



A Predictive Decision Support Model to Address Deindustrialisation and Trade Deficits in Zimbabwe

Simbarashe A. Nyamusa¹; Mr. A. Ndlovu²

^{1,2}Department of Software Engineering, Harare Institute of Technology, Zimbabwe
Email: ¹ simbarashenyamusa@gmail.com; ² Andlovu@hit.ac.zw

DOI: <https://doi.org/10.47760/ijcsmc.2025.v14i06.011>

Abstract: Zimbabwe's economic landscape is challenged by a persistent trade deficit, heavy import reliance, and deindustrialisation [1]. This research proposes a Predictive Decision Support Model (PDSM) leveraging Artificial Neural Networks (ANNs) to aid policy-makers in formulating data-driven strategies for industrial growth. By forecasting key economic indicators and identifying priority sectors for investment, the Model supports reindustrialisation and trade deficit reduction.

Keywords: Machine Learning, Artificial Neural Networks, Economic Forecasting, Trade Deficit, Zimbabwe Industrial Growth

I. INTRODUCTION

A. Background

Zimbabwe continues to face economic instability, characterised by inflation, currency volatility, and a widening trade deficit. In the first quarter of 2023, the trade deficit grew by 26%, from US\$491 million in Q4 2022 to US\$621.8 million. This was driven by a 23.8% drop in exports, mainly due to falling international prices for key minerals such as PGMs, and reduced gold and tobacco output, while imports remained high, particularly for essentials like fuel, maize, and wheat. Power shortages further impacted local production, especially in the mining sector [2]. These conditions underscore the urgent need for predictive tools and innovative policies to support industrial resilience and reduce economic vulnerability.

B. Statement of the Problem

Zimbabwe's economic challenges are deeply rooted in its reliance on imported goods, leading to a significant trade deficit. The lack of robust local industrial capacity exacerbates this issue, resulting in deindustrialisation and economic stagnation. The absence of effective decision-making tools hampers policymakers' ability to implement strategies that could revitalize the industrial sector and reduce dependence on imports.

C. Proposed Solution:

To address Zimbabwe's economic challenges, we propose a **predictive decision support model** that leverages historical economic data and machine learning. The model will analyze key economic indicators to guide trade planning and prioritise industrial development. By enabling data-driven decisions, it aims to reduce import dependency and enhance local industrial output.

II. LITERATURE REVIEW

A. Introduction

In this chapter we look at the available technologies as well as the progress that has been made in the area of industrial advancement. We will discuss the work that is related to the previous work that has been done and the contributions made. This chapter also reviews the relevant literature and establishes the research framework for the study.

A review of the following is going to be carried out:

- Products in demand in the nation
- Products that can be exported with hope to boost the national income
- Industries that should be registered that have a hand on the national income
- Calculations of the following and other economic indicators
 - GDP
 - BOP
 - Export vs. Import
 - GDP per Capita
 - Ideal Growth Rate
- Similar applications that have been developed and their limitations
- Most appropriate tools and methods that can be used to build the application

B. Products on Demand

As stated by the Zimbabwean government, [4] the basic commodities that should be available to everyone are as follows: cooking oil, mealie meal, bread, flour, sugar, rice, salt, chicken, eggs, beef fresh milk, laundry soaps and washing powder. Cement, fuel and energy where also stated as products on demand.

C. Products that are being exported and imported

Zimbabwe also produces products that can be exported to boost the national income; these products come from the mining and agricultural industries. Following a construction from 1998 to 2008, the economy recorded real growth of more than 10% per year in the period 2010 – 2013, before falling below 3% in the period 2014 - 17 due to poor harvest, low diamond revenue and decreased investment. The minerals that are exported by Zimbabwe are gold, platinum, diamond, nickel and coal. [5]

1) The agricultural products that are exported

- Tobacco, manufactured substitutes: \$893.1 million (22.1%)
- Cotton: \$75.4 million (1.9%)
- Sugar, sugar confectionery: \$44.5 million (1.1%)
- Raw hides, skins not fur skins, leather: \$36.8 million (0.9%)

2) The mineral products that are exported

- Gems, precious metals: US\$1.3 billion (33.4% of total exports)
- Nickel: \$525.3 million (13%)
- Ores, slag, ash: \$497.2 million (12.3%)
- Iron, steel: \$251.8 million (6.2%)
- Salt, sulphur, stone, cement: \$55 million (1.4%)
- Mineral fuels including oil: \$40.8 million (1%)

3) Calculations of economic indicators

GDP

Gross domestic product (GDP) is a standard measure of a country's economic health and an indicator of its standard of living. Also, GDP can be used to compare the productivity levels between different countries.

The biggest advantage of GDP is that calculations of the measure are fairly uniform from country to country. Thus, a comparison between countries provides a high level of accuracy. Furthermore, GDP indicates economic expansion or compression and the growth or decline of an economy.

Different organizations such as the World Bank, the United Nations, and the International Monetary Fund collect and publish data on GDP estimates in all countries.

Gross Domestic Product (GDP) = $C + G + I + NX$ where C = private consumption, G is the sum of government spending, I is the sum of the country's investment (including business capital expenditure), NX is the nation's total net export calculated as total exports minus total imports.

$NX = \text{Exports} - \text{Imports}$

GNP

Gross National Product (GNP) is a measure of the value of all goods and services produced by a country's residents and businesses. It estimates the value of the final products and services manufactured by a country's residents, regardless of the production location.

Gross National Product takes into account the manufacturing of tangible goods such as vehicles, agricultural products, machinery, etc., as well as the provision of services like healthcare, business consultancy, and education. GNP also includes taxes and depreciation. The cost of services used in producing goods is not computed independently since it is included in the cost of finished products.

Gross National Product (GNP) = $C + I + G + X + Z$,

Where:

C – Consumption Expenditure

I – Investment

G – Government Expenditure

X – Net Exports (Value of imports minus value of exports)

Z – Net Income (Net income inflow from abroad minus net income outflow to foreign countries)

Alternatively, the Gross National Product can also be calculated as follows:

$GNP = GDP + \text{Net Income Inflow from Overseas} - \text{Net Income Outflow to Foreign Countries}$

Where:

$GDP = \text{Consumption} + \text{Investment} + \text{Government Expenditure} + \text{Exports} - \text{Imports}$

BOP

The Balance of Payments is a statement that contains the transactions made by residents of a particular country with the rest of the world over a specific time period. It is also known as the balance of international payments and is often abbreviated as BOP. It summarizes all payments and receipts by firms, individuals, and the government. The transactions can be both factor payments and transfer payments. There are two accounts in the BOP statement: the Current Account and Capital Account.

Current Account

The four major components of the Current account are as follows:

- Visible trade – This is the net of export and imports of goods (visible items). The balance of this visible trade is known as the trade balance. There is a trade deficit when imports are higher than exports and a trade surplus when exports are higher than imports.
- Invisible trade – This is the net of exports and imports of services (invisible items). Transactions mainly consist of shipping, IT, banking and insurance services.
- Unilateral transfers to and from abroad – These refer to payments that are not factor payments. For example, gifts or donations sent to the resident of a country by a non-resident relative.
- Income receipts and payments – These include factor payments and receipts. These are generally rent on property, interest on capital, and profits on investments.

Capital Account

The Capital account is used to finance the deficit in the current account or absorb the surplus in the current account. The three major components of the Capital account:

- Loans to and borrowings from abroad – These consist of all loans and borrowings given to or received from abroad. It includes both private sector loans, as well as public sector loans.
- Investments to/from abroad – These are investments made by non-residents in shares in the home country or investment in real estate in any other country.
- Changes in foreign exchange reserves – Foreign exchange reserves are maintained by the central bank to control the exchange rate and ultimately balance the BOP.

