

Applicability of Random Forests Forecasting to International Currency Trade: An Investigation Through Language

Kamakshaiah Musunuru*, S.S. Prasada Rao**

Abstract

The goal of this research is to study the performance of foreign exchange trade in both India and China. India and China raised rapidly in recent times and there is abundant of speculation that these countries might reach to the level of few other developed nations as far as international trade is concerned. Whereas there isn't any doubt that these countries emerging as economic powers in the Asia-Pacific region, a lot of effort is required at international platform with respect to trade and commerce. One of such areas of competition is international currency trade. The aim of this study is to understand trends of currency trade in order to predict how likely these countries are going to emerge as best in the region. The study used certain secondary datasets from very reliable and authenticated sources. As far as statistical techniques are concerned, random walk forecasting methods were employed to test the study hypothesis. The study gathered certain evidence that though there are similarities in present and past performance, it is not likely to be the same in the future. However, the study concludes that random forests forecasting as a methodology is highly useful in studying trends in the data.

Keywords: Asia-Pacific Region, Currency Trade, International Trade, Random Walk Forecasting, Time Series Analysis

Introduction

By September 1, 2017 foreign exchange reserves in India increased to USD \$398,120 million. The averaged foreign

exchange reserves were USD \$206,149.58 million between 1998 and 2017. By September of 2017, the same had reached to an all-time high level of USD \$398,120 million and USD \$29,048 million which is a record low in September of 1998 (TE, 2017).

There were certain funny but unfavourable speculations done against American stock markets in 2007. The stock market performance was expected to be low by the very first quarter of 2007; this was the speculation at that time. However, the actual story was different. Many companies surpassed the market speculations and did performed well. Experts show the evidence in support of the fact that about the influence of ever aggrandizing economics of India and China (Bloomberg, 2017). Perhaps, this is one of the reasons for these countries to be viewed as major competitors in the region (Dar and Ahmad, 2014). The other reasons could be ever increasing consumerism, getting redefined through empowerment in education and affluence supported by credit disbursement from financial institutions. These lead to a plausible reasoning that India's and China's growth invincibly affects the US's growth and remain a robust driver for their economies. The other reason that supports the aforementioned idea is the amount of innovation that is happening in these countries. Markets have affordable products driven by innovation that might affect consumerism grow further and economies evolve more robust.

Some of the neoclassical researchers like Jagdish Sheth coined a new strategy known as Chindia. Some of the world's renowned MNCs like Nokia and GE successfully implemented a Chindia strategy. These researchers say that the shift in focus on these two emerging Asian economies

* Assistant Professor, GITAM School of International Business, GITAM University, Visakhapatnam, Andhra Pradesh, India.
Email: kamakshaiah.m@gmail.com

** Director, Academic Affairs & Professor, GITAM Institute of Management, GITAM University, Visakhapatnam, Andhra Pradesh, India. Email: director_academicaaffairs@gitam.edu

is necessary not only for the Western countries, but also Japan and South Korea, who will find themselves in a similar position as their population starts to age (Sangani, P., 2008).

One of the best anecdotes of rising importance of these countries can be illustrated by Coca-Cola. Coca-Cola is one of the world's giant entities that believed in Chindia strategy. While the cola biggie's growth has almost flattened out in the United States, in 2006, 70% of its growth and 80% of its profits came from international business, a significant part of this from China. Similarly, German auto major Volkswagen was among the first to enter the Chinese market and had a market share of 50% in the early 1990's (Tang, R., 2009). A few years ago, this came down to 18%, but the company figures out that it's still a healthy market share in what is soon going to be the world's largest auto market with 60 brands, compared with 37 in the United States.

Albeit the aforementioned facts, the economies of India and China have grown rapidly over the past couple of decades, and it is widely accepted that these two emerging giants will transform the global economy in numerous ways over the coming decades. Despite the importance of these countries, their strengths and weaknesses, the sources of their growth, and the missing ingredients to sustain high growth rates are not widely known.

Indian Economy

Since liberalization began in the 1980s, GDP growth has surged in India. The Elephant metaphor did not reflect the recent speed of India's transformation, which has been more like a tiger. Since 2003–2007, GDP growth has averaged 8.6% (NRC, 2010). India's growth would continue and increase in the coming decade if economic reforms continue, and are expanded, and large-scale structural changes are undertaken to support growth (IBEF, 2017). Exports have doubled in three years, and software exports doubled in the last two years. The exports-to-GDP ratio is "extremely low," he said, even though huge increases in foreign investment—over \$21 billion—is comparable to that seen in China. India can adapt quickly, as evidenced by India's telecommunications revolution. From 5 million telephone lines in 1991, India now has over 200 million lines. India's demography will very likely help sustain

this growth. India's population is younger than China's and exhibiting a rising rate of personal savings.

