

Machine learning approaches for forecasting inflation: empirical evidence from Sri Lanka

W.M.S Bandara* and W.A.R. De Mel

Department of Mathematics, Faculty of Science, University of Ruhuna, Matara, Sri Lanka

*Correspondence: bandarasudarshana009@gmail.com  ORCID: <https://orcid.org/0000-0003-0968-8940>

Received: June 18, 2023, Accepted: March 05, 2024, Published: June 30, 2024

Abstract The aim of this study is to forecast the inflation rate using supervised machine learning models (SMLM). While SMLMs are widely used in various fields, they have also been widely applied to forecast inflation rates. Therefore, the main objective of this study is to identify the best model for forecasting inflation among four different SMLMs: LASSO regression (LR), Bayesian Ridge Regression (BRR), Support Vector Machine Regression (SVR), and Random Forest Regression (RFR) models. To achieve this objective, two different types of cross-validation techniques were used: K-fold cross-validation method (KCV) and walk forward validation (WV) methods. These techniques were used to estimate the parameters and hyper-parameters for each machine learning model with the aid of root mean square error. The mean absolute percentage error (MAPE) was used to compare the performance of the different SMLMs. Empirical evidence from Sri Lanka between 1988 and 2021 was used to test the performance of the SMLMs in forecasting inflation rates. The results show that the LR model with walk forward validation is the best method for forecasting the future inflation rate of Sri Lanka based on the MAPE value. Overall, this study demonstrates the effectiveness of SMLMs in forecasting inflation rates under the critical conditions and highlights the importance of employing appropriate cross-validation techniques when using these SMLM models. The findings of this study can provide valuable insights for policymakers, investors, and researchers who are interested in forecasting inflation rates.

Keywords: Cross-Validation, Hyper-parameter, Inflation Forecasting, Machine Learning Models.

1 Introduction

Inflation is a critical macro-economic indicator that measures the rise in the general price level of goods and services over a given period. It affects individuals, investors, businesses, and the overall economy of a country. High rate of inflation causes a

decrease in the purchasing power of the currency, which can lead to economic instability, social unrest, and political turmoil. Therefore, it is essential to forecast the inflation rate accurately to take preemptive measures to mitigate its adverse effects.

Several factors make forecasting inflation rates a challenging task. For instance, unpredictable events such as natural disasters, political and social conflicts, and global economic crises can impact the economy in unexpected ways. Additionally, traditional forecasting methods may not be adequate to capture the complex and dynamic relationships between economic variables. Therefore, there is a need to develop advanced forecasting models that can handle the complexity of the economic data and provide accurate predictions.

The present research aims to assess the specifications of machine learning models for predicting inflation rates in Sri Lanka, as the approaches used in conventional approaches are quite restricted. Earlier techniques like the univariate models and the Phillips curve are therefore not as effective in capturing on real-time economic data because they do not fully address this dynamic nature of data, especially in the presence of disruptions or shifts and do not take into account local peculiarities. By the same token, in machine learning models, it is possible to enhance the forecast accuracy of inflation modelling, by managing non-linear relationships and extensive datasets. The topic of a particular interest in this regard as it shows that machine learning could offer a lot to Sri Lankan financial and banking space but is not actively embraced by the policymakers or the Central Bank as of now. The potential contribution to knowledge of this research lies in the fact that its results could be useful for policymakers, central banks and investors, primarily, as the results of improved inflation forecasting will contribute to the formation of accurate decisions.

In earlier research, several classical models have been used to predict inflation rates. For example, in the survey, Jesmy (2010) used the Box–Jenkins method to model Sri Lanka's monthly mean inflations using the data from the year 1952 to 2009, and found that the most suitable among all the models chosen by Box and Jenkins based on the adjusted R-squared statistic is the univariate ARIMA (1,1,2). Similarly, Bandara (2011) used Inflation data from 1985 to 2005 to predict inflation using the VAR models. Jere and Mubita (2016) used Holt's Exponential Smoothing to forecast inflation rate in Zambia. Thus, it is still difficult to compare these two models on the account of variation in the political, social and environmental conditions of the developing countries across the globe, differing samples of time under consideration and the factors characteristic of individual countries. This challenge only goes further to show why more flexible and reliable models are needed in different conditions, and this is where we can get better machine learning models.

Standard Phillips Curve Models (PCM), which rely on economic activity, have acted as a basis to the typical forecasting models of inflation. Stock and Watson (1999) also argue that these PCM-based models outperform the traditional inflation forecasting models. Atkeson and Ohanian (2001), however, criticized this claim by

showing that Phillips curve forecast of U.S. inflation over a 15-year period is not better than those obtained from a random walk model. Nevertheless, this instability of forecasting relationships is not limited to traditional Phillips curve based models but extends to other theoretical or ad hoc empirical models used in the literature as well (Marcellino *et al.* 2000, Goodhart and Hofmann (2000), Marcellino (2002). Although forecasting specifications built by adding one indicator of real activity at the time work poorly and tend to be unstable, some improvements have been documented by Wright (2003), Cristadoro *et al.* (2005), Granger and Jeon (2004), and Inoue and Kilian (2006) using methods that combine information obtained from many predictors.

In the literature, various categories of cross-validation methods have also been used in conjunction with conventional methods of forecasting inflation. Cross-validation is another approach to validating the performances of a given model by dividing a data set into training and validation segments. This approach goes a long way in reducing the model overfitting as well as increasing its ability to generalize new data sets. For instance, Bergmeir and Benítez (2012) used various cross-validation methods (i.e., stranded 5-fold, blocked, last block, second block and second) with time series models, among which, the blocked cross-validation predictively outperformed others in terms of Root Mean Square Error (RMSE) statistic making it a more useful method.

Machine learning models are commonly used in forecasting inflation data. Ulke *et al.* (2018) forecast the core and non-core versions of inflation in the USA by using univariate Auto Regressive Distributed Lag (ARDL), Multivariate time series (VAR), SVR, k-nearest neighbour, and artificial neural network models. According to their results, ARDL provided the highest prediction accuracy for forecasting core-CPI inflation, while SVR outperformed the other models in forecasting core inflation. All these machine learning models work better with more volatile and irregular series.

