

ASSESSMENT OF IMPORTANT RISK FACTORS OF HYPERTENSION USING DISCRIMINANT ANALYSIS, LOGISTIC REGRESSION AND NEURAL NETWORKS

¹N. Shehu, ²S. U. Gulumbe and ³H. M. Liman

¹Statistics Unit, School of Health Information Management, Usmanu Danfodiyo University Teaching Hospital, Sokoto, Nigeria

²Statistics Unit, Department of Mathematics, Usmanu Danfodiyo University Sokoto, Nigeria

³Medicine Department, College of Health Sciences, Usmanu Danfodiyo University Sokoto, Nigeria

ABSTRACT

Hypertension (HTN) is the third leading cause of morbidity and mortality worldwide. It represents the single greatest preventable cause of death in humans and one of the most important modifiable risk factor for cardiovascular diseases. Financial and public health consequences of both hypertension and the failure to control it are enormous. Identification of different risk factors is important for the prevention and control of hypertension. This paper examines the most important risk factors of hypertension among age, cigarette smoking, stress, alcoholism, sedentary life style, family history, comorbidity, arteriosclerosis using discriminant analysis, logistic regression and neural networks. It reveals that, among age, cigarette smoking, stress, alcoholism, sedentary life style, family history, comorbidity, arteriosclerosis; age and family history are the most important predictors. Moreover, based on neural networks analysis, the model identified not only age and family history but also stress, alcoholism and comorbidity as important predictors also.

KEYWORDS: Hypertension, Discriminant Analysis, Logistic Regression, Neural Networks

Received for Publication: 22/05/13

Accepted for Publication: 28/07/13

Corresponding Author: nsbozo@gmail.com, nsbozo@yahoo.com

INTRODUCTION

Hypertension is one of the important public health challenges worldwide because of its high frequency and consequent risks of cardiovascular, kidney and other end-organ diseases. It has been identified as a leading risk factor for mortality and ranked third as a cause of disability-adjusted life-years (Tefaye *et al*, 2007). Hypertension has long been established as a strong, independent, and etiologically significant risk factor for cardiovascular disease. Despite progress in prevention, treatment, and control of high blood pressure, hypertension remains a major public health challenge. High blood pressure (BP) or hypertension is a common condition in Nigeria and is a risk factor for heart attack, stroke, left ventricular hypertrophy, renal failure, and blindness. People who have hypertension are usually unaware that they have the condition, unless the BP has been measured at health-care facilities (Ashish *et al*, 2004). It is therefore frequently referred to as a 'silent epidemic' in the world. Consequently, hypertension is universally under diagnosed and/or inadequately treated resulting in extensive target-organ damage and premature death (Ukoli *et al*, 1995). Furthermore, hypertension frequently co-exists with other chronic diseases of lifestyle (CDL), such



All rights reserved

This work by [Wilolud Journals](http://www.wiloludjournal.com) is licensed under a [Creative Commons Attribution 3.0 Unported License](http://creativecommons.org/licenses/by/3.0/)

as diabetes and obesity. These interrelationships of hypertension with other CDL risk factors and the various possible target organs that can be influenced by uncontrolled hypertension result in a diverse picture that has an impact on the different population groups of people. This picture can also be influenced by the impact of socio demographic or genetic factors in different settings.

Despite the availability of a wide range of antihypertensive drugs, hypertension and its complications are still important causes of adult morbidity and mortality in sub-Saharan Africa (Isezuo *et al.* 2000). More than 50% of treated hypertensive patients have a blood pressure level greater than 140/90 mm Hg (uncontrolled hypertension) (Salako *et al.* 2003). Several factors including, among others, poor adherence to therapeutic regimen, ignorance, and poverty (Isezou *et al.* 2000) have been adduced for the high prevalence of uncontrolled hypertension. Recent reports have however focused on the role of health care provider to poor adherence to antihypertensive drugs (Hyman *et al.* 2001). Consequently, compliance with standard guidelines aiding physicians in effective prescription of antihypertensive drugs have been emphasized.

