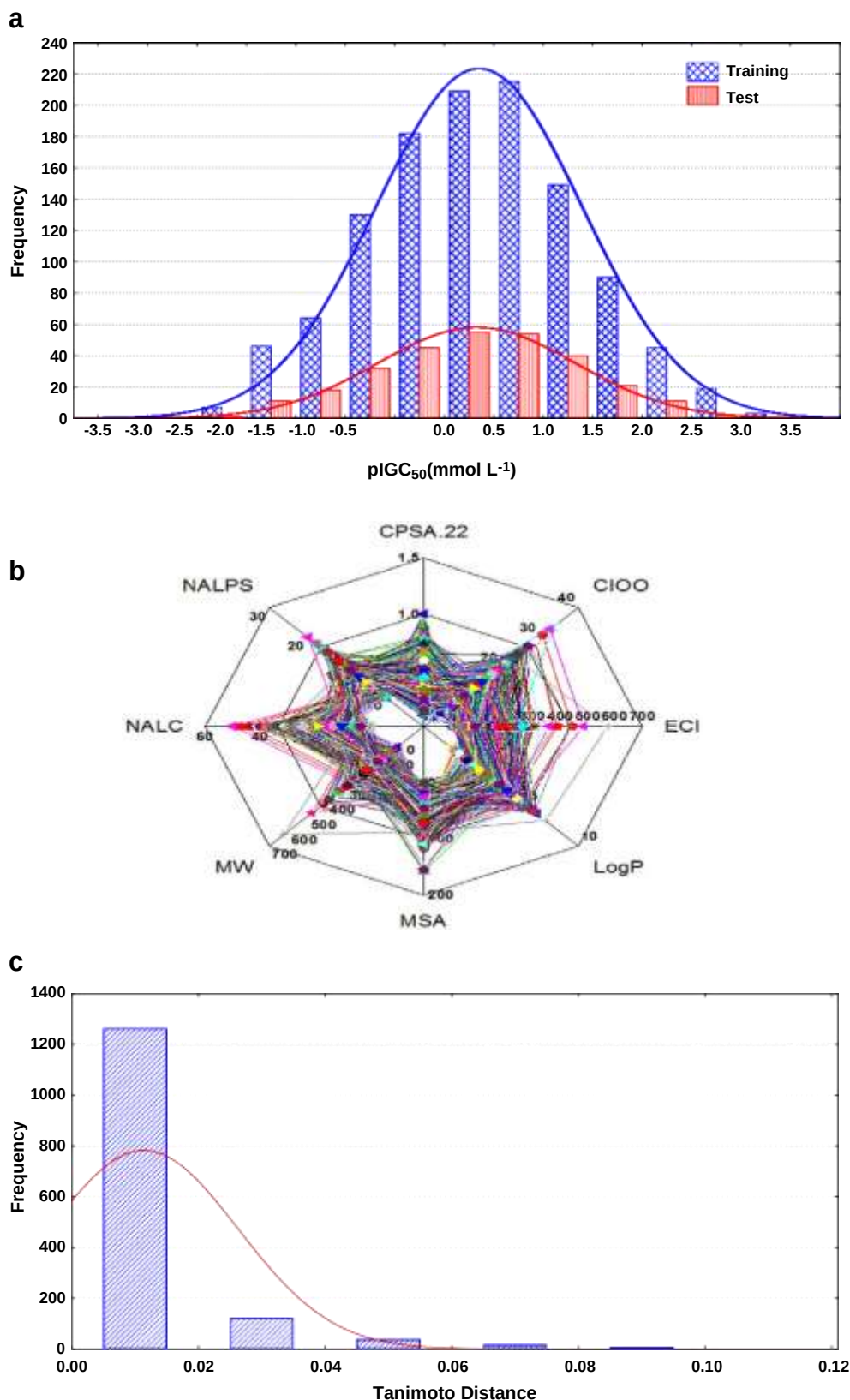


Fig.2. (a) Histogram of the toxicity values (pIC₅₀) of chemicals in *T. pyriformis* in training and test sets, (b) Radar plot of the complete toxicity data (*T. pyriformis*), and (c) Histogram of Tanimoto distance of completed dataset (*T. pyriformis*).

Table2
Misclassificationrate(%)intendifferentfoldsof toxicitydata incrossvalidationof classification(DTB,DTF)models.

Run	2-categoryclassification				4-categoryclassification			
	DTB		DTF		DTB		DTF	
	Training	Validation	Training	Validation	Training	Validation	Training	Validation
1	0.15	7.59	0.15	8.28	0.92	24.14	1.00	26.21
2	0.15	9.66	0.15	11.03	0.77	28.28	0.84	31.03
3	0.08	13.79	0.08	13.79	0.77	26.21	0.77	24.83
4	0.15	8.97	0.15	8.28	1.00	26.21	1.30	27.59
5	0.15	11.72	0.08	13.10	0.84	28.28	1.00	31.03
6	0.15	7.59	0.15	7.59	0.54	27.59	0.61	28.28
7	0.15	6.90	0.15	6.90	0.92	27.59	1.00	26.90
8	0.23	7.59	0.15	7.59	0.69	20.00	2.53	22.07
9	0.08	6.90	0.08	6.21	0.84	24.83	0.84	24.14
10	0.23	8.28	0.15	6.90	0.77	24.14	0.84	24.83
Average	0.15	8.90	0.13	8.97	0.80	25.72	1.07	26.69

The optimal DTB and DTF regression models have the total number of trees in series, maximum depth of tree, and the number of average group splits 435, 190; 10, 21; 973.1, 609.5, respectively. Contribution of the selected descriptors in constructed regression models ranged between 19.36%–100% (DTB) and 24.02%–100% (DTF). The optimal DTB and DTF models were applied to the test and complete datasets. The constructed DTB and DTF models explained 96.05%, 95.60% variance in training, 87.08%, 86.94% variance in test, and 94.37%, 93.98% variance in complete data. Proportion of variance explained by the model variables is the best single measure of how well the predicted values match the actual values. A model predicting exactly matching values with measured ones would explain 100% variance in data. The two models yielded MSE and R^2 values of 0.059, 0.945 (DTB) and 0.064, 0.944 (DTF) in completed data. Values of the performance criteria parameters (R^2 , MAE, and RMSE) yielded by the regression models in training, test and completed data are represented in Table 4b.

Further, a closely followed pattern of variation by the measured and model predicted toxicities of chemicals by the constructed DTB and DTF models in the training and test phases (Fig. 3) suggest that both the models performed reasonably well. Plots of the model-predicted responses (DTB and DTF) and the corresponding residuals for the training and test sets show almost complete independence and random distribution (Fig. 4).

Further, the predictive performance of the proposed EL-based (DTB and DTF) regression models was compared with the KPLS model. In KPLS, the kernel function (RBF) was selected on the basis of R^2 and RMSE values of the validation set. The kernel function parameter (σ) and number of the LVs in the feature space were determined on the basis of the minimum CV error value (0.41). The optimum values of σ and LVs were 0.5 and 5, respectively. The results pertaining to the performance criteria parameters for the KPLS and EL-based regression (DTB, DTF) models are presented in Tables 4a, 4b. The optimal KPLS yielded R^2 and RMSE values of 0.678, 0.59 in training, 0.803, 0.44 in test, and 0.701, 0.56 in complete data. The results suggest that DTB and DTF models performed relatively better than the KPLS in predicting the pIC₅₀ values of the chemicals, which may be attributed to the boosting and bagging algorithms implemented in DTB and DTF approaches.

Testing with external datasets

The constructed DTB and DTF models were also applied to three different external validation datasets collected from literature (Zhanget al., 2010) to further assess their versatility towards using these for predictive purposes. These datasets report experimental toxicities (IG C₅₀) of chemicals in *T. pyriformis* (TSI = 0.26; n = 15), and EC₅₀ values in *V. fischeri* (TSI = 0.04; n = 96), a marine bacterium, and *S. obliquue* (TSI = 0.17; n = 37), an algae. The *T. pyriformis* data was used both for classification and regression, whereas the other two datasets were used for regression alone. Both the proposed DTB and DTF models yielded 100% accuracy of classification (*T. pyriformis*) in two-category scheme. The classification results suggest that the proposed discrimination models performed well with external data.

