

Quantitative trading models based on Sufficient Dimension Reduction and Ensemble Learning

Jiaheng Wang*

School of Statistics, Renmin University of China, Beijing, China, 100872

* Corresponding Author Email: loganlogan0807@163.com

Abstract. With the continuous development of the securities market, constructing quantitative trading strategies with strong generalization ability the adaption to dynamic and changeable market environment is becoming more and more momentous in the field of quantitative finance. Under such circumstances, this paper proposes a stock prediction model based on sufficient dimension reduction theory and ensemble learning, which can be deployed on quantitative trading strategy. The proposed model on the one hand alleviates the curse of dimension by using the sufficient dimension reduction method and maximally retains variation of stock factors, on the other hand employs ensemble learning technique which can weigh model bias and variance to address the overfitting or underfitting problems. Furthermore, since ensemble learning are model free, different sub-models can be applied to improve the generalization ability of the algorithm on predicting occasions. In the section of accuracy comparison and quantitative trading strategy experiments based on real stock datasets, the proposed model demonstrates the best prediction results and adequate robustness compared to other existed methods.

Keywords: Quantitative trading; sufficient dimension reduction; multi-factor model; machine learning.

1. Introduction

Under the current situation of rapid development in information technology and securities market, investors are faced with a more dynamic investment feasible set and more complex market environment [1]. Quantitative trading refers to the process of adopting tremendous amounts of financial data, constructing corresponding machine learning models and algorithms in computers and automatically executing trading orders [2]. Since its introduction, quantitative trading has developed rapidly by virtue of its efficient data processing capability, automated trading technique and other distinguished features, which enable investors to gain higher excess returns in the capital markets stocks and futures [3][4]. Therefore, developing quantitative machine learning models with high efficiency and strong generalization ability becomes a popular direction in financial investment field. Multi-factor model is the most extensively used and maturely-developed model in quantitative investment, which is essentially a model for asset pricing based on modern financial investment theories such as Capital Asset Pricing Model (CAPM), Arbitrage Pricing Theory (APT), Markowitz portfolio theory (MPT), etc. Fama and French conducted an experimental study using historical data of U.S. stock market to determine the factors which influence price-earnings and proposed that the price-earnings is linearly affected by three factors, which are excess market return, HML (high minus low) and SMB (small minus big) [5]. This theory has been improved subsequently with the proposal of the four-factor model containing UMD (up minus down) and Fama-French five-factor model containing CMA (conservative and aggressive) and RMW (robust minus weak) [6].

As data science theory develops, multi-factor quantitative trading models are divided into two parts: factor selection and predictive model construction. There are more than 500 common factors in quantitative field, resulting in the failure of both traditional statistics models such as regression analysis and machine learning models such as decision trees because of the curse of dimension. The practicable methods are to select the significant factors or reduce the dimension of factors before making predictions. In terms of factor selecting, Xie (2017) used regularization methods elastic net and Lasso to impose penalties on regression factors in order to select significant factors and construct

the corresponding portfolios [7]. This type of multi-factor model is a kind of single linear model, which means it is not perfect enough to extract the factor information and it is prone to cause overfitting and underfitting problems. Fang (2021) applied the optimum feature subset selection method and used factor IC indicator to select the quality factors [8]. However, this type of method will inevitably discard a certain amount of information related to the responding variable and has a high computational complexity.

For the dimension reduction method, Ding (2020) extracted the principal components of stock factors as predictors to explore the share price movement [9]. Though with simplicity and convenience for calculation, the principal component analysis does not refer the variation of responding variable when extracting the variation of the factors because of its property being an unsupervised learning method and linear dimension reduction method. Therefore, it is problematic to achieve a better generalization when constructing prediction model. In addition, supervised dimension reduction methods and nonlinear dimension reduction techniques such as partial least squares, kernel principal component analysis, independent component analysis, and linear discriminant analysis could not guarantee the maximal retention of the effective information between the original predictors and the responding variable as well.

For prediction models, Zhao (2023) used statistical learning models ARIMA and GARCH to predict share price in short term [10], ending up with unsatisfactory prediction results due to the inefficiency of China's A-share market, where the share price changes with complex nonlinear features. With the proposal of deep learning and reinforcement learning algorithms, Li (2022) and An et al. (2022) respectively combined LSTM and Q-learning with multifactor model for share price prediction and obtained lower prediction errors [11-12]. However, China's A-share market is immeasurably influenced by policy and other non-certainty impacts, the factors as well as the prediction model face the risk of failure at any time. The poor interpretability of deep learning and reinforcement learning models also leads to the fact that they fail to be the optimal methods for prediction.