In this study, based on historical data, the performance of four supervised learning models, i.e., Lasso Regression (LR), Bayesian Ridge Regression (BRR), Support Vector Regression (SVR), and Random Forest (RF) was tested for predicting inflation rates. In evaluating the proposed models, cross-validation, specifically the Walk Forward Validation (WV) and the K-fold Cross-Validation (KCV) were employed in a bid to fine-tune the hyperparameters of the model for efficiency. We chose the Mean Absolute Percentage Error (MAPE) as the primary evaluation function.

Overall, our study aimed to identify the best machine learning model for forecasting the monthly mean inflation rate in Sri Lanka, considering different cross-validation methods and performance metrics. Our results provide useful insights into the effectiveness of different machine learning models for time-series data and can guide practitioners in selecting the most suitable model for their specific application.

2 Material and Methods

2.1 Models

In this subsection, we briefly explain supervised machine learning models, which are used to forecast inflation data, and two cross-validation techniques.

LASSO Regression (LR)

According to Ogutu *et al.* (2012), LASSO (Least Absolute Shrinkage and Selection Operator) is an L1 regularization technique used in regression analysis. This method applies a penalty to the regression coefficients, effectively shrinking some of them to zero, which allows LASSO to automatically perform feature selection. This means that LASSO can handle models with a large number of variables by identifying and retaining only the most significant predictors.

The LASSO objective function can be defined as follows:

$$\beta'(\text{lasso}) = \operatorname{argmin}_{\beta} (||Y - \beta X||_2^2 + \lambda ||\beta||_1) \quad (1)$$

In this equation, $\beta'(\text{lasso})$ represents the optimal coefficients that minimize the LASSO objective function. The vector β contains the coefficients for the regression model. Y denotes the observed values of the dependent variable, which is the outcome or response that the model aims to predict, while X is the matrix of predictor variables (independent variables) used to predict Y . The parameter λ is the regularization parameter that controls the strength of the L1 penalty applied to the coefficients. This penalty encourages sparsity, meaning it reduces some coefficients to zero, effectively selecting a simpler model with fewer predictors. By using LASSO, we can manage models with many potential predictors, retaining only those that have the most significant impact on the outcome, which is particularly useful in scenarios where model simplicity and interpretability are important.

Bayesian Ridge Regression (BRR)

In Bayesian ridge regression (Hoerl and Kennard 1970), the estimate β' is obtained by using the L_2 norm, and it is given by;

$$\beta'(\text{BRR}) = \operatorname{argmin}_{\beta} ||Y - \beta X||_2^2 + \lambda ||\beta||_2^2 \quad (2)$$

where $||\beta||_2 = \sum_1^n \beta_i^2$ is the L_2 - norm penalty on β and $\lambda \geq 0$. In this case, we obtain the posterior distribution to estimate β with normal likelihood and normal prior distribution.

Support Vector Regression (SVR)

SVR (Smola and Schölkopf 2003) gives the flexibility to define how much error is acceptable in our model. It will compute the parameter estimates by utilizing the following minimization problem. Minimize

$$\min \frac{1}{2} ||\beta||^2 \quad (3)$$

under constraint

$$y_i - W^T x_i \leq \epsilon + \xi_i, W^T x_i - y_i \leq \epsilon + \zeta_i^*, \text{ and } \xi_i, \zeta_i^* \leq 0 \quad (4)$$

The objective of SVR is to minimize half of the squared norm of the weight vector, $\frac{1}{2} ||W||$, which contributes to achieving model simplicity and better generalisation. The constraints ensure that the absolute difference between the predicted value, $W^T x_i$, and the actual value, y_i , for each data point i , is within a predefined margin ϵ . This margin ϵ represents the maximum acceptable error, serving as a critical hyperparameter in SVR. By optimizing ϵ , we can effectively control the trade-off between the model's complexity and its accuracy, thus enhancing the model's ability to generalize well to new data while adhering to specific error tolerance requirements.

Random Forests Regression (RFR)

RFR is a tree-based algorithm with each tree depending on a set of random variables (Cutler *et al.* 2012). Let $X = (X_1, X_2, \dots, X_p)'$ be a p -dimensional random input vector and Y be the response variable. Moreover, we assume that $p_{XY}(X, Y)$ is the unknown joint distribution of X and Y . The objective of the RFR is to find a function $f(X)$ to predict the response variable Y by minimizing the risk function.

$$E_{XY}(L(Y, f(X))) \quad (5)$$

where $L(Y, f(X)) = (Y - f(X))^2$ is the squared error loss function. Here, one can define $f(X)$ as $f(x) = E(Y | X = x)$, and in regression setting, $f(x)$ can be written as

$$f(x) = \frac{1}{J} \sum_{j=1}^J h_j(x) \quad (6)$$

with respect to a collection of basis functions $h_1(x), h_2(x), \dots, h_J(x)$.

2.2 Cross-validation Methods

Cross-validation is a method used to increase the effectiveness of a machine learning model. It is based on re-sampling training data to train and test groups and evaluating

model performance under over-fitting and under-fitting conditions. In this study, we use two types of cross-validation methods, namely, K -fold cross-validation and walk forward validation.

K-folds Cross-validation (KCV)

In regression and classification settings, the K -fold cross-validation (Trevor *et al.* 2009) technique is commonly used due to its simplicity, fairness, and high effectiveness. The dataset is divided into K intervals, and one subinterval is used as test data while the remaining $K-1$ intervals are used as training data. By fitting the model K times and selecting the optimal K value that minimizes the RMSE, we can evaluate the model's performance.

Walk Forward Validation (WFOV)

Walk forward validation (Wickham 2016) also known as time-series validation, is a technique commonly used for evaluating time-series data. In this method, the entire dataset is divided into K intervals. The model is trained on the first interval and tested on the second. Then, the first two intervals are combined to train the model, which is tested on the third. This process is repeated until the first $K - 1$ intervals are used for training, and the remaining interval is used for testing. At each step, the model is fit, and the RMSE is computed. The optimal K value is selected based on the minimum RMSE.

