

Luque Raya, Ignacio Manuel; Luque Raya, Pablo

Article

Machine learning algorithms applied to the estimation of liquidity: The 10-year United States treasury bond

European Journal of Management and Business Economics (EJM&BE)

Provided in Cooperation with:

European Academy of Management and Business Economics (AEDEM), Vigo (Pontevedra)

Suggested Citation: Luque Raya, Ignacio Manuel; Luque Raya, Pablo (2024) : Machine learning algorithms applied to the estimation of liquidity: The 10-year United States treasury bond, European Journal of Management and Business Economics (EJM&BE), ISSN 2444-8451, Emerald, Leeds, Vol. 33, Iss. 3, pp. 341-365,
<https://doi.org/10.1108/EJMBE-06-2022-0176>

This Version is available at:

<https://hdl.handle.net/10419/325573>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



<https://creativecommons.org/licenses/by/4.0/legalcode>

Machine learning algorithms applied to the estimation of liquidity: the 10-year United States treasury bond

Estimation of
liquidity

341

Ignacio Manuel Luque Raya
Universidad de Granada, Granada, Spain, and
Pablo Luque Raya
Credit Suisse Group AG, Zurich, Switzerland

Received 1 August 2022

Revised 22 January 2023

8 March 2023

Accepted 12 March 2023

Abstract

Purpose – Having defined liquidity, the aim is to assess the predictive capacity of its representative variables, so that economic fluctuations may be better understood.

Design/methodology/approach – Conceptual variables that are representative of liquidity will be used to formulate the predictions. The results of various machine learning models will be compared, leading to some reflections on the predictive value of the liquidity variables, with a view to defining their selection.

Findings – The predictive capacity of the model was also found to vary depending on the source of the liquidity, in so far as the data on liquidity within the private sector contributed more than the data on public sector liquidity to the prediction of economic fluctuations. International liquidity was seen as a more diffuse concept, and the standardization of its definition could be the focus of future studies. A benchmarking process was also performed when applying the state-of-the-art machine learning models.

Originality/value – Better understanding of these variables might help us toward a deeper understanding of the operation of financial markets. Liquidity, one of the key financial market variables, is neither well-defined nor standardized in the existing literature, which calls for further study. Hence, the novelty of an applied study employing modern data science techniques can provide a fresh perspective on financial markets.

Keywords Data science, Finance, International markets, Machine learning liquidity, Treasury bond

Paper type Research paper

1. Introduction

The foundation of the present study is the concept of liquidity as a key financial indicator with which to predict the behavior of financial markets.

Liquidity is the flow of capital and credit within the global financial system. A concept that both the Bank for International Settlements and the Federal Reserve System apply, as well as many other financial institutions, as is reflected in the Fed Financial Stability Report. The concept of liquidity is approached in this study, so as to analyze financial stability, to anticipate systemic risk and particularly to analyze capital management among certain private sector investors.

The area of greatest economic significance in relation to liquidity is the central banking community, in which the term “Financial conditions” is also used. This concept is equivalent to the underlying idea behind liquidity. The capability to anticipate financial instability helps policymakers to make decisions on monetary policy, and it is likewise decisive for capital management among certain investors.



© Ignacio Manuel Luque Raya and Pablo Luque Raya. Published in *European Journal of Management and Business Economics*. Published by Emerald Publishing Limited. This article is published under the Creative Commons Attribution (CC BY 4.0) licence. Anyone may reproduce, distribute, translate and create derivative works of this article (for both commercial and noncommercial purposes), subject to full attribution to the original publication and authors. The full terms of this licence may be seen at <http://creativecommons.org/licences/by/4.0/legalcode>

European Journal of Management
and Business Economics
Vol. 33 No. 3, 2024
pp. 341-365
Emerald Publishing Limited
e-ISSN: 2444-8494
p-ISSN: 2444-8451
DOI 10.1108/EJMBE-06-2022-0176

Alessi and Detken (2011) used liquidity as a predictive indicator to study asset prices during boom/bust cycles that can have serious economic consequences. Chen *et al.* (2012) recognized the importance of decision-making and explored the significance of liquidity for the global economy and for policymakers through multiple sets of liquidity indicators.