The constructed EL regression models were also applied to the three different toxicity datasets (*T. pyriformis*, *V. fischeri*, *S. obliquue*) with an

Table3a
Classification results for toxicity prediction of chemicals in TPT and non-TPT categories by ensemble learning models.

Model/sub-sets	Class	Total cases	Sensitivity (%)	Specificity (%)	Accuracy (%)	MCC
Training set						
DTB	TPT	910	100.00	100.00	100.00	1.00
	Non-TPT	250	100.00	100.00	100.00	1.00
	Total	1160				
DTF	TPT	910	100.00	100.00	100.00	1.00
	Non-TPT	250	100.00	100.00	100.00	1.00
	Total	1160				
Test set						
DTB	TPT	230	96.52	86.67	94.48	0.83
	Non-TPT	60	86.67	96.52	94.48	0.83
	Total	290				
DTF	TPT	230	96.92	84.13	94.14	0.83
	Non-TPT	60	84.13	96.92	94.14	0.83
	Total	290				
Complete set						
DTB	TPT	1140	99.30	97.42	98.90	0.97

Table3b
Classification results for toxicity prediction of chemicals by ensemble learning models in four toxicity categories.

Model	Total cases	Sensitivity (%)	Specificity (%)	Accuracy (%)	MCC
Training set					
DTB	1160	100.00	100.00	100.00	1.00
DTF	1160	100.00	100.00	100.00	1.00
Test set					
DTB	290	83.00	93.64	90.69	0.76

Dogo Rangsang Research Journal
ISSN : 2347-7180

DTF	Non-TPT	310	97.42	99.30	98.90	0.97
	Total	1450				
	TPT	1140	99.38	96.81	98.83	0.97
	Non-TPT	310	96.81	99.38	98.83	0.97
	Total	1450				

UGC Care Journal
Vol-10 Issue-04 No. 01 April 2020

DTF	290	83.00	93.61	90.69	0.76
<i>Completeset</i>					
DTB	1450	96.46	98.69	98.14	0.95
DTF	1450	96.48	98.70	98.14	0.95

objective to develop inter-species predictive models. The DTB model yielded MSE of 0.02, 0.53, 0.15, and R^2 of 0.975, 0.637 and 0.741 between the experimental and predicted toxicity values in three test species, respectively, whereas in DTF model the respective diagnostic values were: 0.02, 0.51, 0.17 (MSE) and 0.968, 0.655, 0.691 (R^2). These results suggest that both the EL models fit well to the marine bacteria and algal toxicity data and can be used as interspecies model for predicting toxicity in test species.

Applicability domain of the proposed models

The applicability domain (AD) of predictive models was taken into account in order to consider the scope and limitations of the proposed models, i.e. the range of chemical structures for which the models are considered to be applicable (Netzeva et al., 2005). In this study, a two-dimensional descriptor space is considered; thus, the minimum and maximum descriptor values defined a rectangle in one plane. The ranges of implemented descriptors in classification and regression models are presented in Tables 5a, 5b. It is evident that all the compounds in all datasets considered here are within the AD range.

Discussion

In this study, we have developed and rigorously validated EL-based QSAR models for predicting the toxicity of diverse chemicals in *T. pyriformis*. Accordingly, two-category and four-category toxicity criteria based classification and regression models (DTB, DTF) were constructed. The fourth OECD principle requires a suitable measure of performance. Validation of QSAR models is an important aspect for determination of reliability of such models (Tropsha, 2010). There are different approaches of validation including internal and external validation. To measure the goodness-of-fit of a model, the R^2 is calculated between predicted and experimental values. The optimal architecture of the respective models were determined using 10-fold CV procedure and these were further validated using external datasets. For 10 fold CV, the training data were partitioned into 10 folds and iterations of training and validation were performed such that, within each iteration a different fold of data held-out for validation while the remaining 9 folds were used for learning and subsequently the learned models are used to make predictions about the data in the validation fold. Thus, each time, a model was constructed and tested with an unseen dataset. This procedure prevents the overfitting problem (Singh et al., 2011). A close resemblance between the criteria parameters for classification (MR) and regression (MSE) models in CV and model training phases (Tables 2-4) revealed that the constructed models are robust. Models were further validated using 20% of the data as test set, excluded during the model building phase, and predicted by the constructed model (Tables 3a, 3b, 4a, 4b). Y-randomization was performed to determine the probability of chance of correlation during descriptor selection process both in classification and regression modeling. The chance of

Table 4b

Performance parameters for ensemble learning models in toxicity prediction of chemicals.

Model	Sub-sets	Mean	*SD	MAE	RMSE	R^2
Measured	Training	0.33	1.04	—	—	—
	Test	0.31	1.00	—	—	—
	Complete	0.33	1.03	—	—	—
DTB	Training	0.33	0.97	0.14	0.21	0.963
	Test	0.35	0.90	0.27	0.36	0.874
	Complete	0.33	0.96	0.17	0.24	0.945
DTF	Training	0.33	0.95	0.15	0.22	0.960
	Test	0.36	0.88	0.27	0.36	0.876
	Complete	0.34	0.93	0.18	0.25	0.944
KPLS	Training	0.33	0.85	0.43	0.59	0.678
	Test	0.35	0.89	0.35	0.44	0.803
	Complete	0.34	0.86	0.41	0.56	0.701

*SD—standard deviation.

correlation may occur during descriptor selection especially if the number of descriptors is large (Jouan-Rimbaud et al., 1996). In two- and four-category classifications, a portion of each category agents was randomly selected and interchanged keeping the ratio of the classes unchanged. In regression, values of the response variable were randomly selected and scrambled keeping the independent variables unchanged. Process of scrambling of the dataset was repeated 10 times both in classification and regression modeling. The randomization results suggested that original classification and regression models are relevant and unlikely to arise as a result of chance of correlation. The classification and regression models were also validated using the external datasets (*T. pyriformis*, *S. obliquue*, *V. fischeri*). Both the optimal classification models (DTB, DTF) yielded 100% accuracy for the external data (*T. pyriformis*) in two-category mode.

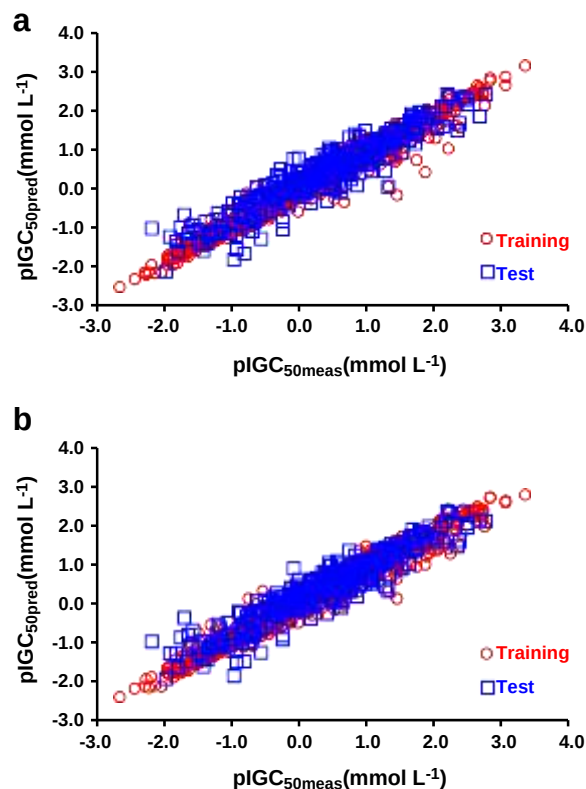


Table 4a

Ten-fold cross validation results for regression models.

Run	DTB-Training	DTB-Validation	DTF-Training	DTF-Validation
-----	--------------	----------------	--------------	----------------

2	0.966	0.037	0.821	0.169	0.960	0.045	0.787	0.202
3	0.968	0.035	0.846	0.155	0.962	0.043	0.828	0.173
4	0.968	0.035	0.776	0.256	0.962	0.042	0.759	0.277
5	0.968	0.034	0.773	0.260	0.964	0.042	0.781	0.255
6	0.968	0.035	0.828	0.218	0.961	0.041	0.835	0.203
7	0.968	0.035	0.854	0.181	0.960	0.044	0.832	0.209
8	0.968	0.035	0.773	0.262	0.962	0.043	0.738	0.261
9	0.968	0.034	0.712	0.279	0.964	0.042	0.717	0.274
10	0.968	0.035	0.810	0.178	0.962	0.043	0.803	0.184
Average	0.968	0.035	0.793	0.220	0.963	0.043	0.786	0.227

Fig. 3. Plot of the measured and model predicted values of the $pIGC_{50}$ for training and testsetsin(a)DTB,and(b) DTFmodel.

