

Learning highly non-separable Boolean functions using Constructive Feedforward Neural Network

Marek Grochowski and Włodzisław Duch

Department of Informatics, Nicolaus Copernicus University,
Grudziądzka 5, Toruń, Poland,
grochu@is.umk.pl; Google: Duch

17th International Conference on Artificial Neural Networks

Introduction

Learning of problems with inherent non-separable Boolean logic is still a challenge (text categorization, medical data, natural language processing, bioinformatics, etc.)

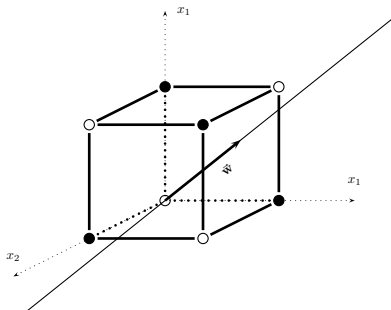
Many learning systems, such as SVMs, decision trees or similarity-based methods cannot deal with such problems - obtained solutions are too complex, generalization is poor

The goal is to find general algorithm capable of solving a large set of problems with the smallest possible number of parameters

N -bit parity problem

Examples of solving n -bit parity problem

- it is trivial with periodic function with a single parameter
- MLP need $O(n^2)$ parameters, learning is difficult
- most of neural network solutions proposed to solve n -bit parity problem will not work for other complex problems (E. Iyoda, H. Nobuhara, K. Hirota (2003); D. Stork, J. Allen (1992), B. Wilamowski, D. Hunter (2003))
- projection on a line with n threshold parameters (k -separable solution), Duch (2006)



Non-separable Boolean functions

The difficulty of learning Boolean functions grows quickly with the minimum k required to solve a given problem.

For many complicated problems often a simple linear mapping exists that leaves only trivial non-linearities that may be separated using window-like neurons.

The main idea is to create constructive network with neurons that realize window-like functions

Parity problem

Example of 10 dimensional parity problem ($2^{10} = 1024$ vectors)
Sequencion of labels of projection on
random direction

```
0110101010010110110110000100110000110111001101010000111000111000
1111110010001100111100001111000011100111000110010000111011000111
1110011111000111000011000101111010111000011000110011001100011001
1001100111000110011000011110100111000000111111000110001110011000
000110011000110011100111110000011010100001111000000111101010001
0011111100000011101100011110011100001100100110001111000001100000
001111111000111111110000110011100001100110001111000000111000111
11000000001111111010001111111111100000000011111000011111100001
1000011111100001111100001111000000111101110100110000111111000011
1101001110000001111000110011001110000110011111100100011111110000
0011011000000011111100000111110000001111100011011100000011111100
1000101011110000001111000010101100000111111001110011000110011000
0001100111000110001111110000001110010111100001100110111000011110
0001100011001100110001100001110101111010001100001110001111100111
1110001101110000100110001110000111001100110011110011000100111111
0001110001100100101011001110110000110000100110110110100101010110
```

Example of 10 dimensional parity problem ($2^{10} = 1024$ vectors)

diagonal direction

[illegible] $k = 11$ (11-separability)

Example of 10 dimensional parity problem ($2^{10} = 1024$ vectors)

diagonal direction

[illegible] $k = 11$ (11-separability)

Error Function

Combination of projection and clustering

Hard-window transfer function

$$\tilde{M}_i(\vec{x}; \vec{w}, a, b) = \begin{cases} 1 & \text{if } \vec{w}\vec{x} \in [a, b] \\ 0 & \text{if } \vec{w}\vec{x} \notin [a, b] \end{cases}$$

Error measure

$$E(\vec{x}; \Gamma) = E_{\vec{x}} \|y(\vec{x}; \Gamma) - c(\vec{x})\|$$

does not give any control over purity of clusters

Error function with purity term

$$E(\vec{x}; \Gamma; a, b, \lambda) = E_{\vec{x}} \|y(\vec{x}; \Gamma) - c(\vec{x})\| + \lambda E_{y \in [a, b]} \|y(\vec{x}; \Gamma) - c(\vec{x})\|$$

λ controls the tradeoff between the covering and the purity

Transfer Functions

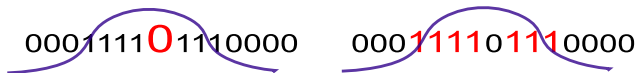
Soft window-like transfer functions

- $M(\vec{x}; \vec{w}, a, b, \beta) = \sigma(\beta(\vec{w}\vec{x} - a)) - \sigma(\beta(\vec{w}\vec{x} - b))$
- $M(\vec{x}; \vec{w}, t, a, \beta) = \sigma(\beta(\vec{w}\vec{x} - t - a))(1 - \sigma(\beta(\vec{w}\vec{x} - t + a)))$
- $M(\vec{x}; \vec{w}, a, b, \beta) = \frac{1}{2}(1 - \tanh(\beta(\vec{w}\vec{x} - a)) \tanh(\beta(\vec{w}\vec{x} - b)))$

slope parameter controls "softness" of transfer function $M(\vec{x}; \Gamma_i) \xrightarrow{\beta \rightarrow \infty} \tilde{M}(\vec{x}; \Gamma_i)$