Export vs. Import

Imports are the goods and services that are purchased from the rest of the world by a country's residents, rather than buying domestically produced items. Imports lead to an outflow of funds from the country since import transactions involve payments to sellers residing in another country.

Exports are goods and services that are produced domestically, but then sold to customers residing in other countries. Exports lead to an inflow of funds to the seller's country since export transactions involve selling domestic goods and services to foreign buyers.

GDP per Capital

The gross domestic product per capita, or GDP per capita, is a measure of a country's economic output that accounts for its number of people. GDP per capita acts as a metric for determining a country's economic output per each person living there. Often times, rich nations with smaller populations tend to have higher per capita GDP

To calculate GDP per Capita the formula is

$$\text{GDP per capita} = \frac{R}{C} \\ = \text{Gross domestic product} / \text{population}$$

D. Similar Applications that have been developed

• **Predicting economic growth using Machine Learning[6]**

In their research they highlighted factors that affect the economy.

Geography. Absolute latitude (distance from the equator) are negatively correlated with growth and certain regions such as sub-Saharan Africa and Latin America under-perform, on average. This then affects the agricultural industry since some climates are not favourable to some crops.

Political institutions. Measures of institutional quality like strong Rule of Law, Political Rights, and Civil Liberties improve growth, while instability measures like Number of Revolutions and Military Coups and War impede growth.

Religion. Predominantly Confucius/Buddhist and Muslim countries grow faster, while predominantly Protestant and Catholic grow more slowly.

Market distortions and market performance. Real Exchange Rate Distortions and Standard Deviation of the Black-Market Premium correlate negatively with growth.

Investment and its composition. Equipment Investment and Non-Equipment Investment are both positively correlated with growth.

Dependence on primary products. Fraction of Primary Products in Total Exports are negatively correlated with growth while the Fraction of gross domestic product (GDP) in Mining is positively correlated with growth.

Trade. A country's Openness to Trade increases growth.

Market orientation. A country's Degree of Capitalism increases growth.

Sala-i-Martin's findings are standard in the growth literature. His econometric techniques cull the immense proliferation of explanatory variables into a tractable and parsimonious list. However, there are several problems with his approach. First, many of the findings say nothing about why these variables matter, or which matter more than others. In fact, we note that if a country's GDP has a large Fraction

of Primary Products in Total Exports it is likely to be a growth laggard, though if it has a high Fraction of GDP in Mining, it is in the high growth category.

This sort of contradiction suggests that maybe the Sala-i-Martin list is not parsimonious enough. It is certainly not completely amenable to good theoretical explanations. Second, econometric techniques focusing on growth and identifying the correlates of growth are not useful for predicting economic growth. Therefore, it is not possible to know whether certain theoretically ordained variables for explaining growth matter more than others. Machine Learning techniques, however, can create parsimonious lists of variables by focusing on the predictive power of variables. This approach can therefore distinguish between different theories of growth, as well as create better models of growth by whittling down the list of variables to those that are the best predictors of growth. Furthermore, Machine Learning has the advantage of not requiring any major assumptions about a variable's underlying distribution, or indeed any prior assumptions about theoretical links.

The predictive power of Machine Learning techniques has practical benefits as well. The policy maker, for instance, mainly needs to know the effect of a current change in policy on a future (out-of-sample) target. From the policy maker's perspective, a parsimonious list from the usual econometric techniques itself does not provide good policy levers for increasing economic growth because existing studies that use these techniques neglect the issues of cross-validation and out-of-sample predictability. Thus, the policy maker has no idea whether moving a country towards a more capitalist direction, for example, will lead to higher growth because the econometric techniques used in finding these sorts of correlations have not usually been validated or tested for out-of-sample predictive ability. Machine Learning approaches address this problem by emphasizing both cross-validation as well as out-of-sample prediction scores for different specifications of the model.

Moreover, policy makers can also get a sense of the relative importance of variables in predicting growth. Thus, they can prioritize policy levers according to which ones may have the greatest impact on economic growth. Of course, both the policy maker and the academic remain concerned about causal links. However, the results of standard econometric techniques designed for establishing causal links, do not satisfy these concerns. This is because the growth literatures focus on growth accounting, and therefore on the correlates of growth, essentially end up generating long lists of possible correlates of growth. Such lists hamper standard econometric techniques since they are plagued by endogeneity problems. Of course, finding parsimonious lists of correlates of growth a la Sala-i-Martin certainly help in specifying econometric models to explore causal links. However, econometric techniques that attempt to solve the problem of endogeneity by using instrumental variables are problematic because the consistency of the estimated marginal effects depends critically on the strength and validity of the instruments. In fact, some of these instruments may actually lead to biased estimates (Bazzi & Clemens, 2013).

Further, it is even possible for more sophisticated parametric methods to do worse than the traditional ordinary least square method. In contrast, Machine Learning techniques can help in this process by winnowing down the list of possible correlates by focusing on how well these variables predict out-of-sample. Thus, variables that appear to be robust correlates of growth but do not predict well out-of-sample, cannot really be causal variables, and can be eliminated. In this sense, Machine Learning can be helpful in exploring causal links to growth (Athey & Imbens, 2015). Another problem in the growth literature is the paucity and unreliability of data for precisely the countries for which growth issues matter most. Standard statistical analyses do not perform well when there is missing data. Machine Learning can address this problem in a scientifically verifiable way by finding "surrogate" variables that can proxy those with missing data. These proxies are chosen by the Machine Learning techniques by their predictive abilities, and to that extent, provide a hard test for the usefulness of a particular proxy variable. Lastly, to the extent that econometric techniques focus on point parametric estimates, they ignore potential non-linear ways in which explanatory variables can affect economic growth. Machine Learning techniques, on the other hand, can provide easily understandable graphs by capturing non-linear "marginal" effects of a particular variable on growth. Thus, we are able to say with some predictive accuracy whether over a certain range, a particular variable has a greater or lesser, negative or positive impact on growth, as well as identify other ranges where the same variable does not affect growth.

They used empirical methods as follows regression analysis, artificial neural network, bootstrap aggregating, boosting, and random forest predictors:

Bootstrap Aggregating (Bagging) Predictor

The bagging predictor proposed by Breiman (1996) takes random resamples $\{L(M)\}$ from the learning sample with replacement to create M samples using only the observations from the learning sample. Each of these samples will contain N observations – the same as the number of observations in the full training sample. However, in any one bootstrapped sample, some observations may appear twice (or more), others not at all. The bagging predictor then adopts the rules for splitting and declaring nodes to be terminal described in the previous section to build M classification trees. To complete steps (3) and (4), the bagging predictor needs a way of aggregating the information of the predictions from each of the trees. The way that bagging predictors (and other ensemble methods) do this for class variables is through voting. For classification trees (categorical target variables), the voting processes each observation through all of the M trees that was constructed from each of the bootstrapped samples to obtain that observation's predicted class for each tree. The predicted class for the entire model, then, is equal to the mode prediction of all of the trees. For regression trees (continuous target variables), the voting process calculates the mean of the predicted values for all of the bootstrapped trees. Finally, the predictor calculates the redistribution estimate in the same way as it did for the single classification tree, using the predicted class based on the voting outcome for each predictor.

Random Forests

Like the bagging predictor, the random forest predictor is a tree-based algorithm that uses a voting rule to determine the predicted class of each observation. However, whereas the bagging predictor randomizes the selection of the observations into the sample for each tree, and then builds the tree using the same procedure as CART, the random forest predictor may randomize over multiple dimensions of the classifier (Breiman, 2001). The most common dimensions for randomizing the trees are the selection of the inputs for node of each tree, as well as the observations included for constructing each of the trees. We briefly describe the construction of the trees for the random forest ensemble below. A random forest is a collection of tree decision rules, $\{d(x, \Theta_m), m = 1, \dots, M\}$, where Θ_m is a random vector specifying the observations and inputs that are included at each step of the construction of the decision rule for that tree. To construct a tree, the random forest algorithm takes to following steps:

- i. randomly select $n \leq N$ observations from the learning sample;
- ii. At the “root” node of the tree, select $k \in K$ inputs from x ;
- iii. Find the split in each variable selected in (ii) that minimizes the mean square error at that node and select the variable/split that achieves the minimal error;
- iv. Repeat the random selection of inputs and optimal splits in (ii) and (iii) until some stopping criteria (minimum improvement, minimum number of observations, or maximum number of levels) is met.