Chinese Economy

There are also counter views that China and India are not comparably sized global giants. China's trade is six times larger than India's (Gathani, B., 2004). Even more striking is the fact that the increase in China's trade level in 2007 (\$433 billion, valued using MER) was greater than India's total trade. India's share of the global economy today is still less than half of what it was at independence in 1948. India's economy is expanding rapidly; but, its trade is still less than 1% of the global total; whereas, China's trade is the second or third largest (TNAP, 2004). A similar disparity exists in foreign investment. For these reasons, Lardy expressed more optimism about China's growth than about India's. The competitive environment in China is more favourable and intense than it is in India, where certain sectors are protected from import competition. In China, with reduced tariffs, domestic firms face competition not just from foreign imports but from foreign firms operating in China. China spends three times as much on infrastructure as India.

China's main challenge is to rebalance its growth strategy, moving towards one that relies more on domestic demand and less on exports. Currently, household consumption is only 36% of GDP; whereas, in India, that figure is 50–60% (CRISIL, 2011). For sustained economic development, India needs more manufacturing, a more liberalized trade environment, and more flexible labour markets.

The conventional wisdom is: "India does software; China does hardware. Those are their paths to expansion." But China's hardware exports are growing much faster than India's software exports, which make up less than 5% of India's GDP. India will need to take advantage of relatively low wage rates to build up its labour-intensive manufacturing sectors.

Comparing the Two Countries

Total Factor Productivity (TFP) growth rates are important. Capital deepening—that is, an increase in capital intensity, usually measured as capital stock per labour hour—also plays a dramatic role in growth,

especially in China, and is the “major explanatory factor” in the differences between the two countries’ per capita annual growth. India averaged 4.8% between 2000 and 2005, about half of China’s 8.1% annual per capita GDP growth rate. This difference is also seen in the R&D expenditure differences: R&D intensity in India is less than 1%; in China, it is 1.4% (Angang, H., et al., 2003).

India has competitive costs and wage levels, but it needs large-scale firms to compete successfully. Labour market restrictions in India are that country’s greatest challenge. At the state level, though, India is deregulating and making labour markets more flexible. In China, where private firms are more productive than public firms, there is a great need to extend privatization (The Economist, 2017). China is restructuring rapidly and deepening regional specializations. India’s financial markets are more developed than China’s but India has a greater need to reduce regulatory restrictions in financial product markets (Merrill, S., Taylor, D., & Poole, R., 2010).

Review of Literature

Rossi, B. (2005) applied newly developed tests for nested models that are robust to the presence of parameter instability. The empirical evidence shows that for some countries we can reject the hypothesis that exchange rates are random walks. This raises the possibility that economic models were previously rejected not because the fundamentals are completely unrelated to exchange rate fluctuations, but because the relationship is unstable over time and, thus, difficult to capture by Granger Causality tests or by forecast comparisons. The authors also analysed forecasts that exploit the time variation in the parameters and find that, in some cases, they can improve over the random walk.

Rossi, B. (2005) did certain study on foreign exchange markets and its speculation. The study finds that the fluctuations in oil/commodity prices have been considered; however, they do not fully explain the volatility. The paper considers various economic indicators, selects the most impactful, and finalizes a model. The variables with the most impact come down to unemployment rate, oil price, consumer confidence index, and wheat prices. Ultimately, through in and out-of-sample forecasting, the authors could find that the model is not fully able to or accurately forecast short-term fluctuations, and yet does well in expressing long-term forecasts.

Zhang, S., et al. (2007) examines the monetary model of exchange rate determination for the US dollar exchange rates against the currencies of Canada, Japan, and the United Kingdom. The study utilized the co-integration technique for testing long-run relationship, and vector error correction model for short-run dynamics and out-of-sample forecasting. The existence of co-integration supports the long-run relationship among nominal exchange rate and a number of fundamental variables. The out-of-sample forecasting indicates that the nominal exchange rate forecasts from the VEC monetary model can be superior to random-walk based forecasts in a projection period of less than one year. This conclusion implies that the monetary model of exchange rate determination is a reliable tool for policy makers to evaluate their currency and the monetary authority should expect a much shortened response time to the monetary policy impulse in the surging trend of international economic integration.

Lam, L., et al. (2008) compares the forecast performance of the Purchasing Power Parity model, Uncovered Interest Rate Parity model, Sticky Price Monetary model, the model based on the Bayesian Model Averaging technique, and a combined forecast of all the above models with benchmarks given by the random-walk model and the historical average return. Empirical results suggest that the combined forecast outperforms the benchmarks and generally yields better results than relying on a single model.