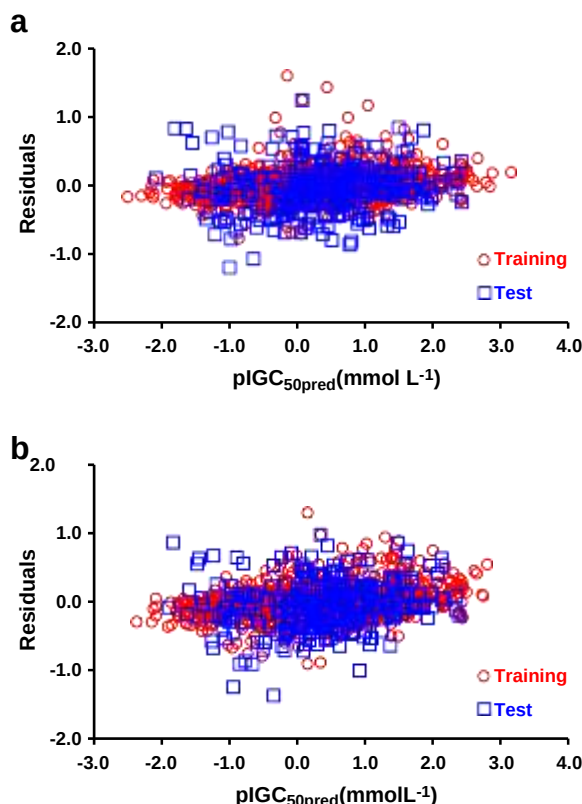


Fig.4. Plot of the model predicted values of the pIC₅₀ and residuals for training and test sets in (a) DTB, and (b) DTF models.

The predictive power of the regression models was evaluated on external validation sets deriving a series of coefficients: R^2 (characterizing the correlation between predicted and experimental values in the validation set) and coefficients Q^2 (Consonni et al., 2009) and r^2 (Roy

et al., 2008). In Q^2 , the denominator is calculated on the training set, and both numerator and denominator are divided by the number of

corresponding elements $\delta Q^2 = \frac{1 - \frac{\sum_{i=1}^{n_{ext}} (y_i - \hat{y}_i)^2}{n_{ext}}}{\sum_{i=1}^{n_{tr}} (y_i - \bar{y}_{tr})^2}$. Consonni

et al. (2009) demonstrated that results obtained by Q^2 are independent of the prediction set distribution and sample size. Recently, r^2_{mtrich} has been proposed as an additional validation parameter (Roy et al., 2008). This metric is calculated based on the correlations between the observed and predicted value with (R^2) and without ($R^2_{intercept}$) for

the least square lines, as $r^2_{mtrich} = \frac{R^2 - R^2_{intercept}}{R^2}$. It does not consider the differences between individual responses and the training set mean, and thus avoids over-estimation of the quality of prediction due to wide

response range (Y-range) (Roy et al., 2009). The results of external validation of the EL-based regression models (DTB, DTF) are presented in Table 6.

The validation threshold criteria values for Q^2 and r^2 are 0.6 and 0.5, respectively. These thresholds of validation measures are applied to all datasets, and then, for every validation criteria, a counting of the accepted models performed. In our case, considerably high correlations (R^2) between the measured and predicted values of response variable both in training and test data were obtained. Model is considered acceptable when the value of R^2 in external set exceeds 0.5 (Ojha et al., 2011). According to the proposed reference criteria (Eriksson et al., 2003), the difference between R^2_{ext} and R^2 should not exceed 0.3. Moreover, R^2 value of ≥ 0.81 for in vitro and ≥ 0.64 for in vivo data can be regarded as good (Kubinyi, 1993). Model validation using the external data yielded criteria parameter (Q^2 , r^2) values (Table 6) well above the

spective thresholds (except for *V. fischeri*). As our models fulfill these criteria and also positively pass internal and external validation, these could be applied to predict the toxicity of new, untested chemicals.

Both the two- and four category classification models (DTB, DTF) yielded high sensitivity, specificity, accuracy and MCC values in training, test and complete data arrays (Table 3a, 3b). Fodorova et al. (2010) suggested that the model for regulatory purposes should be connected with high sensitivity. It may be mentioned that sensitivity is the most important parameter in a classification model. In fact, the low sensitivity value indicates the low ability of a model to recognize the toxicity of diverse compounds. The specificity is another important indicator. High specificity value indicates the high ability of the model to recognize the false positive compounds and it can save the experimental costs (Cheng et al., 2011). Accuracy represents the total number of active and inactive compounds correctly predicted among the total number of tested compounds. MCC value equal to 1 is regarded as a perfect prediction, whereas, 0 is for a completely random prediction. It is evident that both the DTB and DTF models yielded excellent results.

The two-category classification models yielded AUC values close to unity. High AUC values indicate the ranking quality of classification. It can be viewed as a measure based on pair-wise comparisons between classification of two classes and is an estimate of the probability that the classifier ranks a randomly chosen positive example higher than a negative example. With a perfect ranking, all positive examples are ranked higher than the negative ones and AUC = 1. Any deviation from this ranking decreases the AUC, and the expected AUC value for

a random ranking is 0.5 (Fawcett, 2006). In two- and four-category classifications, the constructed models yielded the gain values (in training and validation) ranging between 1.136-2.357, and 1.481-2.707, respectively. The gain values show how much improvement the model provides in picking out the best of the cases. The gain of 1 means non-selective targeting (Abdelaal et al., 2010; Singh et al., 2013a). The performance criteria parameters (sensitivity, specificity, accuracy, MCC), and values of average gain for the classification models suggest that

both these selected models are fully capable to discriminate the chemicals into two and four categories.

Table 5a
Applicability domain of these selected descriptors in two-category and four-category classification models.

Descriptors	Two-categoryclassification				Four-categoryclassification				<i>T.pyrifformis</i> *	
	Training		Test		Training		Test			
	Min	Max	Min	Max	Min	Max	Min	Max	Min	Max
CPSA.22	0.09	1.00	0.14	0.82	—	—	—	—	0.33	0.75
CIOO	2.56	33.45	6.37	21.67	2.56	33.45	5.33	21.67	6.46	10.63
ECI	2.00	573.00	19.00	196.00	2.00	573.00	14.00	196.00	45.00	123.00
LogP	−2.41	6.86	−1.95	4.93	−2.41	6.86	−1.95	4.93	1.84	4.73
MW	32.03	623.74	70.04	390.82	32.03	623.74	55.04	390.82	108.06	263.85

NALC					0.00	50.00	0.00	32.00	0.00	5.00
NALPS	0.00	20.00	0.00	15.00	—	—	—	—	6.00	12.00

*Externaltestset.

Table 5b
Applicability domain of the selected descriptors in regression models.

Descriptors	Trainingset		Testset		<i>T.pyriformis</i> *		<i>S.obliquue</i> *		<i>V.fischri</i> *	
	Min	Max	Min	Max	Min	Max	Min	Max	Min	Max
CIOO	2.56	33.45	2.85	32.67	6.46	10.63	6.46	11.63	2.64	14.57
LogP	-2.41	6.86	-2.01	6.50	1.84	4.73	1.10	4.73	0.19	5.75
MSA	0.00	157.69	0.00	155.39	0.00	86.28	20.23	106.51	0.00	112.30
MW	32.03	623.74	46.04	483.59	108.06	263.85	93.06	326.79	58.04	281.81
NALPS	0.00	20.00	0.00	16.00	6.00	12.00	7.00	13.00	0.00	13.00

*External test set.

An in-depth investigation of the results revealed that the proposed two and four-category EL classification models in completed data misclassified seventeen and fifty four compounds, respectively. These compounds were mainly comprised of phenols, esters, alcohols, aldehydes, acrylates, aromatic amines, and neutral organics. For organic pollutants the mechanism can be distinguished based on the primary processes that cause toxicity (Verhaar et al., 1992). Non-polar narcosis also called baseline toxicity. This type of action results from rather inert chemicals, which act via a non-specific mode of action and finally produce an effect that is characterized as narcosis. The severity of this narcosis type effect and the rate at which the effect occurs depend on the hydrophobicity of the compound. Mainly alcohols (Koleva and Barzilov, 2010), and esters are known to cause non-polar narcosis, whereas phenols and their derivatives cause polar narcosis. Phenols are slightly more toxic than compounds causing non-polar narcosis. The capability of these substances to form H-bond probably enhances their toxicity (Roex et al., 2000). Neutral organic chemicals are non-ionizable and non-reactive and act via a simple non-polar narcosis (Veith and Broderius, 1990).