In conclusion, this paper proposes a multi-factor quantitative trading model based on sufficient dimension reduction and ensemble learning. The innovations and advantages of this method are mainly in the following aspects. Firstly, as a supervised dimension reduction method, sufficient dimension reduction can ensure that the responding variable is conditionally independent of the original stock factors, which theoretically maximizes the retention of dependence information between responding variable and predicting factors [13]. Second, the factors after sufficiently dimension method are involved as the predictors of ensemble learning machine, which can enhance the performance of weak prediction machines and obtain results with extensive accuracy and robustness. Finally, the predictive sub-models of ensemble learning are flexible, which enable us to obtain factor importance and provide a reference basis for decision-making if the sub-model is interpretable [14]. In the analysis of market data, the proposed model gains the highest accuracy in prediction and produces the best result in back testing compared with other models.

2. The basic fundamental of sufficient dimension reduction and ensemble learning

2.1. The theory of sufficient dimension reduction

Sufficient dimension reduction is the process of finding several linear combinations of original predictor variables to replace them without missing any dependence information between original predictor variables and responding variable. For one-dimensional responding variable Y and p -dimensional predictor variables $X = (X_1, X_2, \dots, X_p)$, the idea of sufficient dimension reduction is to find a $p \times d$ matrix that ensures the conditional independence between Y and X given the matrix $B^T X$. This can be expressed by the formula:

$$Y \perp\!\!\!\perp X | B^T X \quad (1)$$

It can be seen that what is desired to be found is the space $Span(B)$ spanned by the column vectors of the matrix B . Due to the non-uniqueness of $Span(B)$, Cook introduced the central reduction subspace $S_{Y|X}$ in 1998, which is defined as the intersection of all reduction space $Span(B)$. Hereinafter this paper introduces the method which will be used later to estimate the central dimension reduction subspace - Slice Inverse Regression. Slice Inverse Regression (SIR) was proposed by Li (1991), which is the earliest and most common classical method for estimating the central dimension reduction subspace [15]. When X obeys the elliptic isometry distribution (from which the linear mean assumption can be inferred), it can be proved that the spanned space of conditional covariance of X is included in $S_{Y|X}$, which can be expressed as:

$$Span(Cov(X|Y)) \subseteq S_{Y|X} \quad (2)$$

Formula (2) indicates that the eigenvectors corresponding to nonzero eigenvalues are a set of bases for the central reduction dimensional subspace $S_{Y|X}$ [16]. Therefore, the essence of SIR is to estimate the conditional covariance array. The specific algorithm is as follows:

(1) Suppose we have the dataset $D = \{(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)\}$, where $Y_i, i = 1, 2, \dots, n$ is one-dimensional and $X_i, i = 1, 2, \dots, n$ is p-dimensional vector;

(2) Standardize matrix $X_{n \times p}$ with Let $\hat{\Sigma} = Var(X)$ and $\bar{X} = E(X)$:

$$Z = \hat{\Sigma}^{-\frac{1}{2}}(X - \bar{X}) \quad (3)$$

(3) Uniformly segmented Y into h intervals and $J_l, l = 1, 2, \dots, h$ represents the h -th interval. Now we can estimate $E(Z|Y \in J_l)$:

$$E(Z|Y \in J_l) = \frac{E(ZI(Y \in J_l))}{E(I(Y \in J_l))}, l = 1, 2, \dots, h \quad (4)$$

(4) Estimate $Var(E(Z|Y))$:

$$\hat{A} = \sum_{l=1}^h E(I(Y \in J_l) E(Z|Y \in J_l) E(Z^T | Y \in J_l)) \quad (5)$$

(5) Determine the eigenvectors corresponding to the first d largest eigenvalues of kernel matrix $\hat{A}$, written as $\hat{v}_1, \hat{v}_2, \dots, \hat{v}_d$, which are a set of bases for the central reduction dimensional subspace $S_{Y|X}$. Obviously, if we let $\hat{\beta}_k^T = \hat{A}^{-0.5} \hat{v}_k, k = 1, 2, \dots, d$, the result of SIR is:

$$\hat{\beta}_k^T (X - \bar{X}), k = 1, 2, \dots, d \quad (6)$$

2.2. The price prediction model based on SDR and ensemble learning

Ensemble learning achieves the purpose to enhance learning results and obtain distinguished generalization by fitting multiple weak learning machines and integrating them together. Under the guidance of this idea, different types of ensemble learning algorithms can be created by changing the training data and the training features of single learning machine, as well as the combination of

learning results. The most common ensemble learning algorithms can be classified as Bagging and Boosting.