Irena, M., Andrius, B., (2013) presents the main fundamental exchange rate forecasting models and discusses the advantages and drawbacks of the mentioned models. The research should help to explain why the forecasts can be not accurate.

Research Methods

The research is basically an exploratory study. The aim of the research is to perform prediction on foreign exchange rates for both India and China. So, besides from being exploratory, the research also falls in the ambit of comparative analysis as in methodology. India and China thought to be racing in the global economy and much of the Asian wealth depends on these two giants. There is more description given in the Introduction vide comparison.

Objectives

The aim of the study is to understand the current status of foreign exchange markets for both India and China. However, since this is a scholarly article and gives more emphasis on certain vigorous research methods. The study uses certain valid forecasting techniques like mean forecasting method in comparison with Random (Walk) Forests. These methods serve as mechanisms for forecasting foreign exchange rates for study sample, i.e., India and China. So, coming to the objectives, the following statements serve as objectives to the study:

- To study and evaluate if there exists any random behaviour of foreign exchange rates of India.
- To study and evaluate if there exists any random behaviour of foreign exchange rates of China.
- To forecast foreign exchange rates and understand the foreign exchange markets of India and China.
- To evaluate linear forecasting models such as random walk and its applicability on foreign exchange rates.

Hypotheses

As the aim of the study is to understand the performance of these two countries with respect to foreign exchange, the following statements deserve to be hypotheses for the study:

H₁: There exist significant differences in performance of currency trade between India and China in past and at present.

H₂: There will be significant differences in performance of foreign exchange trade between India and China.

Data

The data used for this study are secondary and they are retrieved from World Bank repository. World Bank serves abundant of economic yet financial data for variegated needs of academic and research. The data for this study were retrieved from <http://data.worldbank.org/indicator/PA.NUS.FCRF?locations=IN>. Table 1 shows the first few records of study data.

Table 1: First Few Records of Study Data

Country Name	China	India	Country Name	China	India
Country Code	CHN	IND	Country Code	CHN	IND
Indicator Name	Official exchange rate (LCU per US\$, period average)	Official exchange rate (LCU per US\$, period average)	Indicator Name	Official exchange rate (LCU per US\$, period average)	Official exchange rate (LCU per US\$, period average)
1960	2.461809455	4.761900004	2011	6.461461327	46.67046667
1961	2.461809455	4.761900004	2012	6.312332827	53.43723333
1962	2.461809455	4.761900004	2013	6.195758346	58.59784542
1963	2.461809455	4.761900004	2014	6.143434094	61.02951446
1964	2.461809455	4.761900004	2015	6.227488673	64.15194446
1965	2.461809455	4.761900004	2016	6.644477829	67.19531281

The above data show the foreign exchange value for both India and China. The dataset is 61×3 order data matrix. All rows represent years and columns represent year-wise foreign exchange rates for both India and China. The base currency for both exchange rates is USD. For instance, the foreign exchange rate for India for 2016 is 67.19 which means one USD is equal to approximately 67 Indian Rupees.

Research Methods

This study used two important theoretical models of forecasting. The first is forecasting by means method (Means Forecast) and the second is forecasting by random walk (RW) method. The first method, i.e., means forecasting assumes that the underlying data distribution

is identically independent distribution. As such, the means method is defined as follows:

$$Y_t = \mu + Z_t$$

where Y_t is the actual data and μ is the mean of the distribution Z_t is the error term, which is believed to be normally distributed. The forecasting function is given by,

$$Y_{n+h} = \mu$$

where μ is sample estimate of distribution mean.

The concept of random walk belongs to a machine learning methodologies known as random forest. Random forests or random decision forests are an ensemble learning method for classification, regression, and other tasks, which operate by constructing a multitude of decision trees at training time and outputting the class that is the mode of the classes (classification) or mean prediction (regression) of the individual trees. Random decision forests correct for decision trees' habit of overfitting to their training set. So, there will be two sets while fitting the actual data through random forests or walk. The first one is the training (set) data and the other is test (set) data. Usually, it is first the training which is performed over a subspace and later the prediction will be verified with the other subspace of the same set.¹ The random walk with drift model is,

$$Y_t = C + Y_{t-1} + Z_t$$

where Z_t is a normal *iid* error. Forecasts are given by,

$$Y_{n+h} = Ch + Y_n$$

If there is no drift, the drift parameter $c=0$. Forecast standard errors allow for uncertainty in estimating the drift parameter. Interestingly, there exists a close analogy between ARIMA and Random Walk. A general model can be given as,

$$\Delta y_t = \alpha + \beta t + \gamma y_{t-1} + \delta_1 \Delta y_{t-1} + \dots + \delta_{p-1} \Delta y_{t-p+1} + \epsilon_t$$

where α is constant and rest of the parameters are coefficients. If the model is brought under constraints that

¹ In machine learning, the random subspace method, also called attribute bagging or feature bagging, is an ensemble learning method that attempts to reduce the correlation between estimators in an ensemble by training them on random samples of features instead of the entire feature set.