EL-based regression models (DTB, DTF) were constructed to predict the toxicity (pIGC₅₀) of chemicals. The performance results (Tables 4a, 4b) suggest that both the models yielded considerably low RMSE, MAE and high R² values in training, test and complete data. RMSE is a quadratic scoring rule which measures the average magnitude of the error. It gives a relatively high weight to large errors, hence most useful when large errors are particularly undesirable. MAE measures the average magnitude of the error in a set of predictions, without considering their direction. It is a linear score which means that all the individual differences between predictions and corresponding measured values are weighted equally in the average (Singh et al., 2013a). Plots of the residuals and model predicted values of the response variable are known to provide more information regarding the model fitness to a dataset. A random distribution of the residuals suggests that the model fits the data well, whereas, a non-random distribution shows that the model does not fit the data adequately (Singh et al., 2009b). Fig. 4 shows a random distribution of residuals yielded by both the models in training and test sets, suggesting for the adequacy of constructed models in predicting the toxicity of chemicals.

Table 6
Performance parameters of the ensemble learning models in toxicity prediction using external datasets.

Model	Sub-sets	RMSE	R ²	Q _{F3} ²	r _m ²
DTB	Training	0.21	0.963	—	—
	Test	0.36	0.874	0.881	0.790
	Complete	0.24	0.945	0.945	0.915
	<i>T.pyriformis</i>	0.13	0.975	0.984	0.975
	<i>S.obliquue</i>	0.39	0.741	0.855	0.674
	<i>V.fischri</i>	0.73	0.637	0.501	0.561
DTF	Training	0.22	0.960	—	—
	Test	0.36	0.876	0.879	0.771

An investigation of the prediction results of both the EL models revealed that most poorly predicted compounds, for which the difference between the experimental and predicted values are more than one log unit included the aromatic amines, phenols, and esters, and neutral organics. From a mechanistic point of view, the phenol compounds are commonly associated with the weak acid respiratory uncoupling mechanism of toxic action (Terada, 1990), which is known to exert toxicity in excess of narcosis. These are generally bulky and electronegative compounds produce their toxic effect by disrupting ATP synthesis. They induce disruption of hydrogen ion gradient in the inner mitochondrial membrane (Schultz, 1999). Increased acidity of phenols promotes their ability to uncouple the respiratory chain from oxidative phosphorylation (Schuurmann et al., 2003). These two toxic events are related to the ionization ability and unequal charge distribution profile of a compound.

Selection of most relevant molecular descriptors to the toxicity prediction is important for optimizing the prediction models and for elucidating the molecular factors contributing to toxicity. The molecular structures of the chemicals control their activities and descriptors directly encode particular features of molecular structures. Thus, it is possible to shed light on mechanism of toxic action of the compounds by interpreting the descriptors in the predictive models. The toxicity process are very complicated and involve both toxicant transport to the target site of interaction and interaction between the toxicant and receptors at the sites and which may not be fully described by a few descriptors, particularly for a large number of heterogeneous compounds than those covered in previous studies (Xue et al., 2006). A set of eight molecular descriptors (physico-chemical, topological, geometrical, and constitutional) are selected here. These describe the molecular bulk properties (including transport properties such as membrane permeability) and those representing chemical reactivity of the substance under study. Since the most toxic compounds are characterized by low IGC₅₀ (or high pIGC₅₀) values, all descriptors having positive correlation coefficients (except CPSA and NALC) increase the adverse effect, whereas negative correlation coefficients led to a decreased toxicity effect (Katritzky et al., 2009). From the contributions of the descriptors in classification and regression models here, it can be concluded that log P contributed the most significantly to pIGC₅₀ variation (Fig. 5).

The penetration/solubility descriptors (like log P) reflect the ability of a compound to form non-covalent interactions with its environment, to dissolve and persist in water or in a lipidic environment, or to permeate the phase interfaces. Generally, larger log P indicates a stronger ability of the chemical to permeate the cell membrane of an organism and, therefore, to much more easily interact with its target in the organism (Jiang et al., 2011). The second significant molecular descriptor is the MW, which represents the bulk effect and is correlated with lipophilicity of molecules. Schultz et al. (2007) also reported it as one of the most important descriptors in toxicity models. The topological descriptor treat

the structure of the compound as a graph with atom as vertices and covalent bonds as edges. Both the topological descriptors (C100 and ECI) considered here showed positive correlations with toxicity. Stearic property of a molecule can be described by topological descriptors. It can affect membrane transport and specific interactions at reactive sites with implications to its toxicological profile (Xue et al., 2006). Geometrical descriptors (MSA and CPSA) reflect features of the

molecular geometry. These include distance between particular points of molecular surface (the two farthest points, the two closest points) and distances between given chemical groups. Increasing the MSA and thus revealing the sites of H-bond donation should also increase hydrophobicity. The CPSA has a negative contribution to the total toxicity. More likely this descriptor can be related to the hydrophilicity of the molecule (Katritzky et al., 2003). The constitutional descriptors (NALC and NALPS) account for the steric hindrance effect. The size and shape of compounds influence their transport properties through a biological system as well as their steric hindrance at the reactive site. A larger value of the descriptor indicates the larger steric hindrance (Luan et al., 2005). From the above discussion, it can be seen that the physico-chemical, constitutional, geometrical and topological

properties of the molecules are likely major factors in the process of toxicity, and all of the descriptors involved in the model, which have explicit physical meanings, may account for the structural features responsible for toxicological properties of chemicals.

Defining an appropriate applicability domain (AD) is also very important for the application of QSAR models (Gramatica, 2007). The AD of a predictive model defines the boundaries whereby the predicted values can be trusted with confidence. The AD of a QSAR model is the physico-chemical, structural, or biological space, knowledge or information on which the training set of the model has been developed, and for which it is applicable to make predictions for new compounds. Ideally, the model should only be used to make predictions within that domain by interpolation and not extrapolation (Nikolova-Jeliazkova and