Hyper-parameter tuning

In the case of machine learning models, hyperparameters are precise settings which regulate different aspects of the said model. For example, in Lasso or Ridge kind of regression models the most important hyperparameter is the regularization strength (λ) which indicates how much penalty should be imposed on the coefficients of the model. For instance, learning rate, the process of adjusting the weights of a given model in the neural network during the training phase is strongly influenced by a hyperparameters it then has the number of hidden layers which determines the depth of the network. Hyperparameters were determined to significantly affect the model's performance on a given data set and therefore choosing the right hyperparameters is critical for best success. The several approaches like grid search or random search can be employed in order to employ a systematic approach of finding out the best hyperparameters out from the entire search space.

2.3 Error calculation methods

Let n , Y_t , and $\hat{Y}_t$ be the number of fitted points, the actual value of the response variable Y at time t , and the predicted value of Y_t , respectively.

Mean Absolute Percentage Error (MAPE)

The Mean Absolute Percentage Error (Armstrong and Kollopy 1992) can be calculated by using the following formula.

$$MAPE = \frac{100\%}{n} \sum_{i=1}^n \frac{|Y_t - \hat{Y}_t|}{Y_t}. \quad (7)$$

Root Mean Square Error (RMSE)

The Root Mean Square Error (Hyndman and Koehler 2006) is given by;

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (Y_t - \hat{Y}_t)^2} \quad (8)$$

where n is the number of observations, Y_t is the actual value of observation and $\hat{Y}_t$ is the predicted value of observation.

2.4 Dataset

In this study, we consider the monthly inflation rate data in Sri Lanka from January 1988 to December 2023 (Tradingview 2023). Figure 1 depicts this data.

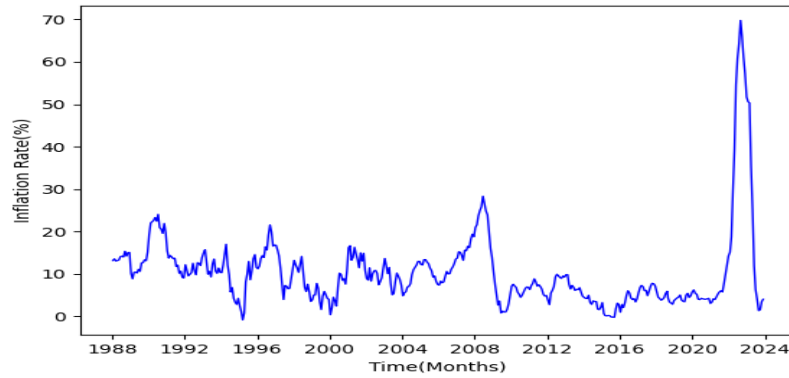


Fig 1. Monthly mean inflation rate of Sri Lanka (1988-2023)

The time series plot in Figure 1 shows a stochastic behaviour of the inflation data, and one can notice that there are a few unusual data points, especially one at June in 2008. The reasoning for this may be that during this time period, the war in Sri Lanka was at a critical stage.

In order to fit the above four machining learning models, we have to convert the inflation rate data into a machine learning data set by introducing a new set of

variables. The response variable, $Y(t)$ is the inflation rate at time t . Here, we use four predictor variables: the first, second, third, and fourth differences of $Y(t)$ and denote them as $Y(t-1)$, $Y(t-2)$, $Y(t-3)$ and $Y(t-4)$, respectively.

3 Results

Using Python 3.8.5, the study embarks on a comprehensive examination of four supervised machine learning models namely: LR, BRR, SVR, and RFR. Each model, distinguished by its unique way of modeling, was meticulously crafted and fine-tuned to forecast the inflation rate, thus building upon our existing corpus of knowledge.

In this part of the paper, we proceed with an empirical analysis that follows an exploration conducted around a set of data collected monthly starting from January 1988 to a pool of 432 data points. This dataset was indeed useful in the studies we conducted in predicting the monthly mean inflation rate in Sri Lanka under a supervised machine learning approach. Once the dataset was obtained, there was pre-analysis in which data was imported and some transformations that saw it change in form to fit a time-series data as well as data visualization in order to have a feel on the inflation trend over time. This phase laid the groundwork for further analysis that could be of higher analytical feel.

One of the key steps of the proposed methodology was the dataset pre-processing to be used in the machine learning algorithm. To capture the time-series aspects of the specified variables, we added first lags ($t-1$ to $t-4$) in our dataset, thereby improving the models' sensitivity to changes in the inflation rate. This preprocessing step was useful in transforming the time series data into a form compliant with supervised learning.

After data preprocessing, the model training and model assessment commenced where more than one model was used to develop the best model; these models include the LR, BRR, SVR, and at last, the RFR. To refer to the validation of our model, a structured approach to model validation was applied via KCV as well as WFV. In particular, the K value of K-Fold Cross-Validation and the number of splits of Time Series Cross-Validation were taken as factors to be tested for each model in terms of the RMSE to assess how their predictability changed relative to the selected configurations.

The exploration included K-Fold splits from 2 to 40 folds as well as the splits for Walk Forward Validation of the same breadth, which provides a reliable view on the stability of the models at issue. We used this kind of extensive study as an effort to identify the best fold that along with the best split number for achieving highest level of forecast accuracy for the monthly mean inflation rate in Sri Lanka which eventually helps to identify the models suitable for this particular forecasting job.

'GridSearchCV' python algorithm in scikit learn library was used in fine-tuning the parameters according to the values determined in a grid form. This process was

important in improving the performance of each model to obtain the configuration that has the least RMSE so that the accuracy forecast made by the models would be refined.

Through this refined simulation study, we aspired to not only assess the predictive performance of each model under various configurations but also to identify the most effective model and configuration for forecasting the monthly mean inflation rate in Sri Lanka. The outcomes of this study are poised to contribute valuable insights into the application of machine learning techniques in economic forecasting, with the potential to guide future research and policy formulation in this vital area.

Figure 2 presents a comparative analysis of the model performance between LR and BRR using two different validation techniques; KCV and WFV. The four subplots graphically illustrate the RMSE as a function of the number of folds used in the validation process. The top two plots show the variability of the LR model's RMSE across an increasing number of folds for both KCV and TSCV, indicating a general trend of RMSE stabilization as the number of folds increases.