The researchers attempted to model the global financial cycle in terms of the interaction of monetary policy and financial stability in markets and they emphasized its relevance to central bank interest rate hikes and quantitative easing policies (Miranda-Agrippino and Rey, 2020, May). They also studied the importance of capital inflows and outflows (international liquidity) to manage the implications of the global financial cycle in emerging markets (Jeanne and Sandri, 2020). Bernanke *et al.* (2011) likewise investigated the effects of capital inflows into the United States on U.S. longer-term interest rates.

The objective that is pursued here is to offer predictions of safe asset prices, through the application of data science techniques, in particular machine learning, identifying the models that yield the most promising results. To do so, the 10-year US treasury bond, considered to be the most representative variable within a typical portfolio of safe assets is used. These predictions are advanced using certain proxy variables, which are considered representative of the concept of liquidity. Various machine learning models are compared with stationary and nonstationary variables.

It is not only that these predictions are of intrinsic value, as they may also serve either to support or to refute the notion that liquidity fluctuations are in some way responsible for the fluctuations of other types of assets, specifically the 10-year US treasury bond and, by extension, economic fluctuations. This can help us to benchmark when working on prediction exercises with liquidity variables.

The use of machine learning techniques to guide monetary policymaking is a novelty with growing interest not only from the perspective of central banks themselves but also from the perspective of academia and, to a lesser extent, independent investors (Guerra *et al.*, 2022).

Despite the widespread use of machine learning, Guerra *et al.* (2022) reiterated that the combination of such factors as risk, safe assets, liquidity and the field of artificial intelligence have rarely been studied together, and in even fewer studies models have been used to forecast economic flows through liquidity variables, Galindo and Tamayo (2000) and Abellán and Castellano (2017) have demonstrated the suitability of machine learning algorithms for such tasks.

Authors such as Hellwig (2021) defend the improvement in predictions of complex financial concepts such as liquidity and the different dimensions by which we model it, that machine learning methods (*i.e. random forest or gradient boosted tree*) attain compared to other traditional econometric approaches. These improvements in predictions are generally obtained due to the limits present in the machine learning models when fitting the data which allow them to explore relationships of variables with a lower risk of overfitting and the increase in precision that ensemble models tend to achieve by averaging the predictions of other models.

One practical reason for using the 10-year US bond is its global importance as a reference price with which many different assets are linked and valued. According to Goda *et al.* (2013), long-term Treasury yields are used as a proxy of safe assets and the same authors also referred to the strong linkages between the Treasury yield and non-Treasury bond yields.

In summary, the state of the art of the liquidity concept is reviewed in Section 2. The different machine learning techniques and the theory behind the algorithms are treated in the same way. Subsequently and arising from that review, a series of research questions are proposed. In the fourth section, the methodological aspects are explained. In the data analysis section, the indicators that are common to the different algorithms are presented to

facilitate their analysis, discussion and comparison. Finally, the conclusions that respond to the research questions are set out, indicating the implications for management and for investors and institutional decision-making.

2. State of the art

2.1 *Liquidity and its different classes*

The concept of liquidity has been broadly investigated and discussed, especially since the global financial crisis of 2008. In particular, it has been and continues to be the business and the concern of the BIS and the IMF. Both financial institutions gather extremely useful macroeconomic data from reliable sources through surveys administered to central and commercial banks. The BIS analyzes liquidity from the point of view of financial stability, in order to minimize systemic risk and vulnerabilities. Its methodology is centered on two basic variables: banking assets and currency-based credits to the nonfinancial private sector.

Borio *et al.* (2011) highlighted the ambiguity of the above definition and laid the foundations of both key concepts to arrive at a standardized definition: separate types of liquidity, *i.e.* public or private, and their application within a context of financial stability. From a macro perspective, they presented the empirical characteristics of the financial cycle and the implications for monetary policy. In effect, significant movements of liquidity are associated with systemic banking crises, hence his proposal for the implementation of mechanisms to anticipate these events, lessening financial distress.