There are many things which contribute to an individual's risk of developing high blood pressure. They are collectively called "risk factors". Many diseases have important risk factors, and high blood pressure is no exception. These include: smoking, obesity, sedentary lifestyle, lack of physical activity, high level of salt intake, insufficient calcium, potassium and magnesium consumption, vitamin D deficiency, high level of alcohol consumption, stress, aging, medicine such as birth control pills, genetics and a family history of hypertension, chronic kidney disease, adrenal and thyroid problems or tumors.

Hypertension is one of the important public health challenges worldwide because of its high frequency and concomitant risks of cardiovascular and kidney disease. Despite progress in prevention, treatment, and control of high of blood pressure, hypertension remains a major public health challenge. This therefore, patient with hypertension often have multiple cardiovascular risk factors, such as diabetes and dyslipidemia and should be evaluated for these other risk factors. High blood pressure management must address all risk factors if the goal is to reduce patient's global cardiovascular risk. Researchers, clinicians and patients must work together to develop creative approaches to achieve adherence, trying different approaches as necessary. Most researches conducted on hypertension adopted conventional approaches which assumed normality and linear relationship between the response variable and predictors. There are possibilities in which the relationship is non-linear, and for this reason conventional methods may not be appropriate. Thus, there is need to adopt alternative method that can capture non-linearity relationship among variables. Classical statistical methods have been applied in industry for years. Recently, Neural Networks (NNs) methods have become tools of choice for a wide variety of applications across many disciplines. It has been recognized in the literature that regression and neural network methods have become competing model-building methods (Smith & Mason, 1997). For a large class of pattern-recognition processes, NNs is the preferred technique (Setyawati *et al.*, 2002). NNs methods have also been used in the areas of prediction and classification (Warner & Misra, 1996).

This research is aimed to examine the important risk factors of hypertension using discriminant analysis, logistic regression and neural networks. This can be achieved through the following objectives: To develop a logistic regression model capable of predicting a patient's hypertension status, develop a linear discriminant model capable of predicting a patient's hypertension status, build an artificial neural networks model capable of identifying a patient's hypertension status.



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

MATERIALS AND METHODS

Data Used

The data used in this research were randomly selected from the patients' case folders in the Medical Out Patient Department and Library Services Section of Usmanu Danfodiyo University Teaching Hospital (UDUTH), Sokoto. Out of 300 patient case folders reviewed, 179 patients were Hypertensive while the remaining 121 patients were Not Hypertensive. The data was subjected to model building with respect to Linear Discriminant, Logistic Regression and Neural Networks.

Discriminant Analysis

Linear discriminant analysis (LDA) is a technique that models both intra-class and inter-class variance as multidimensional Gaussians. It seeks directions in space that have maximum discriminability and are hence most suitable for supporting class recognition tasks. The objective of LDA is to perform dimensionality reduction while preserving as much of the class discriminatory information as possible. It seeks to find directions along which the classes are best separated. The aim of discriminant analysis (DA) is to classify a heterogeneous population into homogenous subsets and further the decision process on these subsets. Using discriminant analysis demands satisfying certain assumption, for example, assumption of normality, linearity, homoscedasticity, but the method is practically healthy to violate these assumptions (Meyers *et al* 2005). At this point, in this study, it is assumed that all required assumptions are satisfied to apply the predictive power of discriminant analysis for classification of patient either as hypertensive or not.

DA involves the determination of a linear equation like regression that will predict which group the case belongs to. The form of the equation or function is:

$$D = b_1X_1 + b_2X_2 + \dots + b_iX_i + C$$

Where:

D = discriminant function

b's = are the discriminant coefficients or weights of the variables

X's = are explanatory variables

C = a constant

Logistic Regression

Logistic regression measures the relationship between a categorical dependent and usually a continuous independent (or several), by converting the dependent variable to probability scores (Hosmer *et al* 2000).