Specifically, Bagging algorithm was proposed by Breiman in 1996 [15] and is one of the earliest ensemble learning algorithms. Though simple in structure, it is remarkable in performance. The algorithm generates new training subsets by resampling methods such as Bootstrap, then trains individual learning machines with different training subsets, and finally assembles them into a whole learning machine. Among all Bagging algorithms, Random Forest Model is the most representative algorithms, whose structure is shown in Figure.1.

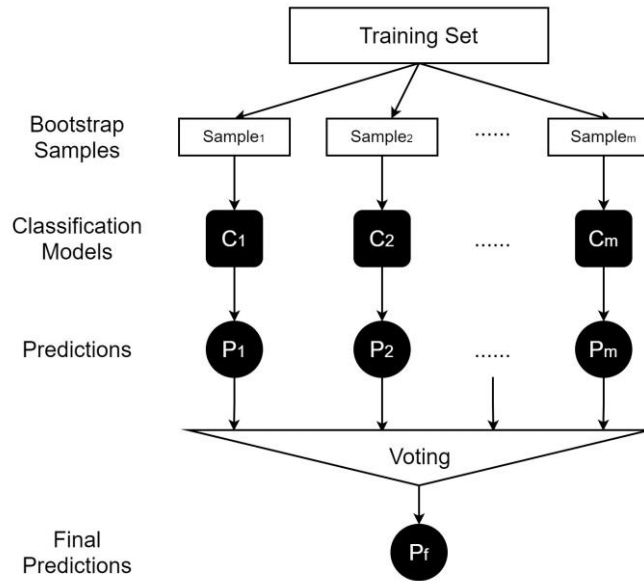


Figure 1. The basic structure of Random Forest model.

Boosting algorithm is an iterative method that converts weak learning machines into strong learning machines. AdaBoost algorithm is a representative method of Boosting, whose stochasticity comes from adjusting the weights of samples and weak learning machines. It prompts each classifier to learn as many samples as possible which were learnt incorrectly by the previous classifier and allows learning machines with high accuracy to have higher weights in the assembling of the final result, as is shown in Figure.2.

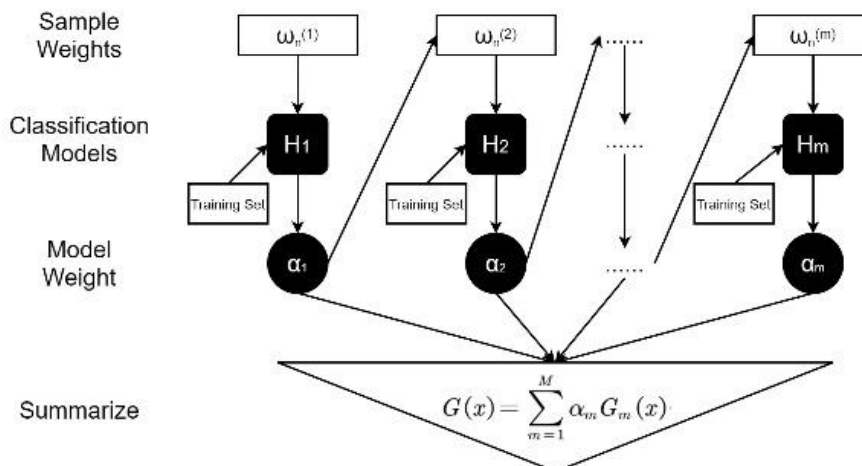


Figure 2. The basic structure of Adaptive Boosting algorithm.

The model proposed in this paper combines sufficient dimension reduction theory and ensemble learning methods. Assuming that there are K types of stock factors and each type of factor contains $m_1, m_2, \dots, m_K$ series of subfactors. For factor i , the data matrix from period1 to t is:

$$F_i = \begin{pmatrix} X_{11}^{(i)} & X_{12}^{(i)} & \dots & X_{1m_i}^{(i)} \\ X_{21}^{(i)} & X_{22}^{(i)} & \dots & X_{2m_i}^{(i)} \\ \vdots & \vdots & \ddots & \vdots \\ X_{t1}^{(i)} & X_{t2}^{(i)} & \dots & X_{tm_i}^{(i)} \end{pmatrix}_{t * m_i} \quad (7)$$