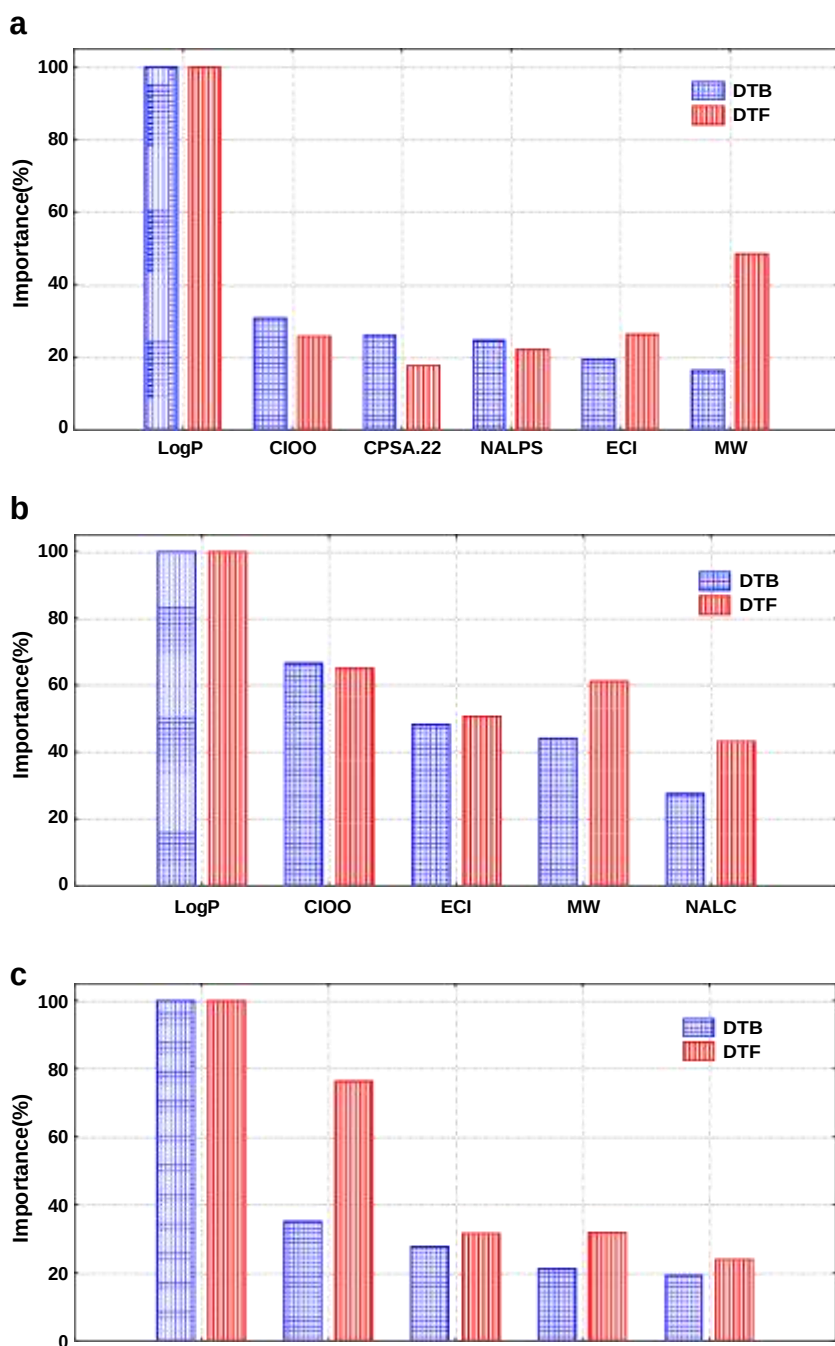


Fig.5. Plot of the contribution of the selected descriptors in toxicity prediction (a) two-category classification, (b) four-category classification, and (c) regression models.

Jaworska, 2005). One of the approaches of defining the AD to estimate the training set coverage in the molecular descriptor space. In mathematical terms, it means estimation of interpolation region in the multivariate space of training set, because the interpolated prediction results are reliable than extrapolated. This approach is especially suitable to those models based on statistical mining techniques (Tan et al., 2010). An interpolation region in one dimensional descriptor space is the interval between the minimum and maximum values of the training dataset. In this study, a two-dimensional descriptor space is considered; thus, the minimum and maximum descriptor values defined a rectangle in one plane. The ranges of implemented descriptors are represented in Tables 5a, 5b. However, this kind of global, chemometric estimation of the AD does not address the chemical space characterized by the descriptor range, so it might happen that the target chemical falls in an area poorly represented in the training set. Moreover, since the AD relies on the chemical descriptors alone, the property under investigation is neglected (Fedorova et al., 2010).

Literature survey shows a number of studies reporting different modeling methods for toxicity prediction in *T. pyriformis*. The results of some research groups on classification and regression modeling for toxicity of chemical compounds are shown in Tables 7a, 7b.

A direct comparison of four results with previous studies is inappropriate, because the number of chemical compounds, nature and number of descriptors, and modeling approaches considered in these studies differ to a large extent. Among various modeling approaches, MLR, PLSR, ANNs, SVMs have commonly been used in QSAR modeling (Tables 7a, 7b). Although, the predictive responses achieved using modeling techniques have been within an acceptable range, these methods have certain limitations. MLR and PLSR are linear methods and do not fit the data with nonlinear structure, a common feature of experimental toxicity data. ANNs, although a universal nonlinear method, it suffers from overfitting in training. SVM uses only a limited data points during model building phase (Singh et al., 2013a). Nevertheless, a simple comparison of the model statistics could provide some basic information about the accuracy of various prediction methodologies. It may be noted that all these studies considered toxicity data of chemicals in single organism (*T. pyriformis*) only. Moreover, most of these studies considered complex and large number of descriptors (including thermodynamic and quantum mechanical) and in several of these, classification and regression accuracies were not satisfactory, thus limiting

the applicability of these models for toxicity prediction in new unknown chemicals. Among these, the present study proposed EL based modeling approaches considering the dataset of structurally diverse chemicals and using smaller number of simple molecular descriptors yielded the best prediction accuracy and correlation for the training, test, complete and external datasets. Proposed EL methods with bagging and stochastic gradient boosting techniques improve the prediction accuracy of weak learners (Breiman, 1996). The bagging minimizes prediction variance by generating bootstrapped replica datasets, whereas, boosting creates a linear combination out of many models, where each new model is dependent on the preceding model (Friedman, 2002). The present study demonstrated applicability of the constructed models (DTB, DTF) in predicting the toxicity of diverse chemicals in marine bacteria and algae and thus, suggesting for the appropriateness of these methods for toxicity prediction of new chemicals in different organisms and can be used as effective tools in risk assessment for regulatory decision making.

Conclusions

In this study, multi-species ensemble learning approach based robust predictive models, such as DTB and DTF were established using large experimental toxicity data containing 1450 structurally diverse compounds, without requiring the knowledge of mechanisms and choices of specific molecular descriptors derived from the chemical's structure. Constructed classification and regression models were validated using internal and external validation procedures, which were directly developed from the molecular descriptors. The statistical results proved that developed DTB and DTF models were efficient algorithms for building two-category and four-category toxicity classification models of TPT prediction. Optimal DTB and DTF models demonstrated excellent predictive abilities in discriminating the TPT and non-TPT, as well as finer toxicity classes of chemicals and can be used as a tool in screening of the new chemicals. The DTB and DTF regression models exhibited excellent ability in predicting the toxicity of the diverse chemicals with the toxicity end points in *T. pyriformis*, marine bacteria (*V. fischeri*), and algae (*S. obliquus*) test species. A notable advantage of the present study may be use of simple non-quantum mechanical descriptors as estimators of the toxicity end-point. Excellent predictive and generalization achieved for proposed EL models may be attributed

Table 7a
Classification accuracies of toxicity models from different studies reported in the literature.

Model	No. of compound	No. of descriptors	Type of descriptors	Accuracy (%)	Reference
SVM	337	3	PD, OH	71.00	Ivanciuc (2004)
SVM	1129	199	CD, PD, GD, CoD, OH	88.90	Xue et al. (2006)
LR	1129	199	CD, PD, GD, CoD, OH	70.10	Xue et al. (2006)
Adaboost	274	9	ED, StD, ThD	92.80 (Test)	Niu et al. (2009)
C4.5	1571	166	MDL	89.20 (Test)	Chen et al. (2011)
k-NN	1571	166	MDL	91.60 (Test)	Chen et al. (2011)
DTB (2-cat)	1450	6	PD, CD, TD, GD	100.00 (Train) 94.48 (Test) 98.90 (Complete) 100.0 (ES)	Present study
DTF (2-Cat)	1450	6	PD, CD, TD, GD	100.00 (Train) 94.14 (Test) 98.83 (Complete) 100.0 (ES)	Present study
DTB	1450	5	PD, CD, TD	100.00 (Train)	Present study

(4-cat)				90.69(Test)	
DTF	1450	6	PD, CD,TD	98.14	
				(Complete)100.0	Presentstudy
(4-Cat)				0Train)	
				90.69(Test)	
				98.14 (Complete)	

PD—Physico-chemicaldescriptor;CD—Constitutionaldescriptors;GD—Geometricaldescriptor;CoD;Connectivitydescriptors;ThD—thermodynamicdescriptor;ED—Electronicdescriptors;StD—Stearic descriptors; TD—topologicaldescriptors;OH—others; ES—T.pyriformisexternal set.

Table7b
Prediction result for toxicity of chemicals from different studies reported in the literature.