Logistic regression can be bi- or multinomial. Binomial or binary logistic regression refers to the instance in which the observed outcome can have only two possible types (e.g., "dead" vs. "alive", "success" vs. "failure", or "yes" vs. "no"). [Multinomial logistic regression](#) refers to cases where the outcome can have three or more possible types (e.g., "better" vs. "no change" vs. "worse"). Generally, the outcome is coded as "0" and "1" in binary logistic regression as it leads to the most straightforward interpretation (Cohen *et al* 2002). The target group (referred to as a "case") is usually coded as "1" and the reference group (referred to as a "non-case") as "0". Logistic regression is used to predict the [odds](#) of being a case based on the predictor(s). The odds are defined as the probability of a case divided by the probability of a non-case. The [odds ratio](#) is the primary measure of effect size in logistic regression and is computed to compare the odds that membership in one group will lead to a case outcome with the odds that membership in some other group will lead to a case outcome. The odds ratio (denoted OR) is simply the odds of being a case for one group divided by the odds of being a case for another group. An odds ratio of one indicates that the odds of a case outcome are equally likely for both groups under comparison. The further the odds deviate from one,



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

the stronger the relationship. The odds ratio has a floor of zero but no ceiling (upper limit) - theoretically, the odds ratio can increase infinitely (Hosmer *et al* 2000).

The dependent variable in logistic regression is usually dichotomous, that is, the dependent variable can take the value 1 with a probability of success θ , or the value 0 with probability of failure $1-\theta$. This type of variable is called a Bernoulli (or binary) variable. Although, applications of logistic regression have also been extended to cases where the dependent variable is of more than two cases, known as multinomial or polytomous. As mentioned previously, the independent or predictor variables in logistic regression can take any form. That is, logistic regression makes no assumption about the distribution of the independent variables. They do not have to be normally distributed, linearly related or of equal variance within each group. The relationship between the predictor and response variables is not a linear function in logistic regression (Hosmer *et al* 2000), instead, the logistic regression function is used, which is the logit transformation of θ :

$$\theta = \frac{\exp(\alpha + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_i X_i)}{1 + \exp(\alpha + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_i X_i)}$$

Where α = the constant of the equation, β_i = the coefficient of the predictor variables. The goal of logistic regression is to correctly predict the category of outcome for individual cases using the most parsimonious model. To accomplish this goal, a model is created that includes all prediction variables that are useful in predicting the response variable.

Neural Networks

Neural Networks (NNs) are non-linear mapping structures based on the function of the human brain (Francesca, 2006). Neural Networks can identify and learn correlated patterns between input data sets and corresponding target values. Neural Networks are made up of interconnected neurons. Each neuron has a certain number of inputs. There is a real number associated with each connection, which is called the weight of the connection, and is denoted by W . Each neuron has its own unique threshold value. The behavior of an NNs depends on both the weights and input-output function (transfer function) that is specified for the units. The combination function and the transfer function make up the activation function of the node. Activation function defines the output of a neuron. The three common transfer functions are sigmoid, linear and hyperbolic functions. Out of these sigmoid function is most common form of activation function used in the construction of neural network (Francesca, 2006). It produces the values between 0 and 1 for any input from the combination function. The hidden neuron act as feature detectors, as such they play a critical role in the operation of multilayer perceptron, the most popular approach to find the optimal number of hidden layer is by trail and error. There are two decisions to be considered with regards to the hidden layers such as number of hidden layers and number of hidden neurons. Hidden layer automatically extracts the features of the input and reduces its dimensionality further. There is no specific rule that dictates the number of hidden layers. Usually, one hidden layer is used. The reason for this is that one hidden layer is sufficient to approximate any continuous function to an arbitrary precision (Sumathi & Santhakumar 2010). This model is used one hidden layer with 4 neurons & sigmoid functions. The number of hidden neurons should be in the range between the size of the input layer and the size of the output layer. If the number of neuron in the hidden layer is more, the network becomes complicated. Results indicate that the present problem is not too complex to have a complicated network routing. Learning is a process by which the parameters of a neural network are adapted through a continuous process of simulation by the environment in which the network is embedded. The type of learning is determined by the manner in which the parameter changes take place (Francesca, 2006). Supervised algorithm is also known as “error-based learning algorithms,” which employ an external reference signal (teacher) and generate an error signal by