The bagging predictor described in the previous sub-section is in fact a special case of the random forest estimator where, for each tree, Θ_m consists of a random selection of $n = N$ observations from the learning

sample with replacement (and each observation having a probability of being selected in each draw equal to $1/N$) and sets the number of inputs to select at each node, k , equal to the full length of the input vector, K so that all of the variables are considered at each node.

- Study on Modelling and Forecasting Inflation in Zimbabwe using Generalized Autoregressive Conditionally Heteroskedastic approach [7]

Data Collection

All the data used in this study was collected from

Testing for Stationarity

The monthly inflation rate series was tested for stationarity using the Augmented – Dickey Fuller (ADF) unit root test as follows:

$$\Delta \text{INF}_t = \mathbf{b}_0 + \sum \mathbf{b}_i \Delta \text{INF}_{t-i} + \text{INF}_{t-1} + \dots$$

Where Δ is the difference operator, b_0 is the drift parameter, β is the coefficient on a time trend, ϵ_t is the error term, and INF_t is monthly inflation at time t .

The null hypotheses of no unit root tests are:

$H_0: = 0$ (if there is a trend [F – test]) and $H_0: = 0$ (if there is no trend [t – test]). If the null hypothesis is rejected, like in this case; the data are stationary and can be used without differencing and if the null hypothesis is accepted, there is a unit root; therefore, we ought to difference the data. Since the probability (p) value was found to be equal to 0.0006835, the null hypothesis was rejected at 1% level of significance and therefore the series is stationary.

Testing for ARCH effects

The Lagrange Multiplier (LM) Test

Run the mean equation or any postulated linear regression of the form:

$$t = 0 + 1\Omega_1 t + 2\Omega_2 t + 3\Omega_3 t + 4\Omega_4 t + zt$$

Where it is the dependent variable, 0 – 4 are estimation parameters of the mean equation, $\Omega 1t - \Omega 4t$ are explanatory variables and z_t is the disturbance term. Save the residuals, $\hat{z}_t$. Square the residuals and regress them on p own lags to test for ARCH of order p, that is; run the regression: $\hat{z}_t^2 = 0 + \hat{\alpha}_1 \hat{z}_{t-1}^2 + \dots + \hat{\alpha}_p \hat{z}_{t-p}^2 + \Upsilon_t$ where Υ_t is an error term. Obtain R^2 from this regression. The test statistic, defined as TR^2 (the number of observations multiplied by the coefficient of multiple correlation); is distributed as a $\chi^2(p)$. The null and alternative hypotheses are:

H0: 1=0 and 2=0 and 3=0 and ... and p=0

H1: $1 \neq 0$ or $2 \neq 0$ or $3 \neq 0$ or $p \neq 0$

In this study, the ARCH / GARCH effects test described above was done and the results are shown below:

[illegible]

The results of the test above indicated that there are (G) ARCH effects in the chosen (optimal) model, since the p – value [which is approximately equal to zero] is significant at 1% level of significance and therefore it is appropriate to estimate a GARCH model.

GARCH (1, 1) Model Results

Variable	Coefficient	Std. Error	z	p – value
0	0.0431191	0.0535928	0.8046	0.4211
AR (λ_1)	0.980045	0.0231912	42.26	0.0000***
d'	0.0371421	0.0188943	1.966	0.0493**

ARCH (α_1)	0.575929	0.109556	5.257	0.000000146***
GARCH (α_1)	0.424071	0.989479	4.286	0.0000182***
$\alpha_1 + \beta_1$	1			

Tab 1.0: GARCH Model Results

The model tabulated above was also presented as follows:

$$INF_t = 0.04 + 0.98 INF_{t-1}$$

$$p: (0.421) (0.000)$$

$$S.E (0.054) (0.023)$$

$$\sigma_t^2 = 0.04 + 0.58 \varepsilon_{t-1}^2 + 0.48 \sigma_{t-1}^2$$

$$p: (0.049) (0.000) (0.000)$$

$$S.E: (0.019) (0.11) (0.99)$$

Strengths and challenges of this method

The estimated model is acceptable simply because the necessary and sufficient conditions (specified in have been met. As theoretically expected, the parameters a_0 and α_1 are greater than zero (0) and β_1 is positive to ensure that the conditional variance σ_t^2 is non – negative and therefore the positivity constraint of the GARCH model is not violated. Thus the estimated GARCH (1, 1) model seems quite good for explaining the behavior of monthly inflation rate volatility in Zimbabwe.

The estimated model can thus be referred to as an IGARCH (1, 1) process. This implies that persistent variance is not inexistent and volatility shocks have a permanent effect and therefore current information remains critical for the forecasts of the conditional variances for all horizons. Caporale *et al* (2003) noted that the high estimated volatility persistence obtained using GARCH models might be attributed to structural changes in the variance process. Tsay (2002) highlights the fact that the actual cause of persistence deserves a careful investigation.

- Study on proposed framework for implementing Artificial Neural Network to enhance decisions in agriculture sector [8]

Research Data and Methods

The data was gathered from Food Security Information Centre which have been gathered from FAO, WHO, CAMPAS, and real studies about the status of food security in the Egyptian governorates about the rate of population growth, the amount of food needed to satisfy the needs of all citizens, the productivity of strategic crops, the rate of imports of each crop with its price.

Artificial Neural Network was used in modelling the needs of the required amounts of food according to the population growth rate. The simplest definition of an artificial neural network (ANN), is a simulation of biological neural system, it is a mathematical model or computational model based on biological neural networks. It consists of an interconnected group of artificial neurons and it processes information using a connectionist approach to computation. Most of times ANN act as an adaptive system that changes its structure based on internal or external or information that flows through the network during the learning phase.

Neural networks are typically organized in layers. Layers are made up of number of interconnected 'nodes' which contain an activation function as shown in Fig below. Patterns are presented to the network via the 'input layer', which communicates to one or more 'hidden layers' where the actual processing is done via a system of weighted 'connections'. The hidden layers then link to an 'output layer'.

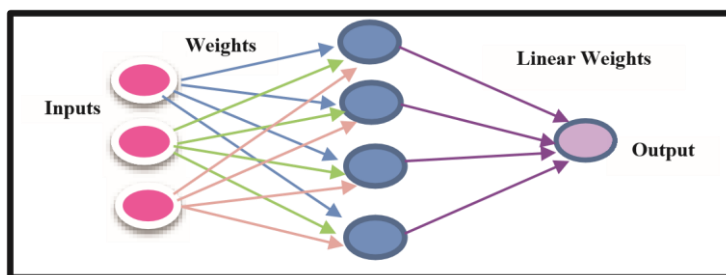


Fig. 1 neural network

Learning steps used:

1. Initialized by setting up all its weights to be small random numbers.
2. Apply input and calculate output that different from target one.
3. Calculate error rate as target – Actual.
4. Used error rate to modify weights to get small error rate.
5. Apply step 2, 3, 4 more and more to get final trusted model.