Howell (2020) presented one of the most up-to-date analyses of liquidity. The main ideas within this field are the classification of liquidity according to the type of source, the cyclic nature of liquidity, and its implications for financial stability. Its principal contribution emphasized the shadow banks (institutions not subject to banking regulation) and how they affect the whole system, especially through collateralized operations and safe assets. Indicators that seek to capture these operations will be taken into consideration in the analysis that is completed in this paper.

To Bruno and Shin (2015a, b), the modeling of global credit flows to anticipate financial distress is the principal example for any study of international liquidity. The works upon which predictive models are constructed on the basis of banking capital and currency flows are used to study the transmission of liquidity between different countries.

On the other hand, Stanley Druckenmiller (Barrow, 2017) based his investment strategy on the analysis of liquidity. “Earnings don’t move the overall market; it’s the Federal Reserve Board . . . focus on the central banks and focus on the movement of liquidity . . . Most people in the market are looking for earnings and conventional measures. It’s liquidity that moves markets”.

Liquidity may be defined as the total quantity of capital and credit existing within the system for use in the real economy (of products and services) and in the financial markets (assets). It is a gross financing concept that represents the overall balance of entities supplying money and credit to the system.

With regard to its source, liquidity may be classified as follows:

- (1) Private liquidity or within the private sector (endogenous according to the literature) covers both financial banking and nonbanking (shadow banks, institutions, large investors, *etc.*) sectors, including data on family credit, growth rate in the volume of financial savings/private credit, and interannual change in consultations of personal credit and small firms.
- (2) Public (or exogenous) liquidity is associated with governmental institutions relating to the source of official liquidity or the set of tools that the central bank can use, principally reference interest rates and operations on the open market (asset purchase

programs), the monetary base (core of the passive monetary policy of the central issuing bank), money in circulation and credits assigned to commercial banks, among others.

- (3) International liquidity or all financial resources available to the monetary authorities of each country that will be used to finance the deficits in their international balance of payments when all the other sources supplying foreign funds are insufficient to ensure a balance in international payments.

There is support in the literature for these dimensions of liquidity, particularly [Howell \(2020\)](#). In it, liquidity was broken down into three sub-components, then explanatory variables were leveraged to explore each one and finally machine learning modeling was applied to the data. The same methodology was also applied, although in a different way, in [Hatzius et al. \(2010\)](#), where the exploratory variables were aggregated and a financial conditions index was constructed.

Liquidity constitutes a time series, *i.e.* a set of observations recorded over a particular time span ([Brockwell and Davis, 2009](#)). What is now proposed is the prediction of multivariate time series, such as metric variables, that measure the different predicted versions within which we will itemize liquidity based on the ideas discussed in the above-mentioned literature.

The choice of the explanatory variables is endorsed in the literature, and the approach toward the measurement of liquidity, and the division between public and private sector entities is based on the work of [Landau \(2011\)](#). He proposed the separation of public and private factors when analyzing liquidity and the inclusion of both price and balance sheet quantities. [Caruana \(2013\)](#) also raised the inclusion of stock variables (*i.e.* amount of debt outstanding) and flow variables (*i.e.* bank credit growth) in the analysis of liquidity. [Chung et al. \(2014\)](#) examined the relationship between financial conditions and the money aggregates. [Shin and Shin \(2011\)](#) explored the link between monetary aggregates and the financial cycle.

[Lane and McQuade \(2014\)](#) suggested the inclusion of an international liquidity component, so as to shed light on financial stability within the domestic financial system. Three determining factors related to liquidity conditions were described in the work of [Eickmeier et al. \(2014\)](#): global monetary policy, global credit supply and global credit demand. Finally, [Cesa-Bianchi et al. \(2015\)](#) utilized bank-to-bank cross-border credit to examine international liquidity.

Given the above definitions of liquidity, we searched for variables representing different forms of capital and credit. We then proposed monetary aggregates as variables for capital and types of loans as variables for credit. The reference interest rate of monetary policy was used as an explanatory variable, because it is at the most fundamental level the price of money, it influences liquidity, and the economy, and it is under government control.