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

comparing the reference with the obtained response. Based on error signal, Neural Network modifies its synaptic connections to improve the system performance. Neural Network have the ability to adapt to any situation and therefore it learns how to solve a problem. Learning with a teacher is also called supervised learning. During the learning process global information may be required. An important issue concerning supervised learning is the problem of error convergence, that is the minimization of error between the desired response and actual response of the network. This adjustment is carried out iteratively in a step-by-step process.

Back propagation Algorithm

Haykin (1999), many types of ANN have been specified in the literature, the most widely used method which is suitable for classification is the back propagation algorithm or multilayer perception as the name suggest, and is a learning rule for multi-layered Neural Networks. Back-Propagation networks are fully connected, layered, feed forward networks, in which activations flow from the input layer through the hidden layer(s) and then to the output layer. Back propagation uses supervised learning in which the network is trained using data for which inputs as well as desired outputs are known. In order to train a neural network to perform some task, the weight of each unit must be adjusted, in such a way that the error between the desired output and the actual output is reduced (Badr *et al* 2003). The algorithm gives a prescription for adjusting the initially randomized set of synaptic weights. Training algorithm of back propagation include four stages as:

- Initialize the weight.
- Feed forward.
- Back propagation of errors.
- Updating the weights and biases

Begins by constructing a network with the desired number of hidden and output units and initializing all network weights to small random values. The main loop of algorithm then repeated over the training examples. For each training examples, calculate the error of the network output, computes the gradient with respect to the error, updates all weights in the network. Repeated until the network performs acceptably well. A unit in the output layer determines its activity by following two step procedures:

- computes the total input I using the formula:

$$I = \sum_{i=1}^n W_i X_i$$

Where $X_1, X_2, \dots, X_n$ are the n inputs to the artificial neuron and $W_1, W_2, \dots, W_n$ are the weights attached to the input links.

- the unit calculates the activity X using some function of the total weighted input, the activation function used is the sigmoid function and is given by

$$\text{Sigmoid}(X) = \frac{1}{1 + e^{-x}}$$

Where X is the sum of weighted inputs to that particular node and e is the base of natural logarithms $e = 2.718 \dots$

The output layer of the network is designed according to need of the application output. Since the output of the neural network is expected to predict the presence or absent of the hypertension. It is assume that the



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

actual output is D_k and expected output is Y_k . So the network error function E will be calculated by the expression

$$E = \frac{1}{2} \left\{ \sum_{k=1}^m (Y_k - D_k)^2 \right\}$$

Where m is the number of neurons in the output layer. So if output is 1 hypertension is present and if it is 0 hypertension is absent (Golden, 1996)