Error Function

$$E(\vec{x}; \Gamma, \lambda_1, \lambda_2) = \frac{1}{2} \sum_{\vec{x}} (y(\vec{x}; \Gamma) - c(\vec{x}))^2 + \underbrace{\lambda_1 \sum_{\vec{x}} (1 - c(\vec{x})) y(\vec{x}; \Gamma)}_{\text{penalty}} - \underbrace{\lambda_2 \sum_{\vec{x}} c(\vec{x}) y(\vec{x}; \Gamma)}_{\text{reward}}$$

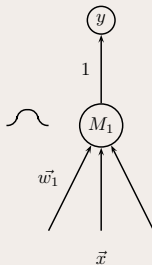


Penalty factor λ_1 increases the total error for vectors x_i from class $c(x_i) = 0$ that falls into group of vectors from class 1 (it is a penalty for "unclean" clusters).

Reward factor λ_2 decreases the value of total error for every vector x_i from class 1 that was correctly placed inside created clusters (it is a reward for large clusters).

Learning

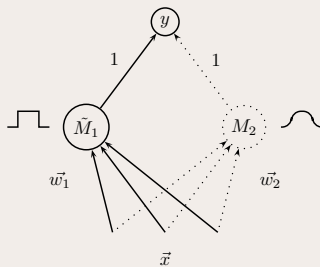
Constructive algorithm



Learning starts with an empty hidden layer and in every phase of the training one new unit is added, initialized and trained using the backpropagation algorithm

Learning

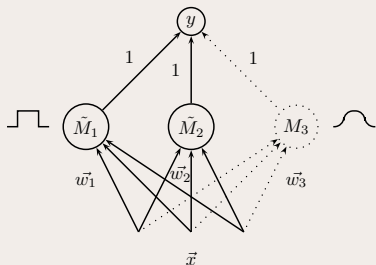
Constructive algorithm



The input vectors correctly handled by the first neuron do not contribute to the error, therefore the weights of this neuron are kept frozen during further learning. Slope β is set to large value to obtain hard boundaries. Next node is added and learning procedure is repeated on the remaining data

Learning

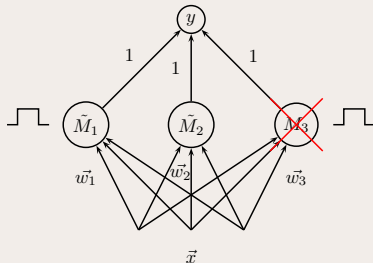
Constructive algorithm



Most nodes and connections are fixed and only weights of one node are modified at each training step

Learning

Constructive algorithm

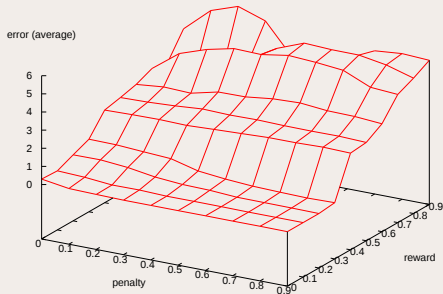


If number of cases correctly classified by a given new node drops below certain minimum the learning procedure stops and this node is removed from the network

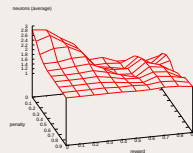
Influence of Penalty and Reward

Average over 192 6-separable 4 dimensional Boolean functions

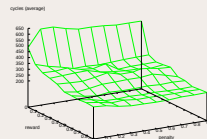
Error



Neurons



Cycles

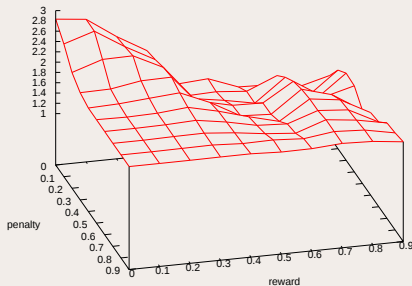


Influence of Penalty and Reward

Average over 192 6-separable 4 dimensional Boolean functions

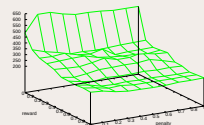
Neurons

neurons (average)



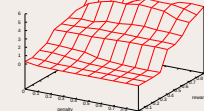
Cycles

cycles (average)



Error

error (average)