α, β are zero, then this general model turned into random walk and the second constraint, i.e., $\beta=0$ makes the model random walk with a drift.

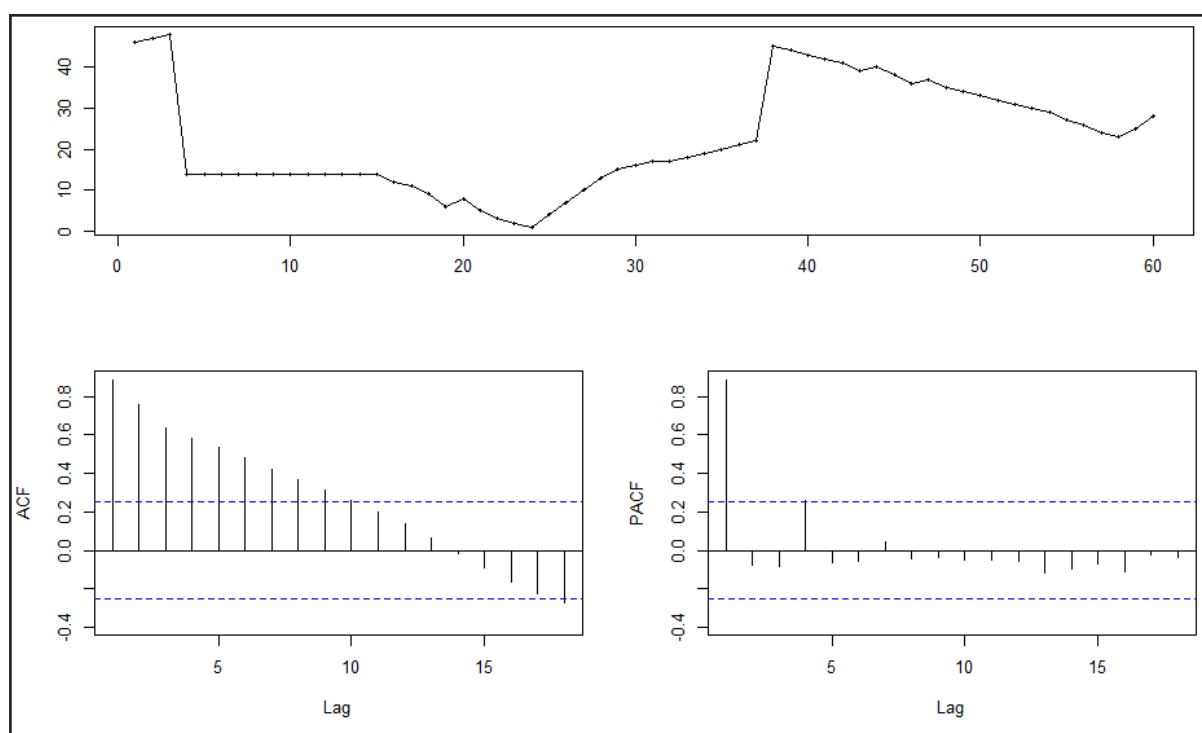
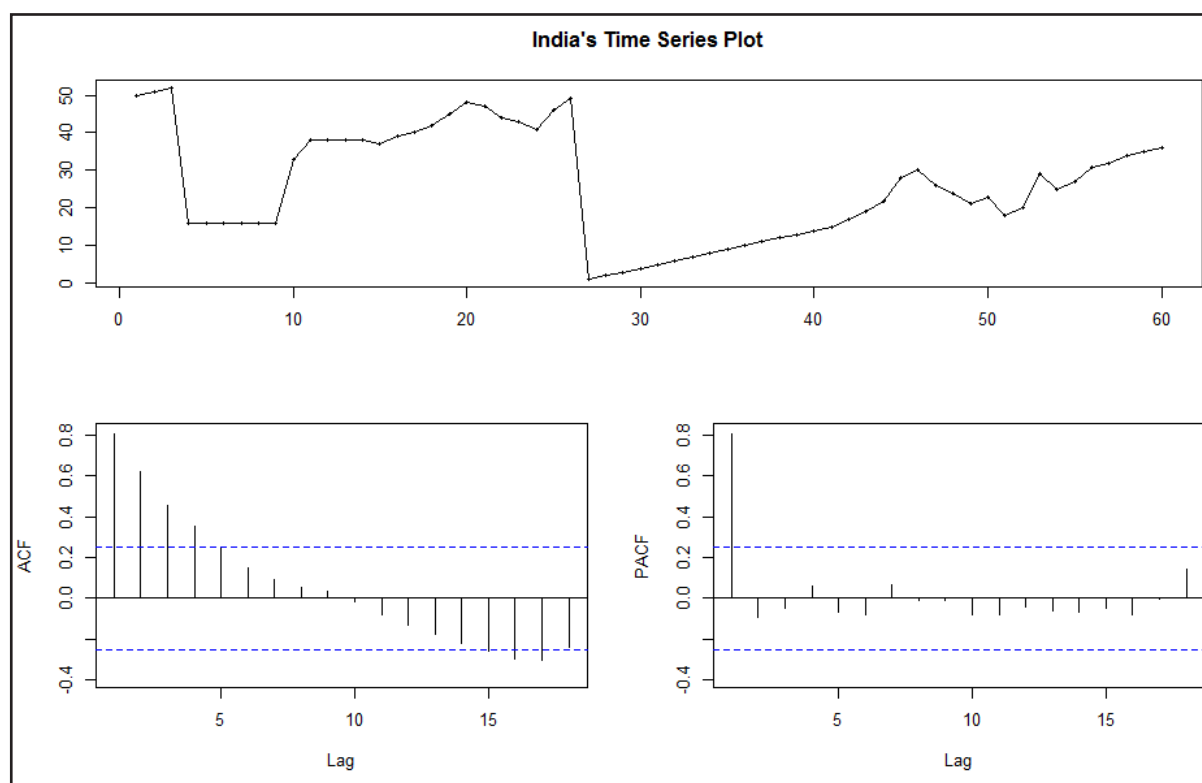
Statistical Tools