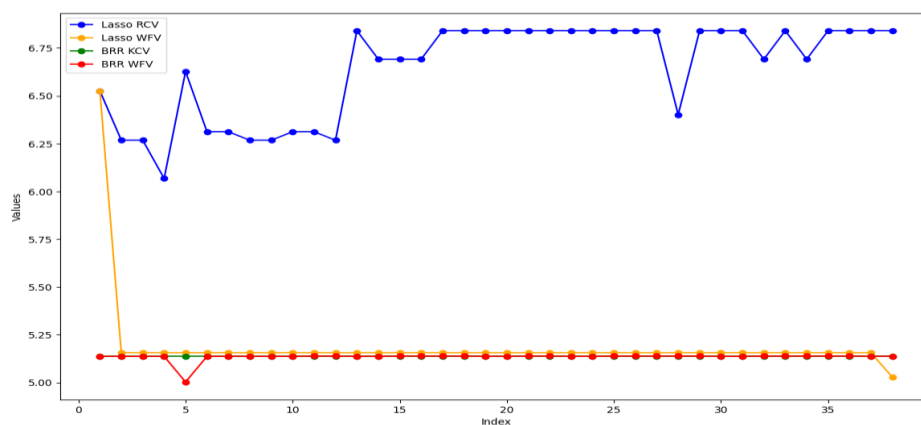


Fig 2. Comparison of LR and BRR models

The bottom two plots exhibit a similar analysis for the BRR model. Notably, the BRR model demonstrates a less volatile, more stable RMSE across different fold counts, suggesting potential robustness in performance with this model. Each subplot's x-axis represents the number of folds ranging from 2 to 40, while the y-axis quantifies the RMSE providing a clear visual depiction of how each model's prediction accuracy is influenced by the fold count in cross-validation.

Figure 3 depicts the performance evaluation of RF and SVR models, utilizing two distinct validation techniques: KCV and TSCV. The subplots, two for each model, demonstrate the RMSE at varying numbers of folds within each validation framework. The top-left and top-right subplots represent the RF model's performance under KCV and TSCV, respectively, and both exhibit significant variability in

RMSE, indicating sensitivity to the number of folds. The bottom subplots showcase the SVR model's performance, with the bottom-left for KCV and bottom-right for TSCV. The SVR model displays a more pronounced variation in RMSE in the KCV approach, while its TSCV performance shows less fluctuation. In all subplots, the x-axis scales from 2 to 40 folds, reflecting the sample size's impact on the model's predictive accuracy, while the y-axis measures the RMSE, providing an assessment of the model's performance over different validation splits.

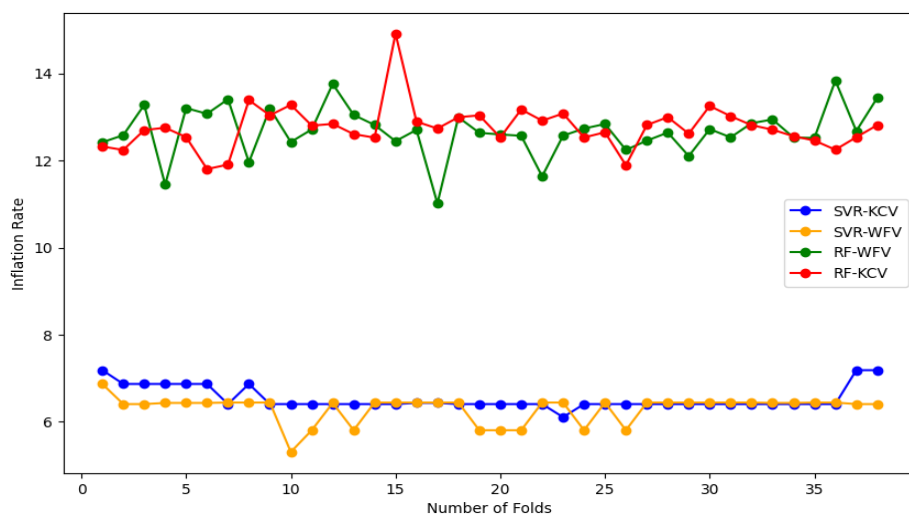


Fig 3. Comparison of Random Forest (RF) and Support Vector Regression (SVR) models' performance

Table 1: Comparison of K-Fold Cross-Validation (KCV) and Walk Forward validation (WFV) techniques in different models

Model	KCV	WFV
LR	5	40
BRR	2	6
RF	7	18
SVR	24	11

Table 1 displays the variation in the RMSE across a range of fold numbers for RF and SVR models, employing both KCV and TSCV techniques. For the RF model, the RMSE exhibits notable fluctuations with different fold numbers in both KCV and TSCV, suggesting a dependency on the chosen fold count, with no clear trend toward

stabilization. In contrast, the SVR model shows a stark drop in RMSE at certain fold numbers within the KCV method, implying specific configurations that markedly improve its performance. The TSCV results for SVR indicate less variability, hinting at potentially better consistency across time-structured data splits. The table accompanying the graphical data summarizes the optimal fold number for each model and validation method, highlighting that the most suitable fold number varies significantly across models and validation techniques. For LR, the optimal fold numbers are 5 for KCV and 40 for TSCV; for BRR, they are 2 for KCV and 6 for TSCV; RF performs best at 7 folds for KCV and 18 for TSCV; and SVR finds its optimal performance at 24 folds for KCV and 11 for TSCV. This information is critical for pinpointing the most effective model configurations for forecasting purposes within the given timeline.

Algorithm 1: Pseudo code for LR

Input: Response ($\mathbf{Y}$) dependent on predictor($\mathbf{X}$) error tolerance ϵ

: LASSO solution for (1)

- Begin
 - Data Preposing: Normalization X and Y
 - Initialization $U = 0, \hat{y} = y - U$
 - $c = X^T Y$
 - $C = \text{Max}_j \{|c_j|\}$
 - $\hat{j} = j \{|C_j|\}, A = \hat{j}$
 - Calculate weight by $R_w = \text{PART_PAC}(X)$ or $R_w = \text{IW}(X)$ or $R_w = \text{CRITIC}(X)$
 - Centralization R_w
 - When $\|\hat{y}\|_L \leq 1$ and $|A| < m$, do β cycle
 - End cycle
 - Return (β)
-

Algorithm 1 is a pseudo code for the LR algorithm, which is used to predict the response (Y) dependent on predictor (X) with an error tolerance ϵ . The algorithm starts with data preprocessing and initialization, followed by weight calculation based on the chosen method. The LR algorithm iteratively cycles through β until the desired result is achieved. Finally, the algorithm returns the calculated β values.