The corresponding share price factor from period 2 to $t + 1$ is $Y = (Y_2, Y_3, \dots, Y_t, Y_{t+1})^T$. Then apply SIR method to factor matrix F_i , setting parameter $h=5$ and $d=1$ before having a t -dimensional factor $F_{iSIR} = (Z_{i1}, Z_{i2}, \dots, Z_{it})^T$. Repeat this process on every stock factor and obtain a reduced factor F_{SIR} matrix ultimately:

$$F_{SIR} = (F_{1SIR}^T, F_{2SIR}^T, \dots, F_{KSIR}^T) = \begin{pmatrix} Z_{11} & Z_{12} & \dots & Z_{1K} \\ Z_{21} & Z_{22} & \dots & Z_{2K} \\ \vdots & \vdots & \ddots & \vdots \\ Z_{t1} & Z_{t2} & \dots & Z_{tK} \end{pmatrix}_{t * K} \quad (8)$$

At last, train an ensemble learning machine with F_{SIR} as predicting variables and price factor Y as responding variables using rolling prediction regulation (presented in Section 4). The details of sufficient dimension reduction and prediction through ensemble learning model are shown in the Figure.3.

$$\hat{Y} = EnsembleLearningMachine(F_{SIR}, \theta) \quad (9)$$

For each certain factor in this model, attain data for the factor with a series of subfactors in it and apply sufficient dimension reduction to it. Train ensemble learning predictive model with shrunken one-dimensional factors and price factor Y .

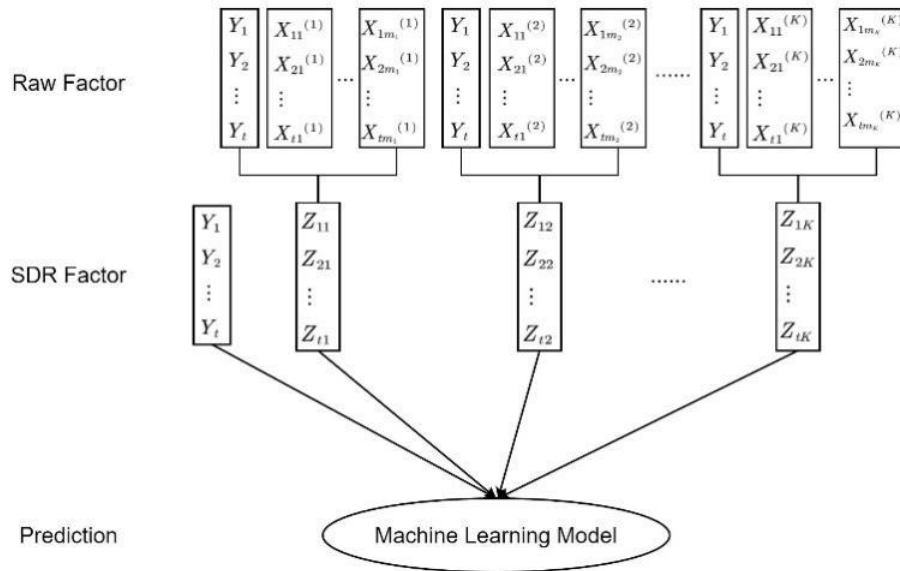


Figure 3. Model details.

3. Results

3.1. Data preparation and experimental Setup

The crucial part of predicting share price lies in the selection of factors. When the selected factors contain more market information related to price fluctuations, the prediction results achieve will be more accurate and robust. The traditional method of selecting valuable factors is to classify and select

the factors according to the Information Coefficient (IC), in which t-test will also be employed supplementarily to verify the independence of the factors. This process will be lengthy and complex in actual application scenarios when there are too many factors with ambiguous market implication, leading to the ignorance and removal of many factors which may be pivotal in some day in the future under such a volatile market. Since the significance of factors are changing all the time, it is also necessary to repeat this process after a period of time to validate the significance of the factors.

The usage of sufficient dimension reduction method omits the process above and only needs to prepare stock factor data for dimension reduction. Not only does this method simplify the process of factor data preparation greatly, but also makes the loss of information to reach a minimum level and avoids factor decay at the same time. This paper uses the quantitative factor data in DataYse! (<https://uqer.datayes.com/>). There are a total of 244 stock factors in the dataset including 7 main categories and 11 sub-categories. Adhering to the principle of computability and accessibility, this paper remove 16 analyst-related factors and the finally there are 228 factors affiliated to 11 sub-categories being adopted in this paper. Finally, there are 7 main factors, which are HML, SMB, UMD, CMA, RMW, technical indicators, and sentiment factor.

In this paper, we obtain the stock trading data of Gree Electric Appliances stock (sz000651) for 729 days from 4 January 2010 to 28 November 2012 for empirical analysis, during which period the company did not operate any business practices such as suspension, ex-rights or dividends. Gree Electric Appliances is the most representative brand of traditional home appliance manufacturing industry, which has been operating well and are worth being invested in the past few decades. This paper aim to construct a single stock daily frequency trading strategy, indicating that the share price should be predicted daily. For this certain stock, use each factor on the previous day and closing price of serval previous days as predicting variables and closing price on the same day as a responding variable. A rolling training prediction mechanism is constructed, which uses 500 days of data as a training set to predict the next 20 days of share prices, as shown in Figure.4.