MS Excel together with R is used for this study. Excel was used only to prune or cure the data *ala* editing cases and fields, merging rows and columns etc. R is one of the programming language and also known as *lingua franca* of statistics. R has wonderful mechanisms to perform all sorts of analyses and its capabilities can be extended to publishing quality visuals. There are approximately 600 packages for all variegated needs of statistical analyses. R is used to perform forecasting on sample data. All the analysis was done using one of the renowned and widely used R package known as forecast. This package has many mechanisms to perform time series analysis on survey or sampled data. Both means forecasting and random walk methods were performed by using forecast package in R.

Analysis

The analysis was done by three main tasks. *First*, importing the datasets to R as in the form of CSV files. *Second*, performing modeling, as in both training and testing, on the sampled data. *Third*, forecasting and plotting. The data are time-bound data which means for R it is time series data. The very first column of the dataset is time variable. R has automatic conversions to deal with time-bound data. The very first method in the analysis is visual investigation. The following is the time series graph which provides certain insights into the dataset.

Fig. 1: (a) is time series plot for China's foreign exchange rate. Figure 1 (b) is time series plot for India's foreign exchange rate. The plots also have graphs for both autocorrelation function (ACF) and partial autocorrelation functions (PACF). It is not very obvious that there is any visible speculative trend. However, there are certain ups and troughs; as such, it might be premature to infer that there exist seasonal changes. The data need to be tested for stationarity and seasonality. However, a careful observation on ACF plots gives us a valuable insights that both datasets are roughly random. The following is the R Output for Augmented Dickey Fuller Tests on sample datasets.

**Fig. 1: (a) Time Series Plot for China Foreign Exchange Rate****Fig. 1: (b) Time Series Plot for India Foreign Exchange Rate**

R Output 1: Tests for Stochasticity.

```
> adf.test(feds[,2])

Augmented Dickey-Fuller Test

data: feds[, 2]
Dickey-Fuller = -1.5061, Lag order = 3,
p-value = 0.7744
alternative hypothesis: stationary

> adf.test(feds[,3])

Augmented Dickey-Fuller Test

data: feds[, 3]
Dickey-Fuller = -2.0542, Lag order = 3,
p-value = 0.5529
alternative hypothesis: stationary
```

The above output was taken from R for ADF test on time series data of China and India. In both cases the p value observed as greater than 0.05. So, at 5% significance level, it is not possible to reject the null hypothesis that the data are not stationary. So, it is clear that the data are stochastic. This gives a notion that the random walk method might be suitable for prediction. The following is the test for the seasonality.