2.2 Machine learning

Machine learning models perform iterative processes on a dataset (divided into a training and a validation or test dataset) that refers to a specific context (in this case, liquidity). On that basis, predictions of the future values of the dependent variable are advanced (in this case, the 10-year US treasury bond), which are tested and validated with both the training and the validation datasets. The results are compared with a reference or benchmarked model.

2.2.1 Models with nonstationary variables. Nonstationary variables follow temporal trends, as they show no constant average or variance over time. In general, the models with nonstationary variables usually present worse results than the models with stationary variables. Overfitting affects these models more than the models with stationary variables, despite applying methods for its reduction that are recommended in the literature.

Bayesian Ridge Model uses probability distributors rather than point estimates. The output Y is extracted from a probability distribution, instead of its estimation as a unique value. In this way, good functioning is guaranteed even with insufficient or poorly distributed data.

2.2.2 Models with stationary variables. It is far easier to model a stationary than a nonstationary series. Having transformed nonstationary into stationary variables, the differences can be compared with the earlier models. That transformation is applied to the nonstationary variables that are identified with the Augmented Dickey–Fuller tests, operating in the following manner.

If $\mathbf{z}$ is a set of differentiated predictions and $\mathbf{Y}$ is the value of the original dependent variable, then

$$Y_{t+1} = Y_t + z^{t+1}$$

$$Y_{t+2} = Y_{t+1} + z^{t+2} = Y_t + z^{t+1} + z^{t+2}$$

The *Orthogonal Matching Pursuit* (OMP) algorithm, as with other “greedy algorithms”, constructs a sequential solution $X_0, X_1, X_2, \dots, X_k$. In this way, it includes an atom that is most closely correlated with the actual residual at each step and, unlike Matching Pursuit, the residual is calculated once again after each iteration using an orthogonal projection within the space of the previously chosen elements. In essence, the algorithm processes all n possibilities to identify a Column, A , that shows the highest correlation with the observations of y in the first step (hence the term “matching”), *i.e.* the best fit of Ax to b . Subsequently, it identifies Column A in each iteration that shows the highest correlation with the actual residual. It therefore seeks the atom with the best fit of Ax to b , on the basis of those selected earlier. In each iteration, the estimation of the vector sign is updated through the most highly correlated column with Column A (Khosravy *et al.*, 2020). The solution at each step is selected in such a way that the new residual is “orthogonal” to all the atoms selected in A .

The *CatBoost Regressor algorithm* is based on the theory behind other algorithms such as Decision Trees and Gradient Boosting. The principal concept of “boosting” involves sequential combinations of multiple models that perform slightly better than random chance. The algorithm is used to create a solid, predictive and competitive model, by applying a “greedy” search (a mathematical process that tests simple solutions to complex problems, through the choice of the subsequent step that provides the most obvious benefit).

In the same way as “gradient boosting” adjusts the decision trees in a sequential manner, the adjusted trees will learn from the errors of the earlier trees and the errors will therefore be minimized. The process continues until the selected loss function can no longer be minimized.

In the growth process of the decision trees, the algorithm produces “unconscious” trees, which means that the trees grow, under the rule that all the nodes at the same level test the same predictor under the same condition. This “unconscious” tree procedure permits simple adjustments and improves computational efficiency, while the structure of the tree operates as a means of regularization to identify an optimal solution and to avoid overfitting (Thiesen, 2020).

The *AdaBoost Regressor Model* has the objective of fitting a sequence of “weak learners” (models that are slightly better than a random estimation) to versions of the data that are repeatedly modified. The predictions of all these “weak learners” are combined through weighted majority voting, *i.e.* a sum in which equal account is taken of the weights, in order to estimate the final predictions.

The data modifications for each “boosting” iteration consist of applying weights $\mathbf{w}_1, \mathbf{w}_2, \dots, \mathbf{w}_n$ to each training sample. Initially, all these weights are established through $\mathbf{w}_i = 1/n$, in such a way that a “weak learner” is trained with the initial data in the first step of the

process. The weights of the sample are individually modified for each iteration that is successively performed and the algorithm is once again applied to the newly modified data.