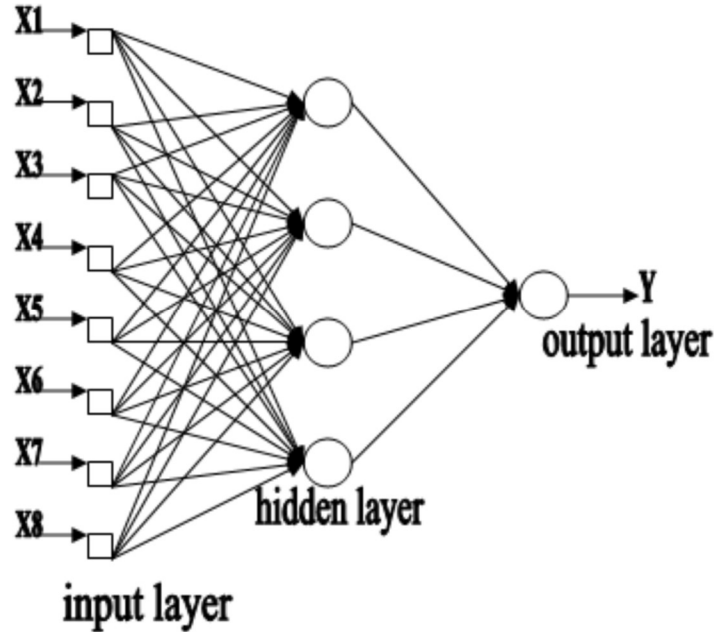


Fig. 1: Artificial Neural Network Architecture

RESULT AND DISCUSSIONS

Discriminant Analysis Model Performance

After development of any predictive model, the most important and appropriate task ahead is to check the usefulness of the model and how significance the independent variables are related to the dependent variable. A discriminant model is declared as useful or good model through model test of "Wilks' lambda", which is a measure of the usefulness of the model. The smaller significance value indicates that the discriminant function does better than chance at separating the groups. Here, Wilks' lambda test has a probability of <0.001 which is less than the level of significance of 0.05, means that predictors significantly discriminate the groups.

Table 1: Wilk's Lambda Test: Checking Usefulness of the Derived Model

Test Function(s)	of Wilks' Lambda	Chi-square	Df	Sig.
1	.862	43.543	8	.000



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

Here, Wilks' lambda test shows a probability of <0.000 which is less than the level of significance of 0.05, indicating that, the predictors strongly discriminate the groups.

Importance of Independent Variables

After checking the usefulness of the derived model, the next task is to identify the importance of independent variables in the predicted model. It can be identified from the "Structure Matrix" and "Standardized Discriminant Functions" in Table 2.

Table 2: Discriminant Analysis model

PATIENT ATTRIBUTES	FUNCTIONS	
	STRUCTURE MATRIX	STANDARDIZED DISCRIMINANT COEFFICIENTS
AGE	0.305	0.547
CIGARETTE SMOKING	0.040	-0.054
STRESS	-0.051	-0.050
ALCOHOLISM	0.093	0.145
SEDENTARY LIFE STYLE	0.133	0.184
FAMILY HISTORY	0.834	0.950
COMORBIDITY	0.057	0.023
ARTERIOSCLEROSIS	-0.057	-0.032

The interpretation of discriminant coefficients (or weight) is like that of multiple regressions. The standardized canonical discriminant function coefficients provide an index of the importance of each predictor variable like the standardized regression coefficients (beta's) did in multiple regression. The variables "Age of the patient" and "Family history" with large coefficients stand out as those that strongly predict allocation to the hypertensive or not hypertensive group. Also, table above provides "Structure Matrix" which is another way of indicating the relative importance of the predictors which holds the same pattern as standardized canonical discriminant function coefficients did. Many researchers use the structure matrix because they are considered more accurate than the standardized canonical discriminant function coefficients. Structure matrix shows the correlation of each variable with discriminate function. These Pearson coefficients are structure coefficients or discriminant loadings, they serve like factor loadings in factor analysis. Generally, just like factor loadings, 0.30 is seen as the cut-off between important and less important variables. Based on the structure matrix above, the predictor variables strongly associated with the discriminant model, are the "Age of the Patient" and "Family History". These predictors possess the loadings of 0.30 or above, which makes them the most important independent variables in the predicted Discriminant model.

Logistic Regression Model Performance

As stated earlier, after model development, the next mission is to check the usefulness of model. This can be done by the use of Omnibus chi-square test. Based on the significant test of model coefficient of Omnibus test in table below:



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

Table 3: Omnibus test of model coefficient

		Chi-square	df	Sig.
Step 1	Step	45.740	8	.000
	Block	45.740	8	.000
	Model	45.740	8	.000

The null hypothesis is that the model coefficient are not statistically significant is rejected since the probability of the model chi-square is <0.001 , less than or equal to the level of significance of 0.05. Also a significance in the Omnibus test indicates that, there is difference between the model with only a constant and the model with independent variables. Since the existence of a relationship is confirmed, the usefulness of the model is affirmed.