Figure 4 illustrates the forecast outcomes from LR and BRR using two validation strategies, namely Cross-Validation with KCV and WFV. In each plot, the 'Training data' line represents the historical data used to train the models. The 'LR_Fitted' or 'BRR_Fitted' lines depict the predicted inflation rates from the respective models. Additionally, 'Test Data' indicates the actual observed values that the models aim to predict. For both LR and BRR models, the left graphs demonstrate the KCV results, while the right graphs display the WFV results. These plots allow us to assess how well each model has learned from the past data (blue line) and how accurately it can

project these insights into the future (black line), especially when contrasted with the real-world data (red line). Notably, we observe that the test data exhibits sharp increases, posing significant challenges for both LR and BRR models, which is reflected in the discrepancies between the predicted and actual values.

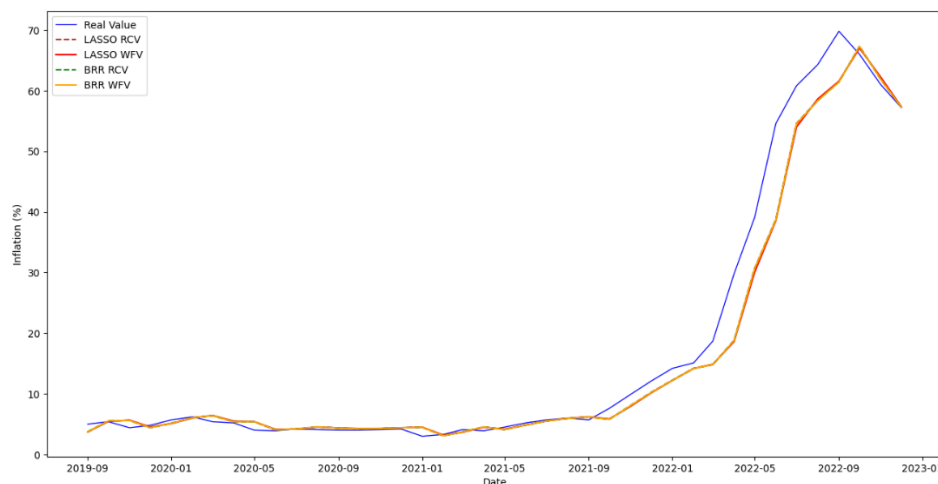


Fig 4. Fitted test inflation rates data for LR and BRR models with KCV and WFV techniques

These discrepancies are particularly noticeable towards the end of the timeline, where an unexpected surge in inflation rates is evident. This surge is not well captured by the models, underscoring the inherent challenges in forecasting outliers and abrupt shifts in economic trends. By comparing the KCV and WFV approaches, we gain insights into each model's ability to adapt to time-evolving patterns, with the WFV validation appearing to provide a more responsive approach to the most recent data trends.

The pseudo-code outlines the implementation of the BRR algorithm, a commonly used method for linear regression analysis. The algorithm begins by taking the response variable Y , dependent on the predictor variable X , along with an error tolerance ϵ as inputs. The data preprocessing step involves normalizing both X and Y , followed by the initialization of variables such as U and $\hat{y}$. The algorithm then calculates the weights using one of the methods—*PARTPAC*, *IW*, or *CRITIC*. These weights, denoted as R_w , are centralized before proceeding to the Gibbs sampling step.

In Gibbs sampling, the algorithm iteratively samples from the posterior distribution of the parameters β and λ . Multiple chains are used to ensure robust sampling, and convergence is assessed using the Gelman and Rubin statistic. The

algorithm continues sampling until the convergence criterion is met, indicated by the Potential Scale Reduction Factor (*PSRF*) being close to 1 for all parameters.

Algorithm 2: Pseudo code for BRR

Input: Response Y dependent on predictor X , error tolerance ϵ

Output: Bayesian Ridge Regression solution for (2)

- Data Preposing: Normalization X and Y
 - Begin Initialization $U = 0, \hat{y} = y - U$
 - $c = (Y - X\beta)^T (Y - X\beta) + \lambda\beta^T \beta$
 - $C = \text{Max}_j \{ |c_j| \}$
 - $\hat{j} = \arg \max_j \{ |C_j| \}, A = \hat{j}$
 - Calculate weight by $R_w = \text{PART_PAC}(X)$ or $R_w = IW(X)$ or $R_w = \text{CRITIC}(X)$
 - Centralize R_w
 - Perform Gibbs sampling to iteratively sample from the posterior distribution of the parameters β and λ .
 - Use multiple chains to run the sampling process.
 - Assess the convergence of the Gibbs sampler using the Gelman and Rubin statistic (Potential Scale Reduction Factor, PSRF).
 - Continue the Gibbs sampling until PSRF values for all parameters are close to 1, indicating convergence.
 - When $\| \hat{y} \|_{L_2} < 1$ and $|A| < m$, perform β cycle.
 - End Cycle:
 - Return (β)
-

After confirming convergence, the algorithm performs the β cycle, which continues until the specified conditions such as the L_2 norm and the size of the active set $|A|$ are satisfied. The algorithm then outputs the estimated β values that minimize the objective function, providing the final BRR solution.

Algorithm 3: Pseudo code for SVR

Input: S, λ, T, k

Initialize Choose: w_i s.t $\|w\| \leq \frac{1}{\lambda}$

For $t = 1, 2 \dots T$

Choose ($A_i \notin S$) Where $\|A\| = k$

Set $A_i^\dagger = \{(x, y) \in A_t : (W, x) < 1\}$

Set $\eta_t = \frac{1}{\lambda_t}$

Set $N_{t+1} = \min \left\{ 1, \frac{1/\sqrt{\lambda}}{\|w_{t+\frac{1}{2}}\|} \right\} w_{t+\frac{1}{2}}$

Output: $w_{T+\frac{1}{2}}$

Algorithm 3 is the pseudo code for SVR, a popular regression algorithm used in machine learning. SVR involves choosing a set of support vectors from the input data and constructing a linear model to minimize the error between the predicted values and the actual values. The algorithm involves several iterations where support vectors are chosen, and the model is updated with the chosen support vectors until convergence is reached.