According to this method, the weights attached to the observations of the training sample that are incorrectly predicted are increased, whereas those that are correctly predicted are, on the contrary, decreased. In doing so, far greater influence is attached to those observations that the model can only predict with difficulty as the iterations continue. Each subsequent “weak learner” is therefore obliged to center on the observations that other “weak learners” had mistakenly predicted earlier on (Wang, 2012).

The *Extreme Gradient Boosting Model* (XGBoost) provides computational rapidity and performance levels that are difficult to equal. This algorithm functions in the same way as other models that use the ensemble methods, in which new models are successively generated from a training set to correct the errors of the earlier models, in a similar way to the above-mentioned AdaBoost algorithm.

The concept of Gradient Boosting entails the design of new models that predict the residuals or errors of earlier models, which are then added together to arrive at a final prediction. It is referred to as Gradient Boosting, because it uses a reduced gradient algorithm to minimize the loss function when adding new models (Brownlee, 2016a).

Extremely Randomized Trees Model (ET) uses an ensemble decision tree method that usually yields better results than those based on simple decision trees.

The Extremely Randomized Trees or Extra Trees (ET) algorithm generates a large number of unpruned decision trees (without removing small-sized branches), based on the training dataset data, fitting each decision tree to the complete training dataset.

In summary, the principal differences that this algorithm has over other decision tree ensembles are as follows: the use of the whole training dataset (instead of a bootstrap replica) to start the growth of the trees as mentioned earlier; the other difference is that ET divides the nodes, randomly selecting the cutoff points (Geurts *et al.*, 2006). At each cutoff point, the algorithm activates a random selection of the different features.

The predictions advanced in a regression problem are prepared through the average of all the decision trees, while a majority voting method among the different decision trees is used for the classification problems (Geurts *et al.*, 2006). In the case of regression, these averages are used to improve the predictions and to check for overfitting.

The choice of this algorithm for the completion of the model was due to the problems with overfitting observed when using the other models, because the random selection of the cutoff points meant fewer correlations between the decision trees (although this random selection increased the variance of the algorithm, increasing the number of trees used in the ensemble can counteract that effect) and the level of overfitting may therefore be reduced in comparison with the levels of other models.

The *Random Forest* (RF) algorithm is a decision tree ensemble method similar to ET. Both are very similar algorithms composed of a large number of decision trees that will influence the final prediction.

The main differences are, on the one hand, that subsamples generated with the bootstrapping method are used in Random Forest, *i.e.* a resampling technique that generates datasets by continuously repeating the sampling of the available data (James *et al.*, 2013). In contrast, the overall sample is used in the et algorithm. On the other hand, the cutoff points were established in the most optimum form in Random Forest, unlike ET in which a higher randomness component was added to its decision-making.

The above-mentioned greedy algorithms have not been widely discussed in the literature on similar liquidity problems, so we wish to delve further into their performance in this area of study. On the other hand, there is ample support for the suitability of both decision tree algorithms and ensemble methods (Galindo and Tamayo, 2000; Abellán and Castellano, 2017; Sahin, 2020), which perform better than conventional approaches and other machine learning

algorithms applied to liquidity-related classification and prediction problems (Guerra *et al.*, 2022).

2.2.3 Models used with both types of variables. The Voting Model consists of combining different sorts of automatic learning algorithms and generating final predictions, through an estimator consensus method (average probabilities when processing a regression problem). The results are intended to yield separate improvements to the predictions of the original method. The aim is therefore to improve the predictions of certain individual models through their combination and by averaging their values.