Model	No. of compounds	No. of descriptors	Type of descriptors	R ²	Reference
RSM	200	2	PD, QD	0.540	Cronin et al. (2002)
kNN	250	173	—	0.646	Guo et al. (2005)
GRNN	202	—	HD, EpD, OH	0.794	Panaye et al. (2006)
RBF-NN	221	25	PD, QD, OH	0.942 (Train) 0.882 (Val)	Melagraki et al. (2006)
SCR-qQNA	200	—	QNAD	0.685	Lagunin et al. (2007)
SVM	983	36	MolconnZ	0.890 (Train) 0.830 (Val)	Zhu et al. (2008)
NN	250	6	PD, CD, OH	0.710 (Train) 0.730 (Val)	Enoch et al. (2008)
MLR	250	168	PD, CD, GD, QD, OH	0.660 (Train) 0.720 (Val)	Enoch et al. (2008)
MDF-SAR	30	4	TD, GD	0.974	Jäntschi et al. (2008)
MLR	313	—	ABSIL	0.730 (Train) 0.697 (Test)	Castillo-Garita et al. (2009)
Linear-QSTR	35	3	VD	N0.972	Vlaia et al. (2009)
ANN	392	7	—	0.822 (Train) 0.815 (Test)	Fatemi and Malekzadeh (2010)
PLS	95	58	PD, OH	0.840	Artemenko et al. (2011)
Stepwise-MLR	250	95	PD, MZD	0.600	Jiang et al. (2011)
DTB	1450	5	PD, CD, GD, TD	0.963 (Train) 0.874 (Test) 0.945 (Complete) 0.975 (ES)	Present study
DTF	1450	5	PD, CD, GD, TD	0.960 (Train) 0.876 (Test) 0.944 (Complete) 0.968 (ES)	Present study

PD—physico-chemical descriptor; TD—Topological descriptors; GD—Geometrical descriptor; CD—constitutional descriptor; QD—quantum mechanical descriptor; EpD—electrophilicity descriptors; HD—hydrophobicity descriptors; VD—vanderwaals spaced descriptors; QNAD—Quantitative neighborhood of atoms descriptors; MZD—Molconn-Z descriptors; ABSIL—atom-based non-stochastic linear indices; OH—others; ES—*T. pyriformis* external set.

to the fact that they make no assumption regarding the underlined distribution of values of predictor variables. Though the need of a minimal testing cannot be completely ignored, the proposed models can be used as effective tools for screening the chemicals for their toxicity profile.

icology: an update on the (Q)SAR models for mutagens and carcinogens. J. Environ. Sci. Health Part C25, 53-97.

Conflict of interest statement

The authors declare that there are no conflicts of interest.

Acknowledgments

The author thanks the Director, CSIR-Indian Institute of Toxicology Research, Lucknow (India) for his keen interest in this work and providing all necessary facilities.

Appendix A. Supplementary data

Supplementary data to this article can be found online at <http://dx.doi.org/10.1016/j.taap.2014.01.006>.

References

- Abdelaal, M.M.A., Farouq, M.W., Sena, H.A., Salem, A.B.M., 2010. Applied classification support vector machine for providing second opinion of breast cancer diagnosis. Online J. Math. Stat. 1, 1-7.
- Aptula, A., Roberts, D., Cronin, M., Schultz, T., 2005. Chemistry toxicity relationships for the effects of di- and trihydroxybenzenes to *Tetrahymena pyriformis*. Chem. Res. Toxicol. 18, 844-854.
- Artemenko, A.G., Muratov, E.N., Kuzmin, V.E., Muratov, N.N., Varlamova, E.V., Kuzmina, A.V., Gorv, L.G., Golius, A., Hill, F.C., Leszczynski, J., Tropsha, A., 2011. QSAR analysis of nitroaromatics toxicity in *Tetrahymena pyriformis*: structural factors and possible modes of action. SAR QSAR Environ. Res. 22, 575-601.
- Benigni, R., Netzeva, T.I., Benfenati, E., Bossa, C., Franke, R., Helma, C., Hulzebos, E., Marchant, C., Richard, A., Woo, Y.P., Yang, C., 2007. The expanding role of predictive toxic

- Benigni, R., Bossa, C., Richard, A., Yang, C., 2008. A novel approach: chemical relational databases, and the role of the *is*ss and database on assessing chemical carcinogenicity. *Ann. Ist. Super. Sanita* 44, 48-56.
- Breiman, L., 1996. Heuristics of instability and stabilization in model selection. *Ann. Stat.* 24, 2350-2383.
- Casalegno, M., Benfenati, E., Sello, G., 2005. An automated group contribution method in predicting aquatic toxicity: the diatomic fragment approach. *Chem. Res. Toxicol.* 18, 740-746.
- Castillo-Garit, J.A., Escobar, J., Marrero-Ponce, Y., Torrens, F., 2009. Atom-based quadratic indices to predict aquatic toxicity of benzened derivatives to *Tetrahymena pyriformis*. 13rd International Electronic conference on Synthetic Organic Chemistry (ECSO-13), 1-30 November, pp. 1-26.
- Cheng, F., Shen, J., Yu, Y., Li, W., Liu, G., Lee, P.W., Tang, Y., 2011. In silico prediction of *Tetrahymena pyriformis* toxicity for diverse industrial chemicals with substructure pattern recognition and machine learning methods. *Chemosphere* 82, 1636-1643.
- Consonni, V., Ballabio, D., Todeschini, R., 2009. Comments on the definition of the Q2 parameter for QSAR validation. *J. Chem. Inf. Model.* 49, 1669-1678.
- Cronin, M., Schultz, T., 2001. Development of quantitative structure-activity relationships for the toxicity of aromatic compounds to *Tetrahymena pyriformis*: a comparative assessment of the methodologies. *Chem. Res. Toxicol.* 14, 1284-1295.
- Cronin, M.T.D., Aptula, A.O., Duffy, J.C., Netzeva, T.I., Rowe, P.H., Valkova, I.V., Schultz, T.W., 2002. Comparative assessment of methods to develop QSARs for the prediction of the toxicity of phenols to *Tetrahymena pyriformis*. *Chemosphere* 49, 1201-1221.
- Cronin, M.T.D., Netzeva, T.I., Dearden, J.C., Edwards, R., Worgan, A.D.P., 2004. Assessment and modeling of the toxicity of organic chemicals to *Chlorella vulgaris*: development of a novel database. *Chem. Res. Toxicol.* 17, 545-554.
- Devillers, J., 2004. Linear versus nonlinear QSAR modeling of the toxicity of phenol derivatives to *Tetrahymena pyriformis*. *SAR QSAR Environ. Res.* 15, 237-249.
- Deweese, A.D., Schultz, T., 2001. Structure-activity relationships for aquatic toxicity to *Tetrahymena*: halogen-substituted aliphatic esters. *Environ. Toxicol.* 16, 54-60.
- Dietterich, T.G., 2000. Ensemble methods in machine learning. *Lect. Notes Comput. Sci.* 1857, 1-15.
- Duchowicz, P.R., Ocsachoque, M.A., 2009. Quantitative structure-toxicity models for heterogeneous aliphatic compounds. *QSAR Comb. Sci.* 28, 281-295.
- Efron, B., 1979. Bootstrap methods: another look at the jackknife. *Ann. Stat.* 7, 1-26.
- Enoch, S.J., Cronin, M.T.D., Schultz, T.W., Madden, J.C., 2008. An evaluation of global QSAR models for the prediction of the toxicity of phenols to *Tetrahymena pyriformis*. *Chemosphere* 71, 1225-1232.
- Eriksson, L., Jaworska, J., Worth, A.P., Cronin, M.T.D., McDowell, R.M., Gramatica, P., 2003. Methods for reliability and uncertainty assessment and for applicability evaluation of classification and regression-based QSARs. *Environ. Health Perspect.* 111, 1361-1375.