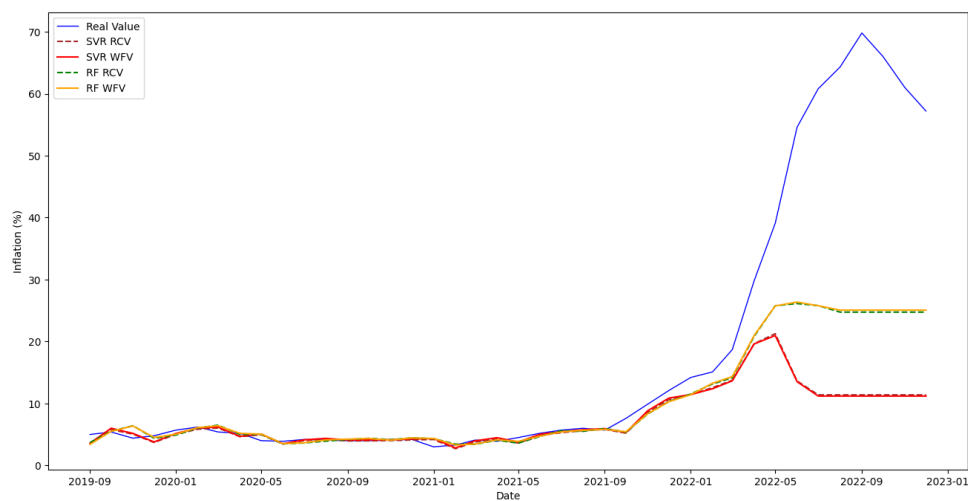


Fig 5. Fitted test inflation rates data for SVR and RFR models with KCV and WFV technique

Figure 5 displays the forecasting performance of SVR and RFR models utilizing KCV and WFV methods. In each graph, the historical data, referred to as 'Training data,' is plotted alongside the predicted values, labeled as 'SVR_Fitted' or 'RF_Fitted,' depending on the model, and the actual test data, marked as 'Test Data.'

The top two plots pertain to the SVR model, with the left showcasing results obtained through KCV and the right through WFV. Similarly, the bottom two plots relate to the RFR model, also separated by the two validation methods. Across all graphs, the models' fitted lines show how the respective algorithms have attempted to capture the underlying pattern in the training data and extend those patterns into the future. Notably, spikes in the test data, signifying sudden increases in the inflation rate, present challenges to the models, with the predicted values deviating significantly from these actual values.

These deviations are particularly pronounced at the tail ends of the plots, where both models struggle to anticipate the dramatic rise in the inflation rate that the actual test data reflects. The stark contrast between predicted and actual values during these periods highlights the difficulty machine learning models face when predicting extreme values that deviate from historical trends. The comparative analysis across different validation techniques demonstrates how each model responds to the

progression of time and changes in data patterns, with the WFV method revealing the models' adaptability to more recent data.

Algorithm 4: Pseudo code for RF

Precondition: A training set $S = (x_1, y_1), \dots, (x_n, y_n)$, features F , and number of trees in forest B

- Start
 - Function Random Forest (F, S)
 - $H \leftarrow \emptyset$
 - for $i \in 1 \cdot B$ do
 - $S^i \leftarrow$ A bootstrap sample
 - $h_i \leftarrow$ Randomized Tree Learn (S^i, F)
 - $H \leftarrow H \cup \{(h_i)\}$
 - End for, return H
 - function: Randomized Tree Learn(S, F)
 - At each node: $f \leftarrow$ very small subset of F
 - Split on best feature in f
 - Return - learned tree
-

Algorithm 4 presents the pseudo code for RFR. RF is a popular ensemble learning method used for both classification and regression problems. The algorithm builds a specified number of decision trees using a bootstrapped sample of the training data and selects a random subset of features at each node to split on. The final prediction is the average of the predictions from all the trees in the forest. The RF algorithm is known for handling high-dimensional data and avoiding overfitting.

The MAPE values presented in Table 2 indicate small differences between the KCV and WFV methods, as well as between the LR and BRR models. Although these differences suggest similar performance among the methods, they may not fully capture the nuances of each model's effectiveness in different contexts.

Table 2: MAPE values in test data of each Table SMLM

<i>KCV</i>		<i>WFV</i>	
Model	MAPE	Model	MAPE
LR	12.54	LR	12.52
BRR	12.56	BRR	13.54
SVR	26.20	SVR	26.05
RF	22.66	RF	21.97

Therefore, while the LR model paired with WFV showed the best overall performance with a MAPE of 12.52%, the slight variations in MAPE between LR-

KCV (12.54%), BRR-KCV (12.56%), and BRR-WFV (13.54%) raise concerns about the robustness of these results.

Given these small differences, it is essential to conduct further analysis by forecasting values using LR-KCV, LR-WFV, and BRR-KCV. This approach allows for a more comprehensive evaluation of model accuracy and ensures that the selection of the best model is based on a more reliable assessment.

The close MAPE values between these methods could be attributed to several factors. One possibility is that the models are capturing the underlying data patterns similarly due to the nature of the dataset or the preprocessing steps applied. Additionally, the choice of cross-validation methods may not have introduced significant variability in the model performance, which often occurs in well-behaved datasets or when the time-series data exhibits relatively stable trends.