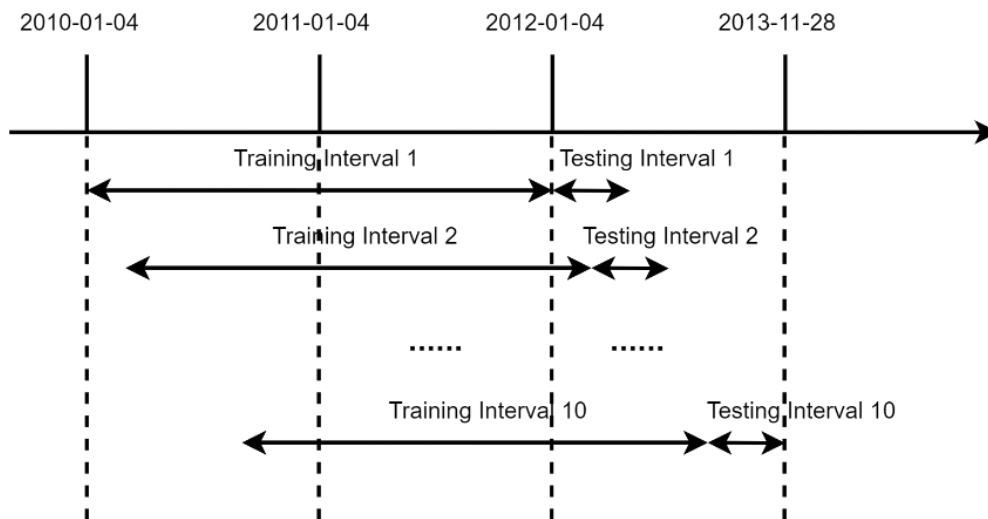


Figure 4. Rolling training regulation.

The effectiveness of model prediction is evaluated by using Root Mean Squared Error (RMSE), which measures the average error of the predictive price for each day, as defined below:

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (Y_i - f(x_i))^2} \quad (10)$$

In back test of the trading strategy constructed, use annualized return, Sharpe Ratio and maximum drawdown to measure the excellence of the strategy. The annualized return is calculated with the compound interest formula, as defined below:

$$Annualized\ Returns = \left(\frac{Total\ Assets}{initial\ Assets} \right)^{250/TradingPeriod} \quad (11)$$

The Sharpe Ratio measures the amount of excess return in each unit of risk, as defined below:

$$Sharpe\ Ratio = \frac{R_p - R_f}{\sigma_p} \quad (12)$$

R_p Is the average profit of each transaction and R_f is Risk-free interest Rate, here taken as 3%.
 σ_p Is standard deviation of each profit and loss? Maximum drawdown is an important measure of the risk of a strategy and can be interpreted as the maximum magnitude of loss that occurred in the trading history, with its value being the maximum value from highs to late lows in capital plot:

$$Max\ Drawdown = \max \left\{ \frac{D_i - D_j}{D_i} \right\}, i < j \quad (13)$$

A satisfactory strategy has high annualized returns and Sharpe ration with maximum drawdown as low as possible.

3.2. Analysis of experimental results

This section applied 6 common ensemble learning models and 2 simple machine learning models to conduct rolling predicting regulation. The result outputs are shown in Table 1.

Table 1. RMSE of different prediction models.

Learning Machine	11 SDR factor	11 SDR factor +3 previous price
Decision Tree Regressor	1.106	0.862
Linear Regressior	0.706	0.602
Random Forest Regressor	0.633	0.470
AdaBoost Regressor	0.548	0.498
GBDT Regressor	0.599	0.530
XGBoost Regressor	0.617	0.498
LigntGBM Regressor	0.584	0.506
Stacking	0.725	0.590

According to Table.1, the following conclusions can be summarized. Firstly, the addition of antecedent price information to the predicting variables reaches even lower RMSE; secondly, the RMSE of ensemble learning machine is much lower than that of the simple machine learning model; lastly, the Random Forest regressor achieves the lowest RMSE (0.455) among all models in experiment. It also can be seen that RMSE of Boosting algorithm is slightly higher than Random Forest, and Stacking has the highest RMSE.