R Output 2: Tests for Seasonality.

```
> PP.test(as.numeric(feds[,2]))

Phillips-Perron Unit Root Test

data: as.numeric(feds[, 2])
Dickey-Fuller = -3.0176, Truncation lag
parameter = 3, p-value = 0.1635

> PP.test(as.numeric(feds[,3]))

Phillips-Perron Unit Root Test

data: as.numeric(feds[, 3])
Dickey-Fuller = -2.4712, Truncation lag
parameter = 3, p-value = 0.3843
```

There is a wide variety of tests to check the significance of seasonality of time series data.² The above R Output

² One of the other ways in testing significance of seasonality is through log-likelihood. The log-likelihood as a ratio follows

shows the Phillips-Perron Test for autocorrelations. The p value is found to be sufficiently greater than 0.05; so, it is not possible to reject null hypothesis that the data under study are sufficiently random. Phillips-Perron Test tests only autocorrelations. So, any decision regarding seasonality is merely based on correlation but not on error variable. It might be better to recheck or reevaluate this decision through any other alternative mechanisms. The forecast package of R has different means to test seasonality of time series data (Hyndman R.J., 2017 & Hyndman R. J., and Khandakar Y., 2008).³

R Output 3: Seasonality check for foreign exchange rates for China and India.

```
> library(forecast)

> fit <- tbats(as.numeric(feds[,2]))

> seasonal <- !is.null(fit$seasonal)

> seasonal

[1] FALSE

> fit <- tbats(as.numeric(feds[,3]))

> seasonal <- !is.null(fit$seasonal)

> seasonal

> fit <- ets(feds[,3])
> ll <- logLik(fit)
> fit$loglik
```

chi-square distribution. Time series data when fitted give log-likelihood of

³ The package **forecast** has sufficient methods to take care of seasonality. Hyndman RJ, *et al.* used exponential smoothing state space model with Box-Cox transformation along with ARMA errors. One way could be through log-likelihood. Both methods ETS and TBATS in **fma** and **forecast** yield log-likelihood. It might be possible to compute chi-square statistic through log-likelihood value based on which the decision can be taken that whether the underlying distribution has seasonality or not. For instance, for Indian time series data can be as follows:

The p value is 1 which exactly means there isn't any whit of seasonality in the Chinese time series data. However, a careful observation could reveal that the Phillips-Perron Test looks to be better than Chi-square test. One of the important weaknesses in Chi-square test is that actually, it is omnibus test, which means the statistic will be testing the sample size instead of sample attribute.

```
[1] -225.2308
> ll
'log Lik.' -225.2308 (df=3)
> 1-pchisq(ll, 3)
'log Lik.' 1 (df=3)
```

[1] FALSE

The seasonal component is absent in both the data variables, i.e., China and India. So, now it seems plausible to conclude that the data are sufficiently random for prediction. Now that it is known that the data are sufficiently random, it might be possible to predict through random walk method.

Prediction Through Random Walk Method

By the definition, random walk is machine learning method which needs two inputs - one, training set and the other, test set. As it was mentioned in the research methods, the entire dataset will be divided into two subsets, one for the training and the other for testing. Such mechanisms are not so naïve in the area of machine learning. The study dataset has 60 rows (records) and 3 columns (fields).

```
> dim(feds)
[1] 60 3
```

So, it is quite rational to make two subsets each with 30 records keeping fields as it is in the dataset. The random walk method is implemented on the first 30 records of the study data. The following is the procedure to perform random walk method in R.

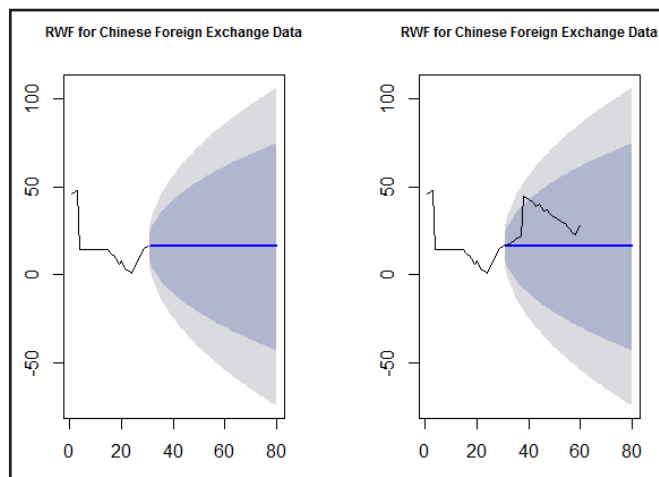


Fig. 2 (a): Random Walk Forecasting Plot for Chinese Foreign Exchange Data.