Importance of Independent Variables:

In the quest of identifying the importance of predictor variables in Logistic Regression Model, the statistical significance test of Wald Statistic is applied. The presence of a relationship between the dependent variable and each of the independent variables is examined using this test. Here, the null hypothesis is that the b coefficient for the particular independent variable is equal to zero.

Table 4: Logistic Regression Model: Important Variables Identified By the Logistic Regression Model

		B	S.E.	Wald	Df	Sig.	Exp(B)
Step 1 ^a	Age	.029	.009	11.319	1	.001	1.029
	CigaretteSmoking	-.270	.772	.122	1	.727	.764
	Stress	-.339	1.278	.070	1	.791	.712
	Alcoholism	.912	1.057	.744	1	.388	2.488
	SedentaryLifeStyle	1.016	.999	1.035	1	.309	2.763
	FamilyHistory	1.846	.332	30.967	1	.000	6.333
	Comobidities	.153	1.431	.011	1	.915	1.166
	Arteriosclerosis	-.230	1.041	.049	1	.825	.795
	Constant	-2.324	1.385	2.817	1	.093	.098

According to Table 4 the independent variables with the probability of Wald Statistics less than or equal to the level of significance of 0.005 holds statistically significant relationship with the dependent variable. Here, the significant important independent variables are Age and Family History of the Patient, while the rest of the independent variables shows insignificance and not important as far as Wald Statistic and the model developed above.

Artificial Neural Networks Model Performance

Performance of Neural Network Model can be checked through model summary table displays by SPSS Software. The model summary table gives information about the results of the neural network training. For example, the error function that the network tries to minimize during training and percentage incorrect prediction.



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

Table 5: Model Summary: Checking Usefulness of the Model

Training	Cross Entropy Error	109.573
	Percent Incorrect Predictions	9.1%
	Stopping Rule Used	1 consecutive step(s) with no decrease in error ^a
	Training Time	0:00:00.286
Testing	Cross Entropy Error	40.277
	Percent Incorrect Predictions	11.7%

From Table 5, the percentage incorrect prediction is equivalent to 9.1% in the training samples. So, percentage of correct prediction is nearer to 90.9%, which is quite high to adopt the usefulness of the model.

Importance of Independent Variables

In neural network, identifying important predictors can be done through the following table, which computes importance and the normalized importance of each predictor in determining the neural network. The analysis is based on the training and testing samples, the importance of an independent variable is a measure of how much the network's model predicted value changes from different values of the independent variable.

Table 6: Independent Variable Importance: Important Predictors Identified By The Neural Networks Model

	Importance	Normalized Importance
CigaretteSmoking	.052	14.0%
Stress	.150	40.6%
Alcoholism	.138	37.4%
SedentaryLifeStyle	.045	12.2%
FamilyHistory	.183	49.4%
Comobidities	.057	15.3%
Arteriosclerosis	.005	1.2%
Age	.370	100.0%

From Table 6, the normalized importance is just the importance value divided by the leading importance value and expressed as percentage. It can be seen from the table that "Age of the Patient" contributes most in the neural network model construction, followed by "Family History", "Stress", "Alcoholism", "Comorbidity". The above mentioned predictors have supreme impact on how the network classifies the prospect of a patient status. But it is not possible to identify the direction of the relationship between the variables. This is one of the famous limitations of the neural network which conventional models will be of assistance in this situation.

CONCLUSION

It is extensively known that the effectiveness of any model is basically dependent on the characteristics of data used to fit the model, suitable predictor variables selection is one of the conditions for successful development of any model. This study is no exception; several considerations regarding predictor variables selection are reviewed. Risk factors of hypertension identified in this study were suggestive and not a



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

causal, however the study has some positive findings. The study found that, among age, cigarette smoking, stress, alcoholism, sedentary life style, family history, comorbidity, arteriosclerosis; age and family history are the most important predictors. Moreover, based on neural networks analysis, the model identified not only age and family history but also stress, alcoholism and comorbidity as important predictors also.