In order to estimate the effectiveness of the random forest regression predictions in real market back test, this paper designed a simple trading strategy. Assuming an initial capital of \$100,000, the strategy uses all capital to buy shares when the predicting price of the next day rises by more than 1% while it sells all of the position when the predicting price of the next day falls by 1%. This is the simplest trading strategy which can estimate the lower bound of profitability of a given quantitative model. Implement the strategy in Python and make a graph WO show detail information such as buy and sell points and asset plot for 200 days of trading, which are shown in Figure.5.

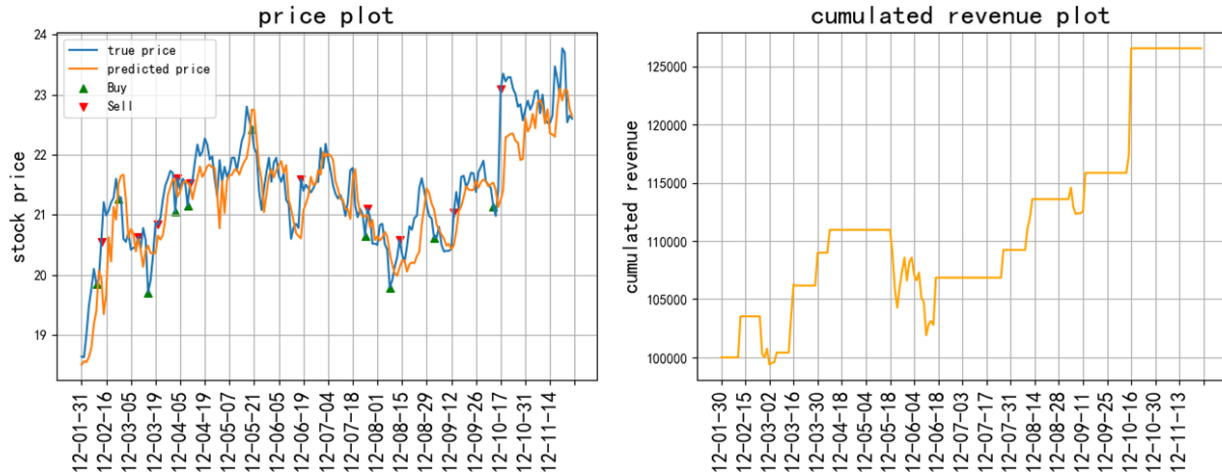


Figure 5. Buying, Selling and cumulated revenue of sz000651 strategy.

Figure 5 shows that the strategy made total revenue of \$26,527 over 200 days, equivalent to an annualized return of 34.19%. The Sharpe ratio of the strategy is 0.65 and the maximum drawdown is only 8.16%, indicating the excellence of this prediction model.

In order to ascertain the superiority of sufficient dimension reduction method, use the traditional method principal component dimension reduction (PCA) to replace SDR in this section with other process unchanged. By comparing the prediction results of PCA and SDR, there will be a more comprehensive understanding of the respective characteristics of these two dimension reduction methods in share price prediction. The result outputs are shown in Table.2.

Table 2. Comparison of PCA factors and SDR factors.

Learning Machine	11 SDR	11 SDR +3	11 PCA	11 PCA +3
Decision Tree Regressor	1.106	0.862	0.808	0.675
Linear Regressor	0.706	0.602	0.583	0.588
Random Forest Regressor	0.633	0.470	0.803	0.620
AdaBoost Regressor	0.548	0.498	0.716	0.508
GBDT Regressor	0.599	0.530	0.644	0.584
XGBoost Regressor	0.617	0.498	0.784	0.601
LightGBM Regressor	0.584	0.506	0.737	0.556
Stacking	0.725	0.590	0.769	0.654

The first two columns in Table 4 have already been shown in Table 2 and they are listed here for better comparison. It can be seen that the RMSEs of SDR factor group are consistently lower than that of PCA factor group in all models trained with four sets of predictors. The Random Forest Regressor reaches the best prediction with an RMSE of 0.470 in SDR factor group, which is lower than the best prediction of PCA factor group with an RMSE of 0.508.

In addition, it is also worth noting that the RMSEs of PCA group are lower than those of SDR group in simple learning machines linear regressor and the decision tree regressor. However, this situation seems to be totally opposite in the ensemble learning machine, where the RMSE of PCA group is higher than that of SDR group. This indicates that the simple learning machine can achieve better prediction results paired with principal component dimension reduction, while the ensemble learning machine prefers factors obtained from sufficient dimension reduction. Now that ensemble learning methods are more widely used due to their high accuracy, strong generalization ability and excellent theoretical properties, so applying the sufficient dimension reduction method in stock factor accords more with realistic application scenarios.