3. Research questions

The importance is evident in both the concept of liquidity and the differentiation of types of liquidity. A variety of indicators are used to measure the different types. Besides, automatic learning offers a broad range of possibilities to make predictions, for which reason the general objective is to predict liquidity through different machine learning algorithms, comparing the results that are provided, in order to identify the best models. In this case, the price of the 10-year US treasury bond is taken as a reference and the following research questions are proposed:

- RQ1.* Will the estimation of the price of the 10-year US treasury bond with machine learning models improve upon the estimations of traditional models?
- RQ2.* Will the best predictions of the 10-year US treasury bond be dependent upon predictions with either stationary or nonstationary variables?
- RQ3.* Which machine learning algorithms will yield better estimations of the 10-year US treasury bond?
- RQ4.* Which voting models will improve upon the estimations of other models?
- RQ5.* Which models will present more problems of overlearning in their predictions?
- RQ6.* Which variables will determine more than any others the price of the 10-year US treasury bond? Will the private, the public or the international liquidity variables be the most decisive?

4. Methodological aspects

4.1 Variables and data sources

The dependent variable is the price of the 10-year US treasury bond. It is the most widely tracked debt indicator and price instrument in the field of financing, being used as a reference benchmark to calculate different values. It is usually perceived as the “safe” asset by antonomasia, attracting large quantities of liquidity especially during times of crisis and uncertainty, in what are referred to as safe havens (Zucchi, 2021).

The contents described in Tables 1–3 were used as the independent variables of private, public and international liquidity. In Table 1, the variables used to describe money and credit circulating between private agents are described. In Table 2, the variables used to describe money and credit circulating between public agents are described and likewise, in Table 3, the variables used to describe money and credit circulating between international agents are described.

In Table 4, the variables used to describe different financial markets (bond market, foreign exchange market, developed economies stock market and emerging economies stock markets) are described. Quantity and price (market-dependent prices or rates) indicators are used.

Table 1.
List of variables
relating to private
liquidity

Variable	Notation in graphs	Definition
Total consumer credit & Change in the total consumer credit	Total_ConsumerCredit & Change_Total_ConsumerCredit	Total consumer credit in property and securitized, nonstationary flows
Total monetary base (M_0)	Total_Monetary_Base	Total quantity of money (in this case, the US Dollar) that is in general circulation in the hands of the public or in the form of commercial banking deposits maintained in the US central banking reserve (M_0)
M_1 for the USA	$M1_USA$	Monetary aggregate resulting from the sum of i) money in circulation in USD; ii) private sector deposits at commercial banks and iii) other liquid and savings deposits
M_2 for the USA	$M2_USA$	Monetary aggregate resulting from the sum of i) M_1 aggregate; ii) savings deposits (<100K) in commercial banks and iii) liquid funds from the money market
M_3 for the USA	$M3_USA$	Monetary aggregate resulting from the sum of i) M_2 aggregate and ii) savings deposits (>100K) in commercial banks
Private credit for the USA	Private_Credit_USA	Total credit (banking and nonbanking) to the nonfinancial private sector for the USA
M_1 for the EU	$M1_EU$	Monetary aggregate resulting from the sum of i) money in circulation in Euros; ii) private sector deposits held in commercial banks and ii) other liquid and savings deposits
M_2 for the EU	$M2_EU$	Monetary aggregate resulting from the sum of i) M_1 aggregate; ii) savings deposits (<100K) at commercial banks and iii) liquid funds from the money market
M_3 for the EU	$M3_EU$	Monetary aggregate resulting from the sum of i) M_2 aggregate and ii) savings deposits (>100K) at commercial banks
Private credit for the EU	Private_Credit_EU	Total credit (banking and nonbanking) to the private sector for the Eurozone
Note(s): In all cases, on a monthly basis. Data preprocessing converted the quarterly variables into monthly figures Source(s): Table by authors		

Table 2.
List of variables
relating to public
liquidity

Variable	Notation in graphs	Definition
Bank credit of US commercial banks	Bank_Credit_Comercial_Banks	Credit in the (asset) balance of US commercial banks
Currency in circulation in the USA	Currency_in_Circulation_USA	Physical amount of money (Dollars) in circulation within the monetary system of the USA
Deposits in US commercial banks	Deposits_Commercial_Banks_USA	Consumer deposits held at US commercial banks
Federal fund rate	FEDFUNDS_%	The federal fund rate is the reference interest rate for the implementation of monetary policy
Reference rate of the ECB (European Central Bank)	Reference_Rate_BCE	The rate that commercial banks and financial institutions receive for placing short-term investments with the ECB
Source(s): Table by authors		

4.2 Data preprocessing

The data were limited to between January 1, 2000 and December 1, 2020, although 2019 and 2020 were reserved to make predictions with “unknown” future dates.