Data obtained from the adapted framework by building process of Artificial Neural Networks (ANNs) via WEKA using Multilayer Perceptron (MLP) function as a data mining predictive technique showed the difference between the available amount of crops and the needed amount to achieve food security according to the population growth up to year 2020, data set trained from the years (2001- 2011) and predict data from the years (2012- 2020) and visualized as follow:

Wheat

Results shown in **Fig. 2** could explain the amounts of production quantities, imports, losses and other uses, the amount of food available for humans and the amounts required to achieve food security according to the calculated food basket proposed by the National Nutrition Institute (NNI) and FSIC to show the difference between the available amounts and the needs to cover the requirements of wheat. The forecasting results show that Egypt will have a shortage to cover the needs of the Egyptian citizens from production and imports of wheat, which is considered a major obstacle since wheat is one of the key elements of bread production, which is indispensable in the Egyptian houses.

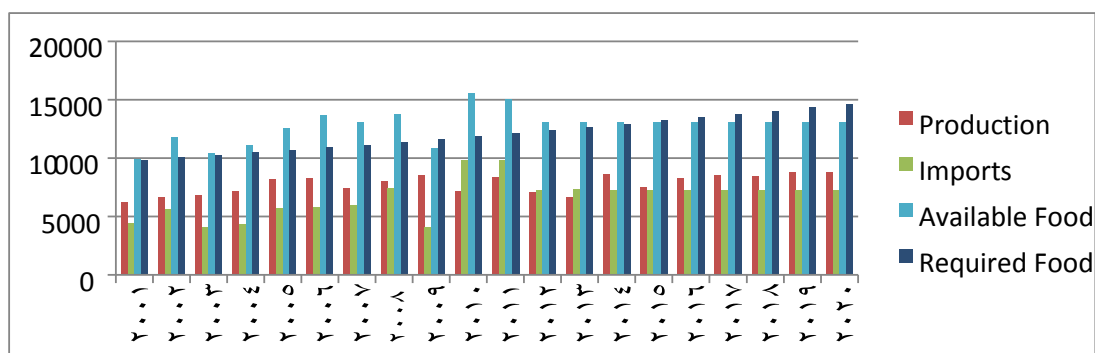


Fig. 2 wheat production

Rice

Results shown in **Fig. 3** could explain the amounts of production quantities, imports, losses and other uses, the amount of food available for humans and the amounts required to achieve food security according to the calculated food basket proposed by the National Nutrition Institute (NNI) and FSIC to show the difference between the available amounts and the needs to cover the requirements of rice. The forecasting results show that Egypt may have a shortage to cover the needs of rice but it could be covered as it is not a large gap between available and required amounts needed to achieve food security to Egyptian citizens.

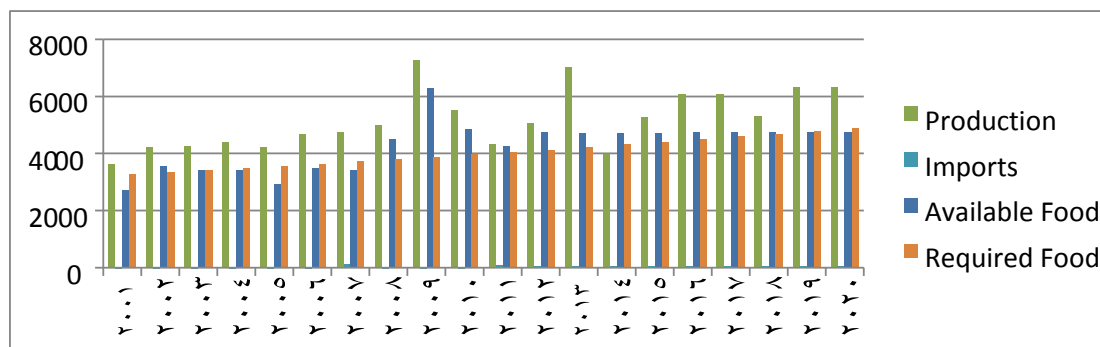


Fig. 3 rice production

• Study on Application of Artificial Neural Networks to Predict Intraday Trading Signals[9]

Research Data and Methods

This research used data taken from the Indonesian Stock Exchange market database. Information on each trading record which was considered useful was trading number, trading date, trading time, price, trade quantity, and transaction volumes.

Researchers focused on the trading data from January 1, 2010 to April 30, 2010 inclusively. The first three-month data was treated as training data set, whereas the following one-month data was used as testing data set to measure the performance of the ANN model.

First, the data was sorted according to the trade number. Then, a macro could run to copy only one leg of the deal and paste it in a separate sheet to ensure that no duplication caused by buying and selling records which are recorded as two instances in the raw data file.

The cleaned data was then split per 15-minutes periods to get the OHLC (Open, High, Low, Close) price for each period, of which then the technical indicators were generated. The technical indicators generated as the inputs were the Price Channel Indicator, the Adaptive Moving Averages, the Relative Strength Index, the Stochastic Oscillator, the Moving Average Convergence-Divergence, the Moving Averages Crossovers and the Commodity Channel Index.

As the model was built with the intention as an agent to intra-day trading predictor, there was a need to balance between the numbers of computation effort with the speed of the execution. Choosing the right amount of hidden layers in an ANN model can be tricky. According to Cybenko, for a continuous function with limited set of discontinuities, a single hidden layer should be sufficient to an ANN model. Because the characteristic of the data which was being used in this study fitted the definition, then only one hidden layer was used in the researchers' model. Figure 4 shows the ANN which was used in the simulation.

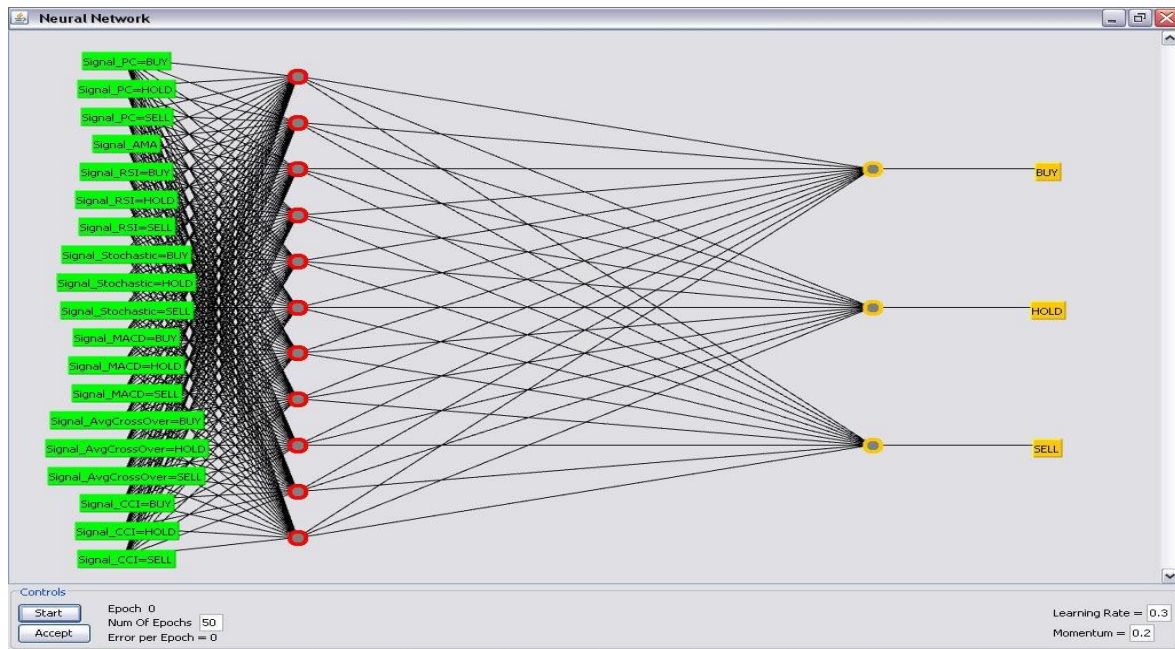


Fig. 4 Artificial Neural Network

The output of the model is computed using the following mathematical expression:

$$y_t = \alpha_0 + \sum_{j=1}^q \alpha_j g \left| \beta_{0j} + \sum_{i=1}^p \beta_{ij} y_{t-i} \right| + \varepsilon_t, \forall t$$

Here y_{t-i} ($i = 1, 2, \dots, p$) are the p inputs and y_t is the output. The integers p, q are the number of input and hidden nodes respectively. α_j ($j = 0, 1, 2, \dots, q$) and β_{ij} ($i = 0, 1, 2, \dots, p; j = 0, 1, 2, \dots, q$) are the connection weights and ε_t is the random shock; α_0 and β_{0j} are the bias terms.