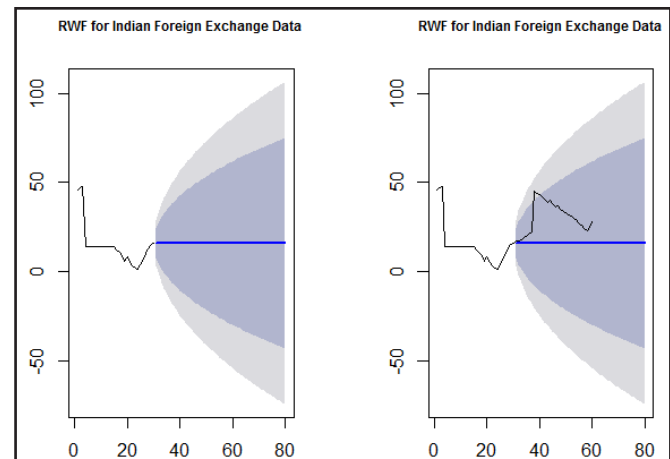


Fig. 2: (b) Random Walk Forecasting Plot for Indian Foreign Exchange Data

The following are the fitted and residuals of the RWF model in R.

Table 2: Fitted Values and Residuals for Random Walk Forecasting for Chinese Foreign Exchange Rate Data

```
> cbind(fit$fitted, fit$residuals,
fit2$fitted, fit2$residuals)
```

Time Series:

Start = 1

End = 30

Frequency = 1

	fit\$fitted	fit\$residuals	fit2\$fitted	fit2\$residuals
1	NA	NA	NA	NA
2	46	1	50	1
3	47	1	51	1
4	48	-34	52	-36
5	14	0	16	0
6	14	0	16	0
7	14	0	16	0
8	14	0	16	0
9	14	0	16	0

10	14	0	16	17
11	14	0	33	5
12	14	0	38	0
13	14	0	38	0
14	14	0	38	0
15	14	0	38	-1
16	14	-2	37	2
17	12	-1	39	1
18	11	-2	40	2
19	9	-3	42	3
20	6	2	45	3
21	8	-3	48	-1
22	5	-2	47	-3
23	3	-1	44	-1
24	2	-1	43	-2
25	1	3	41	5
26	4	3	46	3
27	7	3	49	-48
28	10	3	1	1
29	13	2	2	1
30	15	1	3	1

Figure 2 (a) shows both training and actual series for random walk forecasting on Chinese foreign exchange data. Figure 2 (b) shows both training and actual series for random walk forecasting on Indian foreign exchange data. From both the figures, it seems that the forecasting is same for both countries. It makes little confusion that whether it is due to the fault of the model or method. Table 1 makes it clear about fitted values and residuals for the both countries. From the table, it very clear that the values are not same. For instance, Table 3 shows the actual data; from the data, it is clear that the actual values for China and India are 1.49 and 7.86, respectively, for the year 1980. And from Table 2 it is clear that the fitted value is 2 (1.49) and 43 (7.86) with their respective residuals, i.e. -1 and -2, respectively. So, the model values are different.

Table 3: Tail Part of the Study Dataset

```
> feds[24:30, ]
```

Country.Name	China	India
24 1980	1.498399999	7.862944701
25 1981	1.704533333	8.658522817
26 1982	1.892541666	9.455131933
27 1983	1.975674999	10.09889824
28 1984	2.320041666	11.36258333
29 1985	2.936658333	12.36875
30 1986	3.452791667	12.61083333

The following are the accuracy measures for both China and Indian foreign exchange rate data.

Table 4: Accuracy Measures for Both China and India RWF Method

```
> accuracy(fit1)
```

	ME	RMSE	MAE	MPE	MAPE	MASE	ACF1
Training set	0.2241903	0.6555713	0.4936788	-95.06225	143.5371	0.7063767	0.003487063

```
> accuracy(fit2)
```

	ME	RMSE	MAE	MPE	MAPE	MASE	ACF1
Training set	-1.586207	11.73853	4.758621	-166.3677	181.4842	1	-0.0488535

It shows that the RWF method is better for Chinese foreign exchange rate compared with that of India. RMSE for China is lesser than that of India, and it is true for rest of the other measures. The accuracy measures for China show that the RWF prediction works better for Chinese foreign exchange trade compared with that of India.

CONCLUSION

Indian and China are two robust economies in Asia-Pacific region. There is lot of literature that these two markets have abundant of influence on trade and economy of the region. The aim of this paper is to study the foreign exchange trade of India in contrast to China. The study assumes that the foreign exchange trade in India and China is not significantly different. The analysis was performed on secondary datasets of foreign exchange rates for both countries using random walk forecasting methodology. The analysis shows that though there isn't much difference in foreign exchange status for present and past for both India and China, the analysis shows that the predictions through methodology, i.e., random walk forecasting are not same. While random walk method could do well for China, is failed for India due to unknown reasons. Hence, though there isn't evidence in support of first hypothesis, i.e., the differences are not significant with respect to present and past performance, the same observation can't be applied for the other hypothesis. The study could find that the differences in performance with respect to foreign exchange could be significant in future. So, the study concludes that though the foreign exchange rates could appear to be governed by certain random processes, but there exist differences for predictions. More research is required to know the reasons as to why different methods are needed for prediction in spite of the similarities in the data.