REFERENCES

Afifi A, Clark VA, and May S. (2004) *Computer-Aided multivariate Analysis*. Fourth edition, Champman and Hall/CRC

Ashish Aneja, Fadi El Atat et al (2004) Hypertension and Obesity: Recent progress research. *Endocrine Society Journal* 59:169-205

Badr EA, Nasr GE, and Joun C. (2003) Back propagation neural networks for modeling gasoline consumption. *Energy Conversion and Management*, 44, 893-905

Cohen Jacob, Cohen Patricia, West Steven G, Aiken Leona S. (2002) *Applied Multiple Regression/Correlation Analysis for the Behavioral Sciences* (3rd edition) Routledge. ISBN 978-0-8058-2223-6

Draper DR, and Smith H. (1998) *Applied Regression Analysis, third edition*. Wiley New York

Francesca Viazzi (2006) Predicting cardiovascular risk using artificial neural network in primary hypertension. *Lippincott Williams and Wilkins*.

Golden RM (1996) *Mathematical methods for neural network analysis and design*. USA: Massachusetts Institute of technology

Haykin S. (1999) *Neural network: A Comprehensive foundation* upper saddle river NJ. (2nd ed.) New York Prentice-Hall

Hosmer David W., Lemeshow Stanley (2000) *Applied Logistic Regression* (2nd edition) Wiley. ISSN 0-471-35632-8

Hyman DJ, Pavlik VN (2001). Characteristics of patients with uncontrolled hypertension in the United States. *N Engl J Med*. 2001; 345:479-486.

Isezuo SA, Opara TC. (2000) Hypertension awareness among Nigerian hypertensive in a Nigerian tertiary health institution. *Sahel Medical Journal*. 2000; 3:93-97

Menard Scott W. (2002) *Applied Logistic Regression* (2nd ed.) SAGE. ISBN 978-0-7619-2208-7

Meyers LS, Gamst GC, and Guarino AJ (2005) *Applied Multivariate Research: Design and Interpretation*, Sage Publication

Salako BL, Ajose FA, Lawani E. (2003) Blood pressure control in a population where antihypertensive are given free. *East Afr Med J*. 2003; 80:529-531

Setyawati B. R., Sahirman S., and Creese R.C. (2002) Neural Networks for cost estimation. *AACE International Transaction EST* 13, 13.1-13.8



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)

Smith A. E., and Mason A. K. (1997) Cost estimation modeling: Regression versus neural network. *The Engineering Economist*, 42(2), 137-161

Sumathi B., Santhakumar A. (2010) Artificial Neural Network approach for predicting hypertension. *International Conference on Computer Applications and Information Technology* 2010

Tesfaye F., Nawi NG *et al* (2007) Association between body mass index and blood pressure across three populations in Africa. *Journal of Human Hypertension*

Ukoli FA, Bunker CH *et al* (1995) Body fat distribution and other anthropometric blood pressure correlates in a Nigerian urban elderly population. *Central African Journal of Medicine*; 41(5): 154-61

Warner B., and Misra M. (1996) Understanding neural networks as statistical tools. *The American Statistician*, 50(4), 284-293.

World Health Organization (WHO 2003)/ International Society of Hypertension (ISH). Statement of Management of Hypertension. *Journal of Hypertension* 2003; 21:1983-1992.

World Health Organization (WHO 2010). The World Health Report: Reducing risk, Promoting Healthy Life. Geneva: World Health Organization; 2002 [cited 2010 Mar 31]. Available: http://www.who.int/entity/whr/2002/en/whr02_en.pdf



All rights reserved

This work by [Wilolud Journals](#) is licensed under a [Creative Commons Attribution 3.0 Unported License](#)