The logistic sigmoid function $g(x) = \frac{1}{1 + e^{-x}}$ was applied as the nonlinear activation function.

The following steps were performed in building the ANN model. Firstly, the training data alone were feed into the model, and then testing was done on the training data itself. By doing so whilst increasing the number of epoch, the model learnt and resulting in the reduced error rate. This was done to obtain a proper number of neurons in the hidden layer while making sure that the ANN still could learn from the data.

Secondly, the training data was split into two parts, with one big part of about 66% of the data was treated as actual training data, and the rest were treated as the testing data. Then again, the model was run while varying the number of epoch, to the point that the error rate was going to increase at a certain point. This point was marked as the over-fitting point, which is a character of ANN modeling. There is no correct formula to calculate the exact point where the overfitting occurs, so it was tested in a trial and error basis.

Finally, after finding the number of epoch which resulted in over-fitting, the model was run using a tenfold cross validation to arrive in a robust model which had sufficient training and ready to run the simulation.

E. Strengths and Challenges of this method

From the experiment results, it was concluded that the artificial neural networks are indeed promising methods to develop forecasting tools to generate trading signals for intraday trading. The ANN model developed in this study performed better than the naïve Buy-and-Hold strategy.

F. Research Gaps

- Lack of ML-based forecasting in Zimbabwean industrial policy.
- Limited decision-support systems tailored to African policy-makers.
- No integration of AI and value chain mapping in trade deficit strategies.

G. Conclusion

Machine learning, when integrated into policy tools, provides actionable insights for industrial advancement. This research aims to close Zimbabwe's industrial policy-tech gap through a predictive system.

III. METHODOLOGY

A. Research Framework

Following a hybrid methodology:

- **Design Science Research (DSR)** for iterative model development.
- **AI/ML** for predictive modelling and forecasting.

B. Data Collection

- Sources: **ZimStats**, World Economic Outlook and Zimbabwe news websites
- Key Variables: GDP, trade volumes, inflation, sector output, energy availability

C. Data processing methods

For this model we will identify and select features most relevant to the problem which are economic indicators for instance exports, imports, gross fixed capital investment

D. Model Development (Classification and prediction)

1) Machine learning

Machine Learning aims to solve the problem of having huge amounts of data with many variables and is commonly used in areas such as pattern recognition (speech, images), financial algorithms (credit scoring, algorithmic trading) (Nuti et al. 2011), energy forecasting (load, price) and biology (tumor detection, drug discovery).

We will build a model using ML techniques to provide an objective approach to finding *linear and non-linear* patterns on *how* publicly-available economic, geographic, and industrial variables predict growth (Hand, Mannila, & Smyth, 2001). First, we will identify the Machine Learning approach that best predicts growth. Then, using the best technique, we will identify the variables that form the pattern that best predicts growth. This is important because, at least theoretically, there may be reason to believe that many of the correlates of growth have complex non-linear impacts of growth because strategic complementarities between variables can lead to multiple possible growth equilibria. For example, openness to trade may improve growth on average, but only if a country is not too dependent on natural resources or other primary commodities that may lead to conflict among competing factions over the rents generated by the resource sector. Moreover, Machine Learning can generate partial dependence plots. These graphs can illustrate how variables identified as good predictors of growth relates (perhaps non-linearly) to growth. Further, by identifying the variables that have the most predictive power we could help develop a framework to distinguish between competing theoretical explanations of growth [10]

2) Supervised Machine Learning

We are going to develop our model based on available input and output data. In this model, every input pattern that is used to train the network is associated with an output pattern, which is the target or the desired pattern. A comparison is made between the network's computed output and the correct expected output, to determine the error. The error can then be used to change network parameters, which result in an improvement in performance.

3) Artificial Neural Networks

Artificial neural network is a computational model which was first developed by McCulloch and Pitts in 1943 (Clarence N. W, 2001). ANN helps in building predictive models which are used for predicting a future phenomenon. Artificial Neural Networks can be classified into two broad categories: feedforward and recurrent networks. An artificial neural network is classified as recurrent if the connections between units (neurons) form a direct cycle. On the contrary, a feedforward NN is a class of NN which considers the transfer of information/messages only in the forward direction (i.e. not in a cycle). That is to say; input signals travel in one direction, which in a sense implies that the output in a particular layer has no effect on the same layer.

We will use Multiple Layer Perceptron Neural Network which is a feed-forward multilayer network architecture composed of several layers of neurons, an input layer, an output layer, and several hidden layers (Haykin, 2008).

4) Multilayer Perceptron (MLP)

A perceptron model without a hidden layer is only efficient when dealing with a linearly separable problem (Clarence N. W, 2001). To overcome this limitation (handle nonlinear problems), MLP is used. A multilayer differs from a perceptron model in the sense that, it takes into account additional layer(s) known as hidden layer(s). A multilayer perceptron is also known as a multilayer feedforward network. For a continuous function defined on a compact domain, a single hidden layer feedforward can approximate; irrespective of the underlying activation function (Hornik et al., 1989).

Example of ANNs is presented in Figure below:

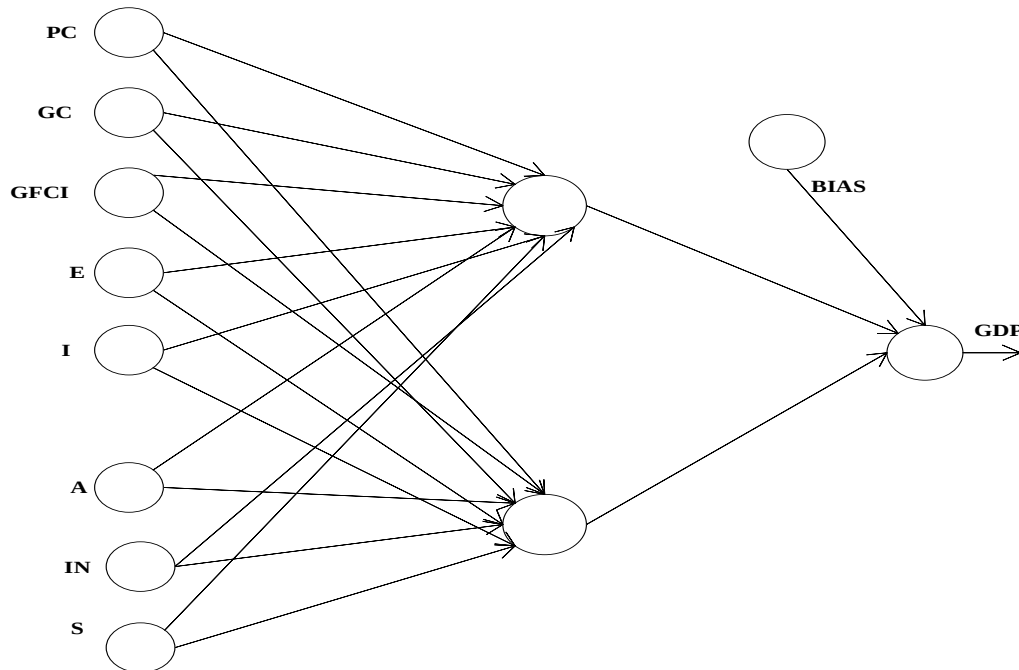


Fig. 5 multilayer perceptron

Where PC is private consumption, GC is government consumption, GFCI is gross fixed capital investment, E is exports, I is imports, A is agriculture, IN is industry and S is services.