Acknowledgements and Declaration of Interest

The authors report no conflicts of interest. The authors alone are responsible for the content and writing of the paper.

References

The Economist. (2017). India Foreign Exchange Reserves 1998-2017. Retrieved from <https://tradingeconomics.com/india/foreign-exchange-reserves>

Bloomberg. (April 07, 2017). Why India could be the winner of a US-China trade war. Retrieved from <https://economictimes.indiatimes.com/news/economy/foreign-trade/why-india-could-be-the-winner-of-a-us-china-trade-war/articleshow/58059316.cms>

Dar, B. A., & Ahmad, S. (2014). Major bilateral issues between China and India. *Arts Social Sci J* 5(64). doi: 10.4172/2151-6200.1000064

Sangani, P. (Feb 08, 2008). India, China to impact global economy. Retrieved from <http://economictimes.indiatimes.com/articleshow/2765861.cms>

Tang, R., (2009). The Rise of China's Auto Industry and Its Impact on the U.S. Motor Vehicle Industry. CRS Report for Congress. Retrieved from <https://fas.org/sgp/crs/row/R40924.pdf>

NRC. (2010). *The Dragon and the Elephant: Understanding the development of innovation*. National Academy of Sciences. ISBN: 978-0-309-15160-3. Pg. No. 6.

IBEF, (2017). About Indian Economy Growth Rate & Statistics. Retrieved from <https://www.ibef.org/economy/indian-economy-overview>

Gathani, B. (2004). People of Indian origin to play larger role in development. Retrieved from <http://www.thehindubusinessline.com/2004/10/08/stories/2004100800540900.htm>

TNAP. (2004). The Dragon and the Elephant: Understanding the Development of Innovation Capacity in China and India. Retrieved from <https://www.nap.edu/catalog/12873/the-dragon-and-the-elephant-understanding-the-development-of-innovation>

CRISIL. (2011). Raising the growth bar. Retrieved from <https://www.crisil.com/pdf/corporate/india-raising-the-bar.pdf>

Baldacci, E., Callegari, G., Coady, D., Ding, D., Kumar, M., Tommasino, P., & Woo, J. (2010). Public Expenditures on Social Programs and Household Consumption in China. International Monetary Fund Working Paper. Retrieved from <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.167.6981&rep=rep1&type=pdf>

Angang, H., Linlin, H., & Zhixiao, C. (2003). China's economic growth and poverty reduction (1978-2002). Retrieved from <https://www.imf.org/external/np/apd/seminars/2003/newdelhi/angang.pdf>

The Economist. (2017). Reform of China's ailing state-owned firms is emboldening them. Retrieved from <https://www.economist.com/news/finance-and>

- economics/21725293-outperformed-private-firms-they-are-no-longer-shrinking-share-overall
- Merrill, S., Taylor, D., & Poole, R., (2010). *The Dragon and the Elephant: Understanding the development of innovation capacity in China and India*. The National Academies Press.
- Rossi, B. (2005). Are Exchange Rates Really Random Walks? Some Evidence Robust to Parameter Instability. Retrieved from <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.497.1697&rep=rep1&type=pdf>
- Palombizio, E., & Morris, I. (2012). Forecasting Exchange Rates using Leading Economic Indicators. Palombizio and Morris, 1:8. Retrieved from <http://dx.doi.org/10.4172/scientificreports.402>
- Zhang, S., & Lowinger, T. C. (2007). The monetary exchange rate model: Long-run, short-run, and forecasting performance. *Journal of Economic Integration* 22(2), 397-406
- Lam, L., Fung, L., & Yu, I.-W. (2008). Comparing Forecast Performance of Exchange Rate Models. Working Paper 08/2008. June 2008. Retrieved from http://www.hkma.gov.hk/media/eng/publication-and-research/research/working-papers/HKMAWP08_08_full.pdf
- Hastie, T., Tibshirani, R., & Friedman, J. (2008). *The elements of statistical learning* (2nd ed.). Springer. ISBN 0-387-95284-5.
- Hyndman, R. J. (2017). Forecast: Forecasting functions for time series and linear models. R package version 8.1. Retrieved from <http://github.com/robjhyndman/forecast>.
- Hyndman, R. J., & Khandakar, Y. (2008). Automatic time series forecasting: The forecast package for R. *Journal of Statistical Software*, 26(3), 1-22. Retrieved from <http://www.jstatsoft.org/article/view/v027i03>