- Eroglu, E., Palaz, S., Oltulu, O., Turkmen, H., Ozaydin, C., 2007. Comparative QSTR study using semi-empirical and first principle methods based descriptors for acute toxicity of diverse organic compounds to the fathead minnow. *Int. J. Mol. Sci.* 8, 1265-1283.
- Fatemi, M.H., Malekzadeh, H., 2010. Prediction of $\text{Log}(IC_{50})^{-1}$ for benzenederivatives toxicity of *Tetrahymena pyriformis* from their molecular descriptors. *Bull. Chem. Soc. Jpn.* 83, 233-245.
- Fawcett, T., 2006. An introduction to ROC analysis. *Pattern Recogn. Lett.* 27, 861-874.
- Fedorova, N., Vracko, M., Novic, M., Roncaglioni, A., Benfenati, E., 2010. New public QSAR models for carcinogenicity. *Chem. Cent. J.* 4, 1-15.
- Friedman, J.H., 2001. Greedy function approximation: a gradient boosting machine. *Ann. Stat.* 29, 1189-1232.
- Friedman, J.H., 2002. Stochastic gradient boosting. *Comput. Stat. Data Anal.* 38, 367-378.
- Gramatica, P., 2007. Principles of QSAR models validation: internal and external. *QSAR Comb. Sci.* 26, 694-701.
- Guo, G., Neagu, D., Cronin, M.T.D., 2005. A study on feature selection for toxicity prediction. *Lect. Notes Comput. Sci.* 3614, 31-34.
- Hancock, T., Put, R., Coomans, D., Vander Heyden, Y., Everingham, Y., 2005. A performance comparison of modern statistical techniques for molecular descriptor selection and retention prediction in chromatographic QSRR studies. *Chemom. Intell. Lab. Syst.* 76, 185-196.
- Huber, P., 1964. Robust estimation of a location parameter. *Ann. Math. Stat.* 53, 73-101.
- In, Y., Lee, S.K., Kim, P.J., No, K.T., 2012. Prediction of acute toxicity of fathead minnow by Local Model based QSAR and global QSAR approaches. *Bull. Korean Chem. Soc.* 33, 613-619.
- Ishwaran, H., Kogalur, U.B., 2010. Consistency of random survival forests. *Stat. Probab. Lett.* 80, 1056-1064.
- Ivanciuc, O., 2004. Support vector machines prediction of the mechanism of toxic action from hydrophobicity and experimental toxicity against *Pimephales promelas* and *Tetrahymena pyriformis*. *Internet Electron. J. Mol. Des.* 2004(3), 802-821.
- Jalali-Heravi, M., Kyani, A., 2008. Comparative structure-toxicity relationship study of substituted benzenes to *Tetrahymena pyriformis* using shuffling-adaptive neurofuzzy inference system and artificial neural networks. *Chemosphere* 72, 733-740.
- Jäntschi, L., Popescu, V., Bolboacă, S.D., 2008. Toxicity caused by para-substituted phenols on *Tetrahymena pyriformis*: the structure-activity relationships. *Electron. J. Biotechnol.* 11, 1-12.
- Jiang, D.X., Li, Y., Li, J., Wang, G.X., 2011. Prediction of the aquatic toxicity of phenols to *Tetrahymena pyriformis* from molecular descriptors. *Int. J. Environ. Res.* 5, 923-938.
- John, M.B., Geoffrey, M.D., Willet, P., 1998. Chemical similarity searching. *J. Chem. Inf. Comput. Sci.* 38, 983-996.
- Jouan-Rimbaud, D., Massart, D.L., de Noord, O.E., 1996. Random correlation in variable selection for multivariate calibration with a genetic algorithm. *Chemom. Intell. Lab. Syst.* 35, 213-220.
- Katritzky, A.R., Oliferenko, P., Oliferenko, A., Lomaka, A., Karelson, M., 2003. Nitrobenzene toxicity: QSAR correlations and mechanistic interpretations. *J. Phys. Org. Chem.* 16, 811-817.
- Katritzky, A.R., Slavov, S.H., Stoyanova-Slavova, I.S., Kahn, I., Karelson, M., 2009. Quantitative structure-activity relationship (QSAR) modeling of EC_{50} of aquatic toxicities for *Daphnia magna*. *J. Toxicol. Environ. Health* 72, 1181-1190.
- Koleva, Y., Barzilov, I., 2010. Comparative study of mechanism of action of fatty alcohols for different endpoints. *Sci. Work. Univ.* 49, 55-59.
- Kubinyi, H., 1993. Methods and Principles in Medicinal Chemistry. In: Mannhold, R., Kroogsgard-Larsen, P., Timmerman, H. (Eds.), vol. 1. VCH, Wiley.
- Lagunin, A.A., Zakharov, A.V., Filimonov, D.A., Poroikov, V.V., 2007. A new approach to QSAR modeling of acute toxicity. *SAR QSAR Environ. Res.* 18, 285-298.
- Luan, F., Zhang, R.S., Zhao, C.Y., Yao, X.J., Liu, M.C., Hu, Z.D., Fan, B.T., 2005. Classification of the carcinogenicity of nitro compounds based on support vector machines and linear discriminant analysis. *Chem. Res. Toxicol.* 18, 198-203.
- Lukacinova, A., Mojzis, J., Benacka, R., Lovasova, E., Hijova, E., Nistiar, F., 2007. *Tetrahymena pyriformis* as a valuable unicellular animal model for determination of xenobiotics. *Proceedings 27th International Symposium "Industrial Toxicology"*, Bratislava, 30 May-1 June 2007, pp. 1-6.
- Mahajan, D.T., Masand, V.H., Patil, K.N., Hadda, T.B., Rastija, V., 2013. Integrating GUSAR and QSAR analyses for antimalarial activity of synthetic prodiginines against multidrug resistant strain. *Med. Chem. Res.* 22, 2284-2292.
- Mahjoobi, J., Etemad-Shahidi, A., 2008. An alternative approach for the prediction of significant wave heights based on classification and regression trees. *Appl. Ocean Res.* 30, 172-177.
- Manley, P.W., Stieff, N., Cowan-Jacob, S.W., Kaufman, S., Mestan, J., Wartmann, M., Wiesmann, M., Woodman, R., Gallagher, N., 2010. Structural resemblances and comparisons of the relative pharmacological properties of imatinib and nilotinib. *Bioorg. Med. Chem.* 18, 6977-6986.
- Masand, V.H., Jawarkar, R.D., Patil, K.N., Mahajan, D.T., Hadda, T.B., Kurhade, G.H., 2010. COM FA analysis and toxicity risk assessment of coumarin analogues as MAO-A inhibitors: attempting better insight in drug design. *Der. Pharm. Lett.* 2, 350-357.
- Melagraki, G., Afantitis, A., Makridima, K., Sarimveis, H., Igglessi-Markopoulou, O., 2005. Prediction of toxicity using a novel RBF neural network training methodology. *J. Mol. Model.* 8, 1-9.
- Melagraki, G., Afantitis, A., Makridima, K., Sarimveis, H., Igglessi-Markopoulou, O., 2006. Prediction of toxicity using a novel RBF neural network training methodology. *J. Mol. Model.* 12, 297-305.
- Netzeva, T., Schultz, T., 2005. QSARs for the aquatic toxicity of aromatic aldehydes from *Tetrahymena* data. *Chemosphere* 61, 1632-1643.
- Netzeva, T., Schultz, T., Aptula, A., Cronin, M., 2003. Partial least squares modelling of the acute toxicity of aliphatic compounds to *Tetrahymena pyriformis*. *SAR QSAR Environ. Res.* 14, 265-283.