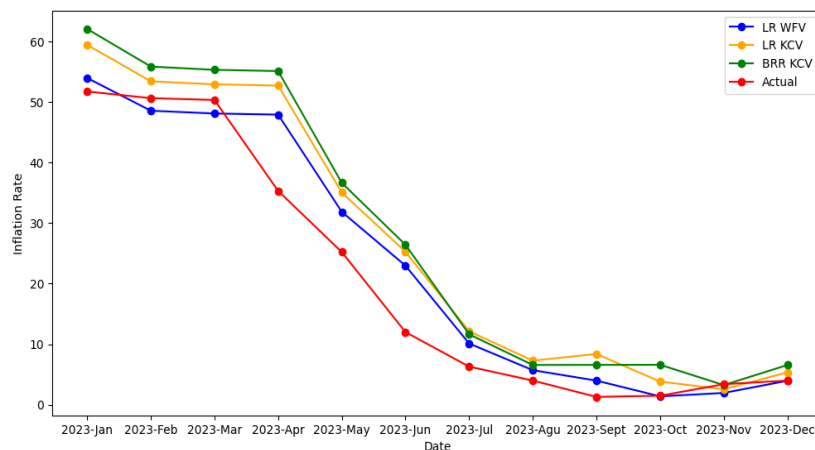


Fig 6. Forecasted values validation data via real data- LR with WFV and KCV and BRR with KCV

Figure 6 offers a visual representation of the actual values and the predicted values of inflation over an year. The variable in question is represented as Y_t while the forecasted values have been generated using a LR with WFV, LR with KCV and BRR with KCV. The graph is composed of four distinct lines: the red line traces the actual observed values, while the blue, orange, and green lines depict the forecasts generated by the LR-WFV, LR-KCV, and BRR-KCV models, respectively.

The actual values are marked as individual points, offering a clear benchmark against which the model predictions can be evaluated. The forecasted values are shown as continuous lines, facilitating a direct comparison with the real data. This visualization highlights how closely the predictions from each model align with the actual values over time.

Upon examining the graph, it is evident that while all three models generally perform well in their predictions, certain discrepancies arise. In particular, the LR-

KCV and BRR-KCV models occasionally overestimate or underestimate the actual values more noticeably than the LR-WFV model. These discrepancies underscore the inherent difficulties in economic forecasting, even when utilizing advanced machine learning techniques. The slight variations among the models also reflect the impact of different cross-validation methods on the accuracy of predictions.

Table 3 provides a comprehensive overview of the forecasting performance of the LR and BRR models, using both WFV and KCV techniques over the course of a year from January to December 2023. The table compares the actual values of an economic indicator, Y_t with the forecasted values generated by each model.

Table 3: Foretasted values for Lasso Regression (LR) and Bayesian Ridge Regression (BRR) model with Walk-Forward Validation (WFV) and K-Fold Cross-Validation (KCV) technique

Date	Actual	LR WFV	LR KCV	BRR KCV
2023-Jan	51.7	53.93551	59.39069	62.02584
2023-Feb	50.6	48.53096	53.38405	55.81061
2023-Mar	50.3	48.06587	52.87261	55.27575
2023-Apr	35.3	47.88063	52.67569	55.07012
2023-May	25.2	31.83928	35.02321	36.61591
2023-Jun	12.0	22.95797	25.25337	26.40100
2023-Jul	6.3	10.11301	12.12581	11.63087
2023-Agu	4.0	5.71609	7.28760	6.57341
2023-Sept	1.3	3.99010	8.38901	6.58851
2023-Oct	1.5	1.39323	3.83255	6.60232
2023-Nov	3.4	1.95886	2.15445	3.25237
2023-Dec	4.0	3.97016	5.36714	6.565702
MAPE		43.90	96.55	102.15

A close examination of the table reveals that the LR model using WFV delivers a relatively strong predictive performance, as indicated by the close alignment between the forecasted and actual values throughout most of the year. This alignment suggests that the LR model, when validated with WFV, effectively captures the underlying patterns in the data. However, there are notable discrepancies in certain months, particularly in April and May, where the LR model's forecasts significantly overestimate the actual trends. These deviations point to areas where the model's accuracy could be improved.

In contrast, the LR and BRR models validated with KCV exhibit more pronounced deviations from the actual values, particularly in the latter half of the year. The larger discrepancies in the KCV-based models highlight the potential limitations of this validation technique in accurately predicting the economic indicator. This is further reflected in the MAPE values, where the LR-WFV model achieves a MAPE of

43.90%, while the LR-KCV and BRR-KCV models record significantly higher MAPE values of 96.55% and 102.15%, respectively.

These MAPE values indicate that while the LR-WFV model provides a reasonable level of accuracy, the KCV-based models struggle to maintain the same level of precision, emphasizing the importance of selecting the appropriate validation technique. The analysis underscores the necessity for continual refinement of these models to enhance their forecasting accuracy, especially in the context of economic indicators where precision is crucial.

In summary, Table 3 not only illustrates the comparative performance of the LR and BRR models under different validation techniques but also highlights the critical role of model selection and validation in achieving reliable forecasts. The findings suggest that the LR model with WFV is more robust in this scenario, though ongoing optimization remains essential to minimize forecasting errors.

4 Discussion

The analysis presented in this study reveals significant insights into the challenges of forecasting inflation during periods of extreme economic volatility. The performance of Lasso Regression (LR) and Bayesian Ridge Regression (BRR) models was evaluated using both Walk-Forward Validation (WFV) and K-Fold Cross-Validation (KCV) techniques, providing a detailed comparison of actual inflation rates against model predictions.

The observed Mean Absolute Percentage Error (MAPE) values highlight the difficulties inherent in forecasting during such turbulent times. The LR-WFV model achieved a MAPE of 43.90%, demonstrating relatively better performance compared to the LR-KCV and BRR-KCV models, which recorded significantly higher MAPE values of 96.55% and 102.15%, respectively. These elevated MAPE values are not surprising, given the volatile economic conditions characterized by sharp rises and falls in inflation rates. This period's unpredictability underscores the limitations of traditional forecasting models when applied to highly unstable economic environments.

A key factor contributing to the LR model's superior performance with WFV is its ability to adapt to changing data patterns. The WFV technique enables the model to continuously update and refine its predictions as new data becomes available, making it particularly effective in capturing the rapid changes typical of a crisis period. This adaptability contrasts with the static nature of KFCV, which does not accommodate real-time data changes and, therefore, struggles to keep up with the dynamic nature of the economic crisis.

Furthermore, the LR model's inherent feature selection capability plays a crucial role in its performance. By penalizing less relevant variables, the model focuses on the most significant predictors, which is particularly beneficial in a scenario where

the underlying data is subject to sudden and severe fluctuations. This targeted approach allows the LR model to maintain a higher degree of accuracy compared to the other models evaluated.