Python was used for the processing work. In order to palliate the effect of outliers 2.5% of the values on each side of the tail of the distribution were removed.

Interpolation was used to standardize the quarterly data on a single monthly timescale. To do so, the Python programming language was used to establish the date (frequency in quarters in this case) as the data frame index. We then applied the resample function of the Pandas package to change the frequency to monthly, and we filled in the NA values (or missing data) obtained through the interpolation function of the same package, using a linear method. Finally, the data was converted into a column instead of an index format, so that the data could be merged with the other Pandas DataFrames.

The daily frequency data in the case of the “European Central Bank reference rate” variable were added in months using the average interest rate of that variable throughout the month.

It is necessary to analyze the specific characteristics of the series referring to autocorrelation: seasonality and stationarity (Chhugani, 2020). A moving average was applied, in order to analyze these tendencies with a temporal window of 12 months, within which the rolling average was calculated, and exponential softening was also applied, the results of which are shown in Figure 1.

The US bond, as may be observed, followed a descending trend. The Augmented Dickey–Fuller test was applied to each variable to test for stationarity. The transformation of the first difference was applied to the nonstationary variables (the preceding observation was

Variable	Notation in graphs	Definition
Global liquidity	Global_Liquidity(%)	BIS Indicators of global liquidity. Index composed of two indicators: a) international banking assets, essentially bank loans throughout the work, both from other banks and firms and consumers and b) credit to firms and consumers by denomination of currency (USD, EUR, JPY)
DXY index	DXY_Index	DXY (US Dollar) index v. a basket of currencies (EURO, YEN, POUND STERLING, etc.). The DXY is a weighted geometric measure

Note(s): Monthly frequency except for the global liquidity indicators that were presented quarterly and were preprocessed for the monthly figures

Source(s): Table by authors

Table 3.
List of variables
relating to
international liquidity

Variable	Notation in graphs	Definition
Price of the 10-year German bond	Price_Bond_German_10Y	10-year German government bond
Exchange rate Euro against Dollar	Variation_EU_VS_USD	Rate of variation of exchange rate with regard to the US Dollar
Rate of variation of S&P 500	Variation_SP_500	Rate of variation of the US S&P 500 share index
Closing price of the VEIEX	Closing_Price_VEIEX	Vanguard Emerging Markets Stock Index Fund Investor Shares (VEIEX), monthly prices of the index of emerging economy shares

Source(s): Table by authors

Table 4.
List of variables
relating to financial
markets used as
predictors



Figure 1.
Softened exponential of
the dependent variable

Source(s): Figure by authors

subtracted from the current observation over time) converting all the variables that were not stationary into stationary variables, as shown in [Figure 2](#). Having obtained the predictions of the model, this transformation was reversed for proper interpretation of the results.

4.3 Modeling

- (1) *Data selection.* The observations from 2019 to 2020 were set aside to make predictions. Although the extraordinary circumstances of 2020, due to COVID-19, complicated the predictions, the training dataset extended from January 2000 until December 2011, while the validation or test dataset extended from January 2012 until December 2018.