To explore the extensiveness of the above conclusions, choose another stock in the Chinese A-share market with different time periods to construct the model. Similarly, the stock factors of China Southern Airlines (sh600029) for seven hundred days from 25 February 2014 to 31 December 2016 are obtained from the DataYes! To conduct the above experiment, with results showing in Table.3.

Table 3. Prediction results of sz000651 and sh600029.

Learning Machine	sz000651		sh600029	
	11 SDR +3	11 PCA +3	11 SDR +3	11 PCA +3
Decision Tree Regressor	0.862	0.675	0.363	0.329
Linear Regressor	0.602	0.588	0.203	0.223
Random Forest Regressor	0.470	0.620	0.328	0.213
AdaBoost Regressor	0.498	0.508	0.182	0.227
GBDT Regressor	0.530	0.584	0.245	0.206
XGBoost Regressor	0.498	0.601	0.223	0.234
LightGBM Regressor	0.506	0.556	0.224	0.221
Stacking	0.590	0.654	0.470	0.224

The first two columns of Table.3 demonstrate the prediction results of the SDR factor and PCA factor for the Gree Electric Appliances stock under each model, and the third and fourth columns demonstrate the SDR factor and PCA factor for China Southern Airlines stock. The prediction performance of both stocks at different times can be illustrated:

- (1) The prediction results of SDR factors outperform PCA factors consistently;
- (2) The AdaBoost model with SDR factors reaches the lowest RMSE in predicting the share price of China Southern Airlines;
- (3) Regardless of the dimension reduction method, the prediction results of ensemble learning are generally better than simple machine learning models;
- (4) Simple machine learning models with PCA factors outperform those with full SDR factors, while ensemble learning models with SDR factors achieve better results than those with PCA factors.

The RMSEs of the two stocks under the same learning machine differ significantly due to the difference of average price and the volatility of fluctuation pattern. Next, conduct parameter tuning of AdaBoost model to obtain more accurate price predictions and run back test.

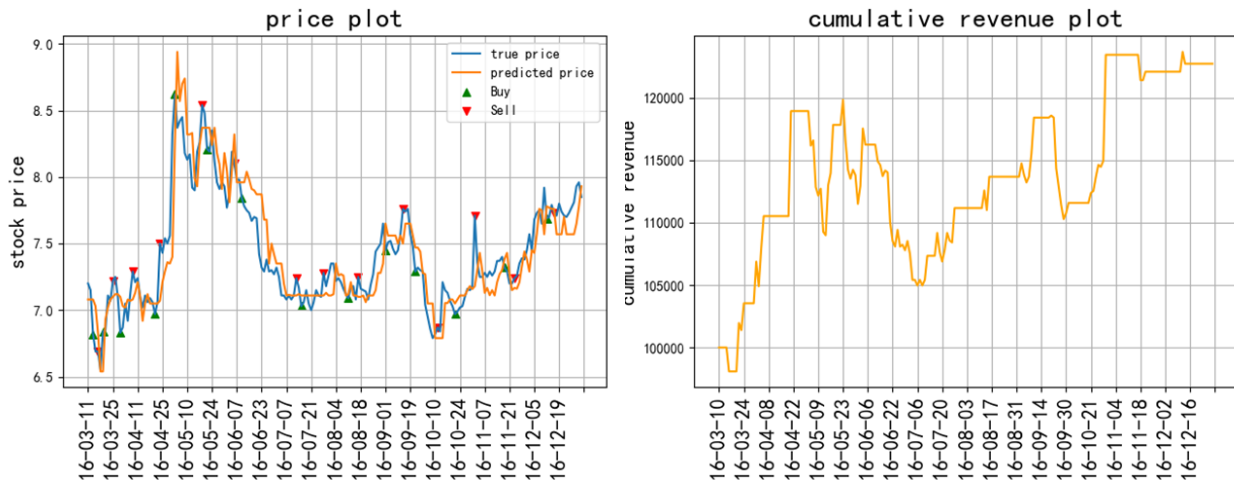


Figure 6. Buying, Selling and cumulated revenue of sh600029 strategy.

Figure.6 shows that the strategy made a total revenue of \$22,709 over 200 days, which is equivalent to an annualized return of 29.15%. The Sharpe ratio of the strategy is 0.47 and the maximum retracement is 12.4%, which makes the predicting model effective. It can be seen that the prediction results obtained by the ensemble learning model produced remarkable performance on simple strategies on both two stocks with SDR factors, indicating a strong adaptability to the Chinese A-share market.