In contrast, models such as SVR and RF, while robust in more stable conditions, demonstrate limitations in their ability to handle non-linear and abrupt changes. Similarly, the BRR model, though effective under normal circumstances, lacks the flexibility required to manage the extreme variability observed during the economic crisis.

The findings from this study indicate that while traditional models may provide some level of insight, their effectiveness is significantly compromised during periods of economic instability. The LR model with WFV stands out as the most reliable approach in this context, though it is not without its limitations. The high MAPE values observed across all models suggest that forecasting during an economic crisis remains a formidable challenge, and there is considerable room for improvement in developing models that can more accurately predict inflation under such conditions.

In summary, the discussion highlights the critical need for adaptive and flexible forecasting models in the face of economic volatility. While the LR model with WFV demonstrates promising results, future research should focus on enhancing this approach by integrating additional macroeconomic indicators and exploring advanced techniques for handling outliers and abrupt data shifts. This will be essential for improving the accuracy and reliability of inflation forecasts in highly volatile economic environments.

5 Conclusion

This study has provided a comprehensive evaluation of various machine learning models in forecasting inflation rates during a period of significant economic turbulence. The analysis focused on the performance of LR and BRR models, employing both WFV and KCV techniques. The results revealed that the LR model, particularly when used with WFV, outperformed other models, including SVR, RF, and BRR, in terms of forecasting accuracy.

The MAPE values, although high across all models, were lowest for the LR-WFV model, indicating its relative robustness in handling the volatile economic conditions characterized by rapid and unpredictable changes in inflation rates. The study underscores the importance of adaptability in forecasting models, with the WFV technique enabling the LR model to adjust to new data in real-time, thereby improving its predictive performance during periods of crisis.

The findings also highlighted the value of feature selection in the Lasso Regression model, which allowed it to focus on the most relevant predictors, thereby enhancing its accuracy in a complex and unstable environment. In contrast, models like BRR,

which do not incorporate such mechanisms, struggled to maintain accuracy under the same conditions.

In conclusion, while the LR model with WFV has demonstrated superior performance in this study, the overall high MAPE values indicate that forecasting during an economic crisis remains a challenging task. The study suggests that future research should aim to refine these models further, possibly by incorporating additional macroeconomic indicators and advanced techniques for handling extreme data variability. Such improvements could lead to more accurate and reliable forecasts, which are crucial for effective economic planning and decision-making in times of crisis.

Conflict of Interest

The authors have no conflicts of interest to declare relevant to the content of this article.

Acknowledgements

The authors acknowledge the comments from anonymous reviewers.

References

- Atkeson A, Ohanian LE. 2001. Are Phillips curves useful for forecasting inflation? *Quarterly Review* 25(1): 2-11.
- Armstrong JS, Collopy F. 1992. Error measures for generalizing about forecasting methods: Empirical comparisons. *International Journal of Forecasting* 8(1): 69-80.
- Bandara R. 2011. The Determinants of Inflation in Sri Lanka: An Application of the Vector Autoregression Model. *South Asia Economic Journal* 12(2): 271-286.
- Bergmeir C, Benítez J. 2012. On the use of cross-validation for time series predictor evaluation. *Information Sciences* 191: 192-213.
- Cristadoro R, Mario F, Reichlin L, Veronese G. 2005. A core inflation indicator for euro area. *Journal of Money, Credit and Banking* 37(3): 539-560.
- Cutler A, Cutler D, Stevens JR. 2012. *Ensemble Machine Learning: Methods and Applications*. New York: Springer.
- Goodhart C, Hofmann B. 2000. Do asset prices help to predict consumer price inflation? *The Manchester School* 68(s1):122-140:55-67.
- Granger C, Jeon Y. 2004. Thick modeling. *Economic Modelling* 21(2): 323-343.
- Hoerl AE, Kennard RW. 1970. Ridge Regression: Biased Estimation for Nonorthogonal Problems. *Technometrics* 12(1).
- Hyndman RJ, Koehler AB. 2006. Another look at measures of forecast accuracy. *International Journal of Forecasting* 22(4): 679-688.
- Inoue A, Kilian L. 2006. On the selection of forecasting models. *Journal of Econometrics* 137(2): 273-306.
- Jere S, Mubita S. 2016. Forecasting Inflation Rate of Zambia Using Holt's Exponential Smoothing. *Open Journal of Statistics* 6(2): 363-372.

- Jesmy A. 2010. Estimation of future inflation in Sri Lanka using ARIMA model. *Kalam* 5(1): 21-27.
- Marcellino M. 2002. Forecast Pooling for Short Time Series of Macroeconomic Variables. *IGIER Innocenzo Gasparini Institute for Economic Research* :Working Paper No 213.
- Marcellino M, Stock J H, Watson M W. 2000. A Dynamic Factor and Neural Networks Analysis of the Co-movement of Public Revenues in the EMU. *Italian Economic Journal* 8(2): 289-338.
- Ogutlu J, Schulz S, Torben P. 2012. Genomic selection using regularized linear regression models: Ridge regression, lasso, elastic net and their extensions. *BMC Proceedings* 6:10.
- Smola AJ, Schölkopf B. 2003. A tutorial on support vector regression. *Statistics and Computing* 14(3): 199-222.
- Stock JH, Watson MW. 1999. Forecasting inflation using a new index of aggregate activity. *Journal of Monetary Economics* 44(2): 293-335.
- Tradingview. 2023. Sri Lanka Inflation Rate YoY. SL: *Trading View*.
- Trevor H, Robert T, Jerome F. 2009. *The Elements of Statistical Learning*. 2 ed. New York: Springer.
- Ulke V, Sahin A, Subasi A. 2018. Comparison of Time Series and Machine Learning Models for Inflation Forecasting: Empirical Evidence from the US. *Neural Computing and Applications* 30(1): 1519-1527.
- Wickham H. 2016. *ggplot2: Elegant graphics for data analysis*. 2 ed. Melbourne: Springer.
- Wright JH. 2003. Forecasting U.S. Inflation by Bayesian Model Averaging. *International Finance Discussion Paper* 9: 1-33.