- Netzeva, T.I., Worth, A.P., Aldenberg, A., Benigni, R., Cronin, M.T.D., Gramatica, P., Jaworska, J.S., Kahn, S., Kloman, G., Marchant, C.A., Myatt, G., Nikolova-Jeliazkova, N., Patlitz, G.Y., Perkins, R., Roberts, D.W., Schultz, T.W., Stanton, D.P., van de Sandt, J.J.M., Tong, W., Veith, G., Yang, C., 2005. Current status of methods for defining the applicability domain of (quantitative) structure-activity relationship. *ATLA* 33, 155-173.
- Nikolova-Jeliazkova, N., Jaworska, J., 2005. An approach to determining applicability domains for QSAR group contribution models: an analysis of SRCKOWWIN. *Altern. Lab. Anim.* 33, 461-470.
- Niu, B., Jin, Y., Lu, W., Li, G., 2009. Predicting toxic action mechanisms of phenols using AdaBoostLearner. *Chemom. Intell. Lab. Syst.* 96, 43-48.
- Ojha, P.K., Mitra, I., Das, R.N., Roy, K., 2011. Further exploring r^2 metrics for validation of QSPR models. *Chemom. Intell. Lab. Syst.* 107, 194-205.
- Panaye, A., Fan, B.T., Doucet, J.P., Yao, X.J., Zhang, R.S., Liu, M.C., Hu, Z.D., 2006. Quantitative structure-toxicity relationships (QSTRs): a comparative study of various non linear methods. General regression neural network, radial basis function neural network and support vector machine in predicting toxicity of nitro-and cyano-aromatics to *Tetrahymena pyriformis*. *SAR QSAR Environ. Res.* 17, 75-91.
- Ren, S., 2003. Ecotoxicity prediction using mechanism and non-mechanism base QSARs: a preliminary study. *Chemosphere* 53, 1053-1065.
- Ren, S., Schultz, T., 2002. Identifying the mechanism of aquatic toxicity of selected compounds by hydrophobicity and electrophilicity descriptors. *Toxicol. Lett.* 129, 151-160.
- Roex, E.W.M., Gestel, C.A.M.V., Wezel, A.P.V., Straalen, N.M.V., 2000. Ratios between acute aquatic toxicity and effects on population growth rates in relation to toxicant mode of action. *Environ. Toxicol. Chem.* 19, 685-693.
- Roy, D., Parthasarathi, R., Maiti, B., Subramanian, V., Chattaraj, P., 2005. Electrophilicity as a possible descriptor for toxicity prediction. *Bioorg. Med. Chem.* 13, 3405-3412.
- Roy, P.P., Leonard, J.T., Roy, K., 2008. Exploring the impact of size of training sets for the development of predictive QSAR models. *Chemom. Intell. Lab. Syst.* 90, 31-42.
- Roy, P.P., Paul, S., Mitra, I., Roy, K., 2009. On two novel parameters for validation of predictive QSAR models. *Molecules* 29, 1660-1701.
- Schultz, T.W., 1997. Tetraox: *Tetrahymena pyriformis* population growth impairment endpoint—as surrogate for fish lethality. *Toxicol. Methods* 7, 289-309.
- Schultz, T.W., 1999. Structure-toxicity relationships for benzenes evaluated with *Tetrahymena pyriformis*. *Chem. Res. Toxicol.* 12, 1262-1267.
- Schultz, T.W., Netzeva, T.I., 2004. Development and evaluation of QSARs for ecotoxic end-points: The benzene response-surface model for *Tetrahymena toxicity*. In: Cronin, M.T.D., Livingstone, D.J. (Eds.), *Modeling environmental fate and toxicity*. CRC Press, Boca Raton, FL, USA, pp. 265-284.
- Schultz, T., Netzeva, T., Roberts, D., Cronin, M., 2005. Structure toxicity relationships for the effects of *Tetrahymena pyriformis* of aliphatic, carbonyl-containing, R, α -unsaturated chemicals. *Chem. Res. Toxicol.* 18, 330-341.
- Schultz, T.W., Hewitt, M., Netzeva, T.I., Cronin, M.T.D., 2007. Assessing applicability domains of toxicological QSARs: definition, confidence in predicted values, and the role of mechanisms of action. *QSAR Comb. Sci.* 26, 238-254.
- Schuurmann, G., Aptula, A., Kuhne, R., Ebert, R., 2003. Stepwise discrimination between four modes of toxic action of phenols in the *Tetrahymena pyriformis* assay. *Chem. Res. Toxicol.* 16, 974-987.
- Serra, J., Jurs, P., Kaiser, K., 2001. Linear regression and computational neural network prediction of *tetrahymena* acute toxicity for aromatic compounds from molecular structure. *Chem. Res. Toxicol.* 14, 1535-1545.
- Singh, K.P., Basant, A., Malik, A., Jain, G., 2009a. Artificial neural network modeling of the river water quality—acase study. *Ecol. Model.* 220, 888-895.
- Singh, K.P., Ojha, P., Malik, A., Gain, J., 2009b. Partial least squares and artificial neural network modeling for predicting chlorophenol removal from aqueous solution. *Chemom. Intell. Lab. Syst.* 99, 150-160.
- Singh, K.P., Basant, N., Malik, A., Jain, G., 2010. Modeling the performance of “up-flow anaerobic sludge blanket” reactor based wastewater treatment plant using linear and nonlinear approaches—acase study. *Anal. Chim. Acta* 658, 1-11.
- Singh, K.P., Basant, N., Gupta, S., 2011. Support vector machine in water quality management. *Anal. Chim. Acta* 703, 152-162.
- Singh, K.P., Gupta, S., Kumar, A., Shukla, S.P., 2012. Linear and nonlinear modeling approaches for urban air quality prediction. *Sci. Total Environ.* 426, 244-255.
- Singh, K.P., Gupta, S., Rai, P., 2013a. Predicting acute aquatic toxicity of structurally diverse chemicals in fish using artificial intelligence approaches. *Ecotoxicol. Environ. Saf.* 95, 221-233.
- Singh, K.P., Gupta, S., Rai, P., 2013b. Predicting dissolved oxygen concentration using kernel regression modeling approaches with nonlinear hydro-chemical data. *Environ. Monit. Assess.* <http://dx.doi.org/10.1007/s10661-013-3576-6>.
- Snelder, T.H., Lamouroux, N., Leathwick, J.R., Pella, H., Sauquet, E., Shanker, U., 2009. Predictive mapping of the natural flow regimes of France. *J. Hydrol.* 373, 57-67.
- Steinbeck, C., Han, Y., Kuhn, S., Horlacher, O., Luttmann, E., Willighagen, E., 2003. The Chemistry Development Kit (CDK)—an open-source Java library for chemical and bio-informatics. *J. Chem. Inf. Model.* 43, 493-500.
- Tan, N.X., Li, P., Rao, H.B., Li, Z.R., Li, X.Y., 2010. Prediction of the acute toxicity of chemical compounds to fathead minnow by machine learning approaches. *Chemom. Intell. Lab. Syst.* 100, 66-73.
- Terada, H., 1990. Uncouplers of oxidative phosphorylation. *Environ. Health Perspec.* 87, 213-218.

- Toropov, A.A., Toropova, A.P., Benfenati, E., Manganaro, A., 2010. QSAR modeling of the toxicity to *Tetrahymena pyriformis* by balance of correlations. *Mol. Divers.* 14, 821-827.
- Tropsha, A., 2010. Best practices for QSAR model development, validation, and exploitation. *Mol. Inf.* 29, 476-488.
- Veith, G.D., Broderius, S.J., 1990. Rules for distinguishing toxicants that cause type I and type II narcosis syndromes. *Environ. Health Perspect.* 87, 207-211.
- Verhaar, H.J.M., Van Leeuwen, C.J., Hermens, J.L.M., 1992. Classifying environmental pollutants. *Chemosphere* 25, 471-491.
- Vlaia, V., Olariu, T., Ciubotariu, C., Medeleanu, M., Vlaia, L., Ciubota, D., 2009. Molecular descriptors for quantitative structure-toxicity relationship (QSTR). II. Three ovality molecular descriptors and their use in modeling the toxicity of aliphatic alcohols on *Tetrahymena pyriformis*. *Rev. Chim. (Bucuresti)* 60, 1357-1361.
- Wang, Y., Zheng, M., Xiao, J., Lu, Y., Wang, F., Lu, J., Luo, X., Zhu, W., Jiang, H., Chen, K., 2010. Using support vector regression coupled with the genetic algorithm for predicting acute toxicity to the fathead minnow. *SAR QSAR Environ. Res.* 21, 559-570.
- Wang, G., Hao, J., Ma, J., Jiang, H., 2011. A comparative assessment of ensemble learning for credit scoring. *Expert Syst. Appl.* 38, 223-230.
- Xue, Y., Li, H., Ung, C.Y., Yap, C.W., Chen, Y.Z., 2006. Classification of a diverse set of *Tetrahymena pyriformis* toxicity chemical compounds from molecular descriptors by statistical learning methods. *Chem. Res. Toxicol.* 19, 1030-1039.
- Yang, P., Yang, Y.H., Zhou, B.B., Zomaya, A.Y., 2010. A review of ensemble methods in bioinformatics. *Curr. Bioinforma.* 5, 296-308.
- Zhang, X.J., Qin, H.W., Su, L.M., Qin, W.C., Zou, M.Y., Sheng, L.X., Zhao, Y.H., Abraham, M.H., 2010. Interspecies correlations of toxicity to eight aquatic organisms: theoretical considerations. *Sci. Total Environ.* 408, 4549-4555.
- Zhao, C.Y., Zhang, H.X., Zhang, X.Y., Liu, M.C., Hu, Z.D., Fan, B.T., 2006. Application of support vector machine (SVM) for prediction toxic activity of different data sets. *Toxicology* 217, 105-119.
- Zhu, H., Tropsha, A., Fourches, D., Varnek, A., Papa, E., Gramatica, P., Oberg, T., Dao, P., Cherkasov, A., Tetko, I.V., 2008. Combinatorial QSAR modeling of chemical toxicants tested against *Tetrahymena pyriformis*. *J. Chem. Inf. Model.* 48, 766-784.