4. Conclusions

This paper proposes an ensemble learning model based on sufficient dimension reduction to predict stock price, and illustrates its effectiveness and robustness by comparing the prediction results

with the principal component dimension reduction method and conducting back testing on two A-share stocks. This paper creatively uses the idea of sufficient dimension reduction in stock factor information extraction and stock price prediction, reducing the dimension of predictor variables while theoretically retaining the information of all factors, which solves the difficulty in factor significance measurement, imperfect utilization of factor information, and the factor decay problem. This paper experimentally illustrates that the prediction results of SDR factors is better than that of PCA factors and the combination of ensemble learning model and sufficient dimension reduction method can achieve better prediction accuracy. In addition, the quantitative trading strategy based on sufficient dimension reduction and ensemble learning achieves excellent results, all of which are able to meet an annualized return of more than 30% with the maximum retraction controlled within 15%.

This paper only explores the use of sliced inverse regression as sufficient dimension reduction method, however, there are many sufficient dimension reduction methods. Different methods hold different properties and spontaneously are suitable for different stock factor datasets. For future research directions, on the one hand, exploring the differences in prediction results obtained by different sufficient dimension reduction methods can be a key point of future research. On the other hand, sufficient dimension reduction is only applicable to continuous responding variable at most of the time, which makes the adjustment to the model becoming a problem demanding prompt solution when the responding variable is discrete (e.g., binary). Finally, when facing more realistic application scenarios and use more complex trading strategies, how to take the advantage of SDR information extraction to obtain greater profits and take less risk will also be an important research topic.

References

- [1] Chen Weizhong, Jin Yiping, Wang Yingluo. Research on dynamic pricing model and mechanism of securities assets [J]. Journal of Xi'an Jiaotong University, 1997, (10): 105-112.
- [2] Shuo Sun; Rundong Wang;Bo An. Reinforcement Learning for Quantitative Trading[J].ACM Transactions on Intelligent Systems and Technology,2023, Vol.14(3): 1-29.
- [3] Luo last night. An Empirical Study on Multifactor Stock Selection Model and Investment Effect Based on Industry Rotation Strategy [J]. Science and Technology Information, 2022, Vol. 20 (20): 148-152.
- [4] Li Shuchao. Quantitative investment effect appeared enhanced index fund excess return outstanding [N]. China Fund News, 2017.06.12.
- [5] Eugene F. Fama and Kenneth R. French. The Cross-Section of Expected Stock Returns [J].The Journal of Finance, 1992, Vol.47 (2): 427-465.
- [6] Eugene F. Fama; Kenneth R. French. A five-factor asset pricing model(Article)[J].Journal of Financial Economics,2015,Vol.116(1): 1-22.
- [7] Xie H-liang, Hu Di. Application of multi-factor quantitative models in investment portfolios-A comparative study based on LASSO and Elastic Net [J]. Statistics and Information Forum, 2017, Vol.32 (10): 36-42.
- [8] Fang Xin. Research on multi-factor stock selection strategy based on quality factors [D]. Shanghai University of Finance and Economics, 2021.
- [9] Ding Qi. Research on multi-factor quantitative investment strategy of stock based on principal component analysis [J]. Times Finance, 2020, (17): 74-76.
- [10] Zhao Ning. Short-term prediction of Vanke stock price based on ARIMA model [J]. Financier, 2023, (1): 34-36.
- [11] Li Gangcheng. Research on quantitative strategy of multi-factor stock selection based on improved LSTM model [D]. University of International Business and Economics, 2022.
- [12] Bo an; Shuo Sun; Rundong Wang. Deep Reinforcement Learning for Quantitative Trading: Challenges and Opportunities [J].IEEE Intelligent Systems, 2022, Vol.37 (2): 23-26.
- [13] J. D M. Sufficient Dimension Reduction: Methods and Applications with R [J]. Journal of the American Statistical Association, 2020, 115(530): 21-23.

- [14] J.W. Xu, Y.Y. Yang. Integrated learning methods: a review of research [J]. Journal of Yunnan University (Natural Science Edition), 2018, Vol. 40(6): 1082-1092.
- [15] Tailen Hsing and Raymond J. Carroll. An Asymptotic Theory for Sliced Inverse Regression [J]. The Annals of Statistics, 1992, Vol. 20(2): 1040-1061
- [16] Tingyou Zhou, Yuexiao Dong and Liping Zhu. TESTING THE LINEAR MEAN AND CONSTANT VARIANCE CONDITIONS IN SUFFICIENT DIMENSION REDUCTION [J]. Statistica Sinica, 2021, Vol. 31(4): 2179-2194
- [17] Leo Breiman [1]. Bagging Predictors [J]. Machine Learning, 1996, Vol. 24 (2): 123-140.