E. TRAINING METHODOLOGY

Validation and Testing of Predictive Accuracy

Once we have built our dataset and imputed the missing values, the next issue is to evaluate the validity of our error estimates and the predictive strength of our models. Error estimates ($R[d]$) can sometimes be misleading if the model we are evaluating is over-fitted to the learning sample. These error estimates can be tested out-of-sample and also cross validated using the learning sample. To test the out-of-sample validity, we simply split the full dataset into two random subsets: the first, known as the *learning sample* (or training sample) contains observations that will build the models; the second, known as the *test sample*, will test the out-of-sample predictive accuracy of the models. The out-of-sample error rates will indicate which models and specifications perform best, and will help reveal if any of the models are over-fitted.

To validate the error rates, machine learning uses a method known as cross-validation. Cross-validation creates V random sub-samples of the data, predicts the observations in those sub-samples using the full model, and calculates the error estimate for each of them. It then averages the error estimates of all V random subsamples to give the cross-validated error estimate

F. Back-propagation networks

Though several network architectures and training algorithms are available, the back-propagation (BP) algorithm is by far the most popular. The network is considered a feed-forward and the learning is supervised. Multilayer perceptron's neural networks trained by BP consist of several layers of neurons, interconnections,

and weights that are assigned to those interconnections. Each neuron contains the weighted sum of its inputs filtered by a sigmoid (S-shaped) transfer function. The error of the output relative to the desired output is propagated backwards through the network in order to adjust the synapse weights.

The learning process of back-propagation neural network is actually an error minimization procedure (Rumelhart, Hinton, and Williams, 1986). A typical BPN model uses three vectors: input vector, one or more hidden vectors, and an output vector. After the input and output vectors are fed into the BPN model, the network first selects parameters randomly and processes the inputs to generate a predicted output vector. After calculating the error between its predicted outputs and the observed outcomes, the network adjusts the parameters in ways that will reduce the error, generates a new output vector; calculates the errors, adjust its parameters again, and so on. The iteration or learning process continues until the network reaches a certain specified error. Figure 3 is a sketch of the stages of a typical BPN model with r inputs, two hidden or intermediate units, and one output unit.

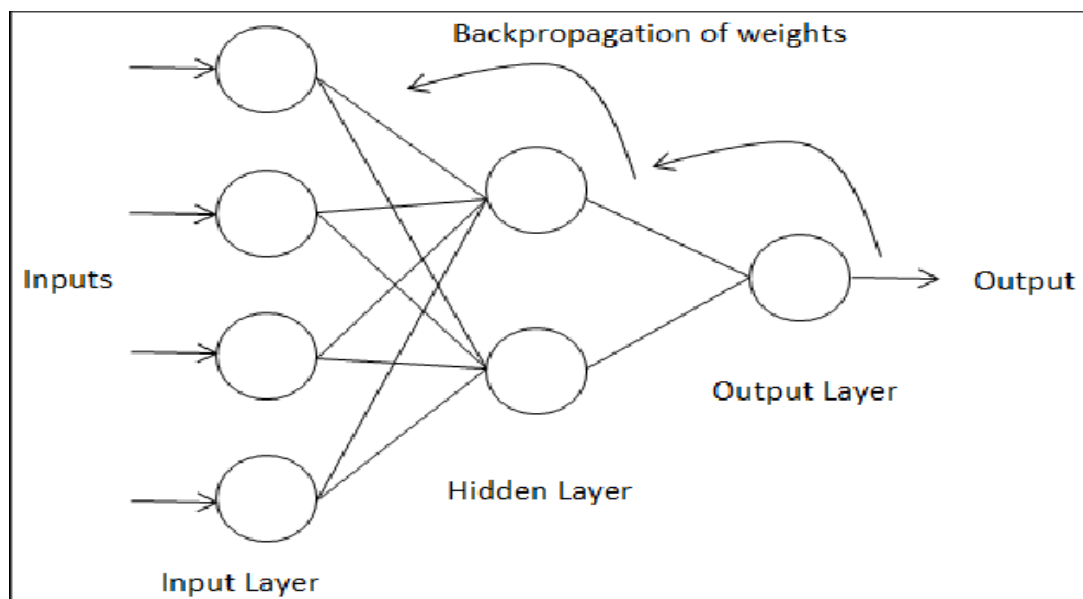


Fig. 6 A Back-Propagation network

G. Model loss

A loss function is used to optimize a machine learning algorithm. The loss is calculated on training and validation and its interpretation is based on how well the model is doing in these two sets. It is the sum of errors made for each example in training or validation sets. Loss value implies how poorly or well a model behaves after each iteration of optimization.

An accuracy metric is used to measure the algorithm's performance in an interpretable way. The accuracy of a model is usually determined after the model parameters and is calculated in the form of a percentage. It is the measure of how accurate your model's prediction is compared to the true data.

Example- Suppose you have 1000 test samples and if your model is able to classify 990 of them correctly, then the model's accuracy will be 99.0%

If the model's prediction is perfect, the loss is zero; otherwise, the loss is greater. The goal of training a model is to find a set of weights and biases that have low loss, on average, across all examples. Higher loss is the worse (bad prediction) for any model.

The loss is calculated on training and validation and its interpretation is how well the model is doing for these two sets. Unlike accuracy, a loss is not a percentage. It is a sum of the errors made for each example in training or validation set.[11]

H. Conclusion

While more traditional models will continue to be the gold standard for teasing out causal effects of individual variables on the margin, we have found that Machine Learning models do a better job than some traditional models for the purpose of predicting. One important function that Machine Learning is that it can be a useful a tool for sifting through the glut of data available to study and predict growth. With the ever-growing number of variables and indices available from governments, the World Bank, private think tanks, and individual researchers, it can be hard to tell which variables serve as the best proxies for a given economic concept. Machine learning also deals more constructively with the problem of missing data.

IV. RESULTS

A. Correlation Analysis

The relationship between variables in the economic environment and the impact to which they are affecting the Gross Domestic Product (GDP) if they are at all. Once we know that two variables are closely related, we can estimate the value of one variable given the value of another.

Correlation Interpretation Table:

Variable Correlation Coefficient	Interpretation Category	Depiction
1 to 0.4	Weak Correlations	Variable has limited capability to affect the target
0.5 to 0.7	Strong positive correlations	Variable can to an extent effect the target
0.7 and above and less than 1	Strong positive correlations	The target is highly dependent on the variable. They mimic each other, as target increases variable also increase and vice-versa
(-0.5) to 0	Weak negative correlations	The variable has limited capability to affect the target
(-0.5) to (-0.7)	Strong Correlations	Variable can to an extent affect the target variable
-0.7 and less and greater than -1	Very strong Negative correlations	Variable with this correlation coefficient oppose the target variable linearly. They increase as the target decrease and vice-versa