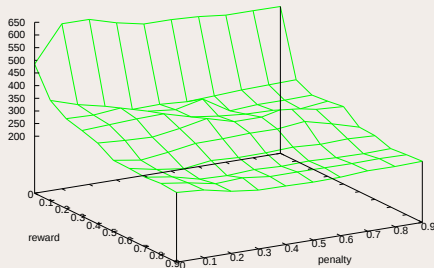


Influence of Penalty and Reward

Average over 192 6-separable 4 dimensional Boolean functions

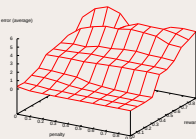
Cycles

cycles (average)



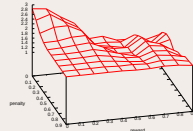
Error

error (average)



Neurons

neurons (average)



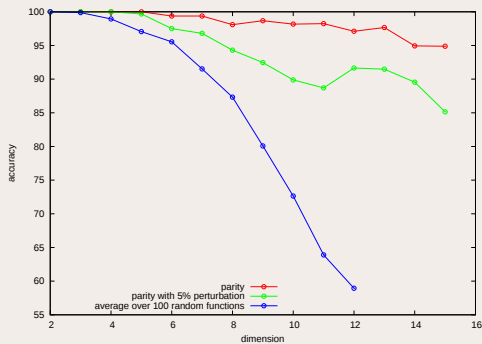
Boolean functions

parity

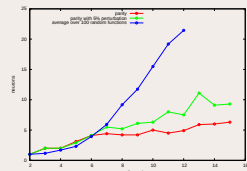
parity functions with small perturbation of labels

random Boolean functions

Average training accuracy



Average no. of neurons



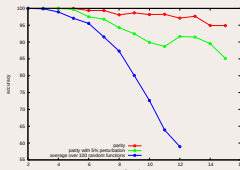
Boolean functions

parity

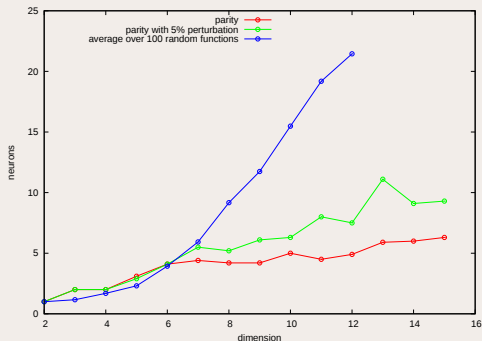
parity functions with small perturbation of labels

random Boolean functions

Average training accuracy



Average no. of neurons



Real World Data

Test for datasets from UCI repository

10x10 CV test accuracy

dataset	1-NN	Naive Bayes	SVM	c3sep
Appendicitis	81.3 ± 1.5	85.3 ± 1.0	86.5 ± 0.3	85.3 ± 1.0
Australian	78.0 ± 0.2	80.0 ± 0.3	85.0 ± 0.1	86.3 ± 0.6
Flag	50.1 ± 1.1	41.1 ± 1.1	51.1 ± 1.1	53.6 ± 1.8
Glass	68.2 ± 1.7	47.4 ± 1.7	66.7 ± 0.9	61.1 ± 1.3
Ionosphere	85.2 ± 1.2	82.2 ± 0.2	85.2 ± 0.2	85.1 ± 1.5
Iris	95.9 ± 0.5	94.9 ± 0.2	95.5 ± 0.3	95.7 ± 1.0
Pima-diabetes	70.5 ± 0.5	75.2 ± 0.5	70.3 ± 1.0	76.3 ± 0.4
Promoters	78.5 ± 1.8	85.81 ± 1.3	93.1 ± 1.5	74.7 ± 5.6
Sonar	86.8 ± 1.8	67.8 ± 1.2	84.2 ± 1.1	77.9 ± 2.4
Wine	95.1 ± 0.8	98.1 ± 0.3	95.1 ± 0.2	97.1 ± 0.8

Real World Data

Test for datasets from UCI repository

Comparison of complexity SVM vs. c3sep

	support vectors	neurons (total)	average number of neurons per class									
Appendicitis	32.1	4.2	2.2	2.0								
Australian	207.2	8.8	4.5	4.3								
Flag	315.2	26.7	5.1	6.1	4.0	2.1	2.7	3.6	1.1	1.6		
Glass	295.8	14.0	1.4	3.9	1.0	2.8	1.9	2.8				
Ionosphere	63.9	7.9	3.0	4.9								
Iris	43.4	5.0	1.0	2.0	2.0							
Pima-diabetes	365.3	9.1	4.2	4.8								
Promotores	77.2	3.7	1.9	1.8								
Sonar	109.7	8.5	4.5	4.0								
Wine	63.3	4.0	1.0	2.0	1.0							

Summary and Further Work

Some conclusions

- First steps towards efficient learning of Boolean functions have been made here
- The approach presented here has been able to learn quite difficult Boolean functions using a very simple model
- Great advantage is small computational costs of the constructive network training

Further work

- Global minimization algorithms instead of backpropagation
- looking for k -separable solution -e.g. network with weight shering
- different error function
- additional transformation of input data
- Experiments with high dimensional real world problems
- Many more ideas ...