Tab. 2 Score Interpretation of the Correlation Analysis

Core industrial Activities

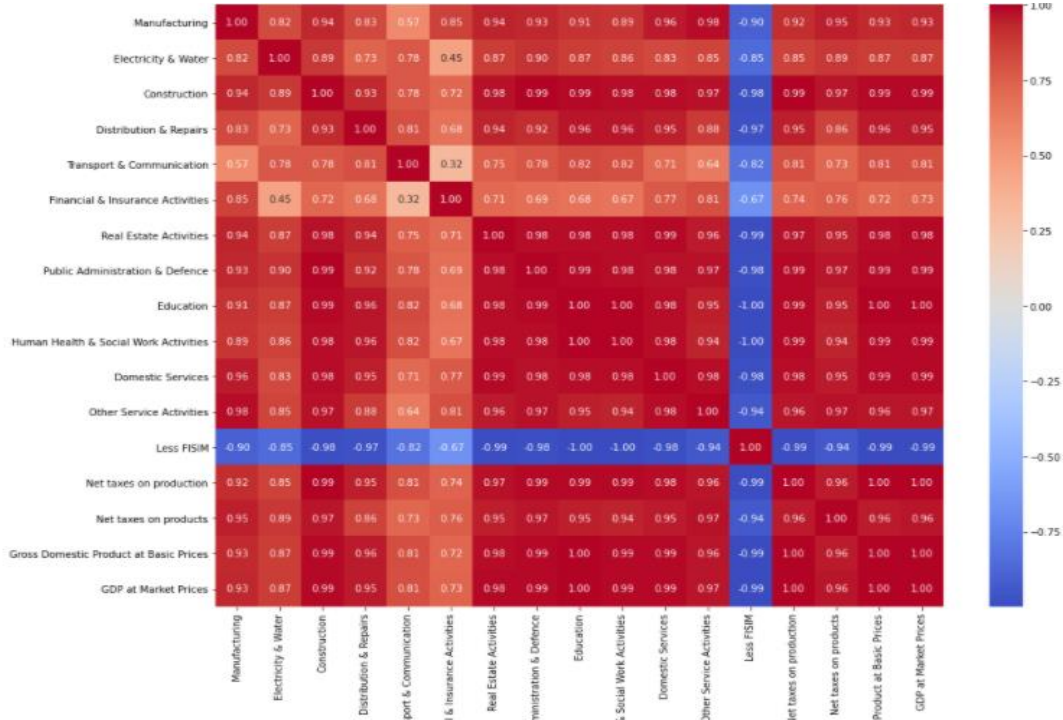


Fig. 7 Correlation of industrial services

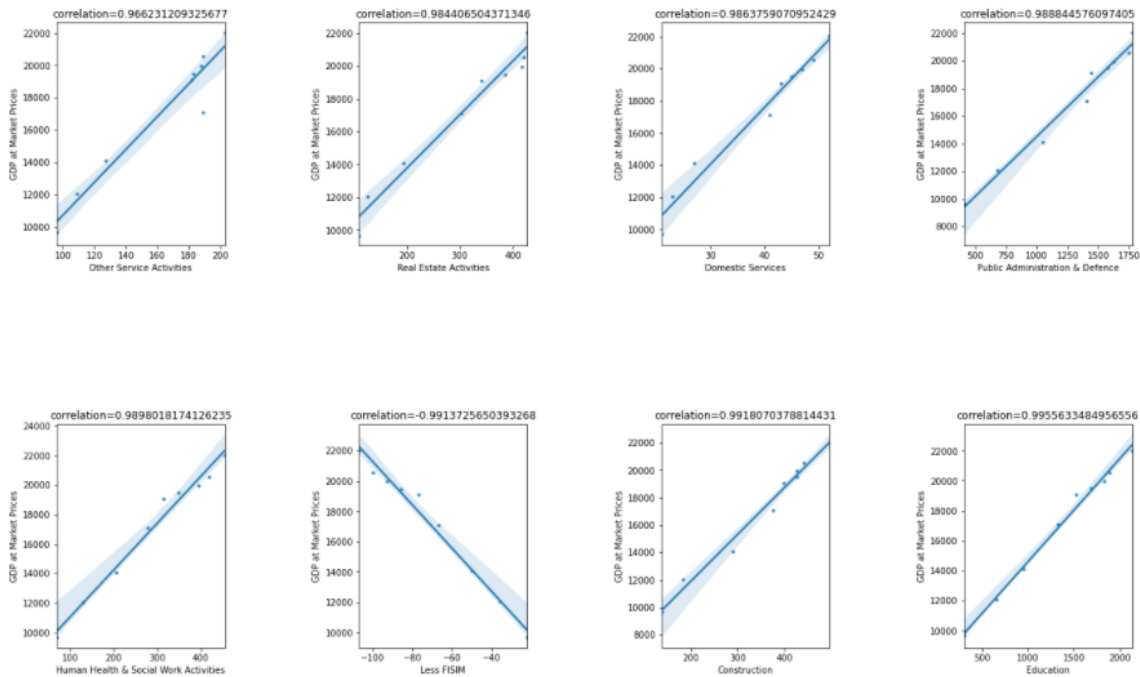


Fig. 8 correlations of industrial activities

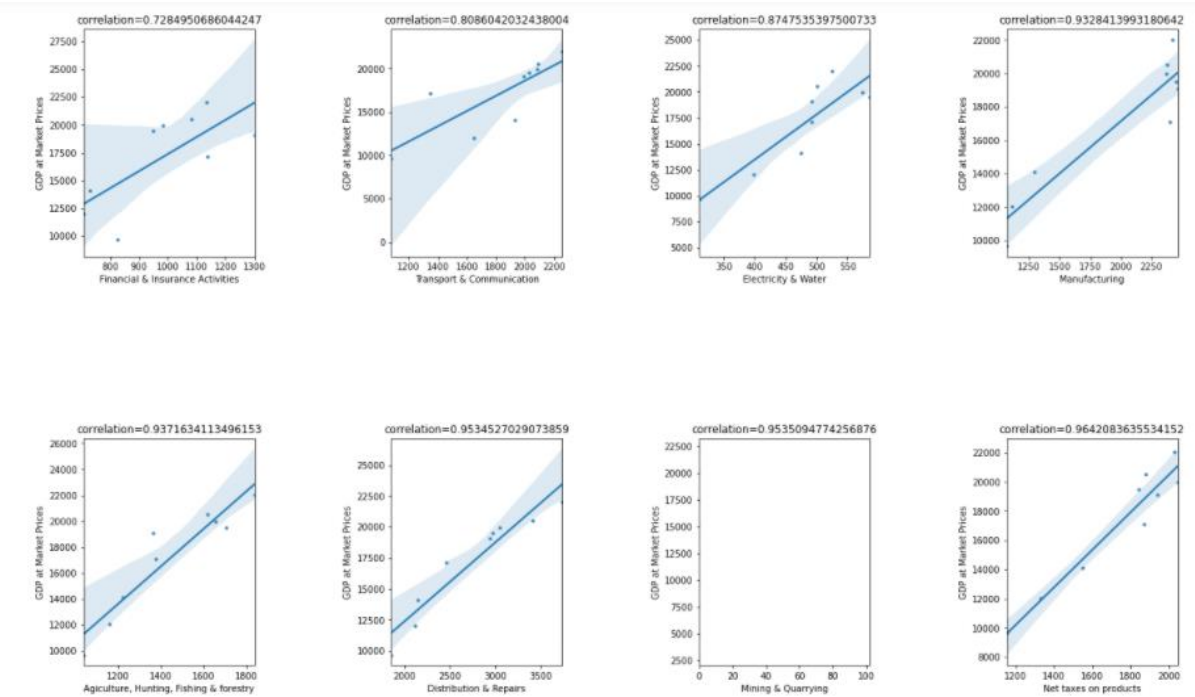
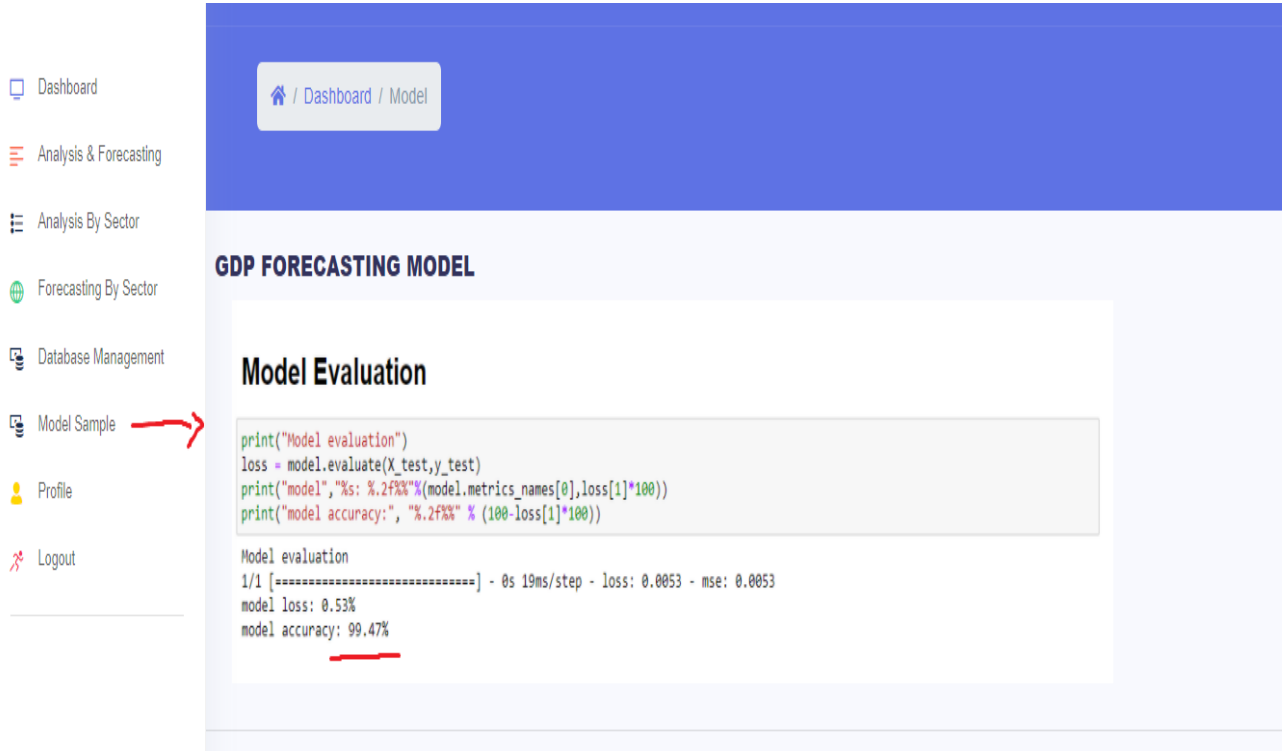


Fig. 9 correlations of industrial activities

B. Forecasting



Dashboard

Analysis & Forecasting

Analysis By Sector

Forecasting By Sector

Database Management

Model Sample

Profile

Logout

Dashboard / Forecasting

PRIMARY SECTOR SECONDARY SECTOR SUPPORTIVE SECTOR

GDP FORECASTING (US\$m)

Year choose year to forecast and enter outputs for previous year in all fields

2023

Agiculture, Hunting, Fishing & forestry

80

Mining & Quarrying

89

Manufacturing

78

Electricity & Water

50

Construction

67454

Distribution & Repairs

45

Transport & Communication

61

Finance & Insurance

78

Real Estate

57

Fig. 10 GDP forecasting

Dashboard

Analysis & Forecasting

Analysis By Sector

Forecasting By Sector

Database Management

Model Sample

Profile

Logout

Dashboard / Forecasting

GDP FORECASTING SECONDARY SECTOR SUPPORTIVE SECTOR

FORECASTING PRIMARY SECTOR

Agriculture forecast for 2023 is: US\$49m

Year choose year to forecast

2023

Choose Primary Activity

Agriculture

Previous Year Output (US\$m)

50

Reset

Forecast

Fig. 11 Primary sector forecasting

References

- [1]. Dr Ngwenya B (2016). Economic Focus: Paying the price of de-industrialization...Its causes and implications [Online]. Available at: www.sundaynews.co.zw/economic-focus-paying-the-price-of-de-industrialisation-its-causes-and-implications/ (Accessed:07/06/2025)]
- [2]. Kanyemba, B., 2023. Zimbabwe Trade Balance Worsens by 26%, What it Means for the Economy. Equity Axis. Available at: <https://equityaxis.net/index.php/post/17496/2023/5/zimbabwe-trade-balance-worsens-by-26-what-it-means-for-the-economy> [Accessed 07/06/2025].
- [3]. Lloyds Bank Trade (2024). Foreign Trade Figures of Zimbabwe.
- [4]. bulawayo24.com 06 Nov 2017
- [5]. Study prepared for FAO by Dr. Moses Tekere (with the assistance of James Hurungo and Masiwa Rusare), Trade and Development Studies Centre
- [6]. James T. Bang, St. Ambrose University, Atin Basuchoudhary, Virginia Military Institute and Tinni Sen, Virginia Military Institute researched on new tools for predicting economic growth using Machine Learning
- [7]. Nyoni, T. (2018). Modelling and Forecasting Inflation in Zimbabwe.
- [8]. Khedr, Kadry, & Walid, 2015
- [9]. I. Putra, Study on the Application of Artificial Neural Networks to Predict Intraday Trading Signals,n.d. https://www.researchgate.net/publication/262172258_Application_of_artificial_neural_networks_to_predict_intraday_trading_signals
- [10].James T. Bang, Atin Basuchoudhary, Tinni Sen 2015 Predicting Economic Growth Using Machine Learning.
- [11].[<https://intellipaat.com/community/368/how-to-interpret-loss-and-accuracy-for-a-machine-learning-model>]
- [12].Trading Economics (2020). Zimbabwe Inflation Rate.
- [13].<https://www.thebalance.com/gross-national-income-4020738>
- [14].<https://corporatefinanceinstitute.com/resources/knowledge/economics/gross-domestic-product-gdp/>
- [15].<https://corporatefinanceinstitute.com/resources/knowledge/economics/gross-national-product-gnp/>
- [16].African Development Bank Group (2024). Zimbabwe Economic Outlook.
- [17].Khedr, A. et al. (2015). Forecasting Agriculture Sector Needs using ANN.
- [18].Putra, R. (n.d.). Application of ANN in Intraday Trading.
- [19].Athey, S. & Imbens, G. (2015). Machine Learning Methods for Estimating Heterogeneous Causal Effects.
- [20].Athey, S., & Imbens, G. (2015). Machine Learning Methods for Estimating Heterogeneous Causal Effects. arXiv Preprints, 1-9. Retrieved from <http://arxiv.org/pdf/1504.01132v2.pdf>
- [21].Khedr et al, 2015. Financial time series forecasting using support vector machines. Neurocomputing 55 (1-2), 307–319.
- [22].Kurniady and Kosala, 2011. Forecasting with artificial neural networks: The state of the art. International Journal of Forecasting 14 (1), 35–62.

- [23].An Introductory Study on Time Series Modeling and Forecasting Ratnadip Adhikari R. K. Agrawal. (n.d.).
- [24].Madhunguyo, C. (2017, March 28). emeraldinsight. Retrieved from <https://www.emeraldinsight.com/doi/abs/10.1108/AJEM>
- [25].Nyoni T (2018h). Population Growth in Zimbabwe: A Threat to Economic Development? DRJ – Journal of Economics and Finance, 2 (6): 29 – 39.
- [26].<https://www.researchgate.net/publication/318211505>
- [27].<https://corporatefinanceinstitute.com/resources/knowledge/economics/>
- [28].<https://www.macrotrends.net/countries/ZWE/zimbabwe/gnp-gross-national-product>
- [29].<https://www.macrotrends.net/countries/ZWE/zimbabwe/gdp-per-capita>
- [30].<https://www.thebalance.com/gross-national-income-4020738>
- [31].<https://corporatefinanceinstitute.com/resources/knowledge/economics/gross-domestic-product-gdp/>
- [32].<https://corporatefinanceinstitute.com/resources/knowledge/economics/gross-national-product-gnp/>
- [33].<https://corporatefinanceinstitute.com/resources/knowledge/economics/imports-and-exports/>
- [34].<https://www.researchgate.net/publication/318211505>