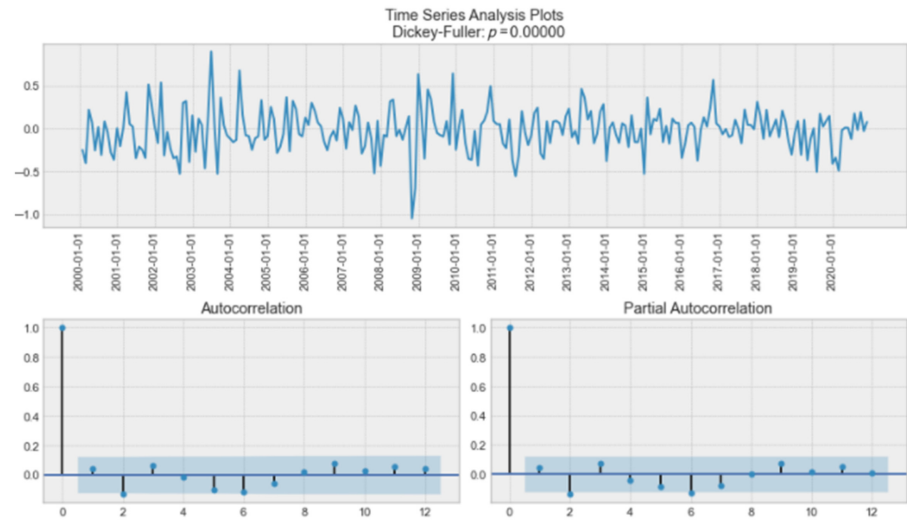


Figure 2.
Dickey-Fuller test,
autocorrelation and
partial autocorrelation
of the dependent
variable with
transformation of the
first differences

Source(s): Figure by authors

- (2) *Cross-validation* strategy. The time series cross-validator, a variation of the k-fold, was used. It returns the first **k**-folds within the **k** subsample as a training set and the **k-fold + 1** as the validation dataset. The successive training datasets are supersets of the previous datasets. In this way, problems of data leakage and overfitting are avoided. Data leakage occurs when the data used to train an algorithm hold the information that the algorithm is attempting to predict (Gutierrez, 2014), in other words, when the model is trained with data that it is meant to predict and that should not be available to it. Overfitting occurs when the model memorizes the noise or the random fluctuations in the training data, which implies a negative impact when generalizing (Brownlee, 2016b). After different tests, it was concluded that the ideal number of folds was **k** = 10.
- (3) Normalization. The most widely used option in the literature for the normalization of temporal series, minmax, was applied.
- (4) Multicollinearity. There is no reason for high correlations to affect the model in a negative way. Although some indications pointed to the liquidity-related variables as a factor in the price variations of the assets that were considered to be safe, the predictions were based on correlations rather than causality. Underlying causes (liquidity) within the financial markets for the variation in the price of safe assets can be intuitively guessed.

In general, a threshold is established to avoid multicollinearity, according to which if the correlation between two variables is higher or equal to a particular value, one of them is removed to minimize problems.

- (1) Evaluation of the models. The goodness of the model predictions is usually evaluated through a comparison with a series of base models (linear regression), rather than with metrics that are dependent on such scales as the Root Mean Square Error (RMSE) and the Measure of Absolute Percentage Error (MAPE). It could therefore be verified that the most complex models contributed greater value to the predictions than the simpler and more easily interpretable metrics such as MAPE, mentioned above, and the Mean Absolute Error (MAE).

5. Data analysis

The results were expressed with standard measurement metrics, in order to compare the different algorithms, following the same evaluation process for each algorithm. The individual performance of each model is graphically represented for a clearer understanding, commenting on the differences between the selected metrics, the adjustments of the models and, in certain cases, the distribution of errors to study the aforementioned adjustments. The set of results of the models with both stationary and nonstationary variables are summarized and compared in Table 6.

Two linear regression models were prepared with the previously explained predictors. Their hyperparameters were adjusted through a Random Grid Search, which was the benchmark with which other more advanced models were compared. One with nonstationary variables and another with variables converted into stationary values. Metrics were obtained from the stationary model that improved upon those of the nonstationary model, as can be observed in Table 5.

A variety of machine learning models with different regression algorithms were applied. The R^2 statistic, considered of little or no use in the literature (Dunn, 2021) in the context of prediction-centered automated learning, was not applied. Instead, metrics that were not

dependent on the scale were principally used, in this case RMSE. The models whose metrics improved upon the pre-established base models may be highlighted.

5.1 Models with nonstationary variables

The *Bayesian Ridge Model* yielded the best results when using the initial variables: MAE = 0.2751, RMSE = 0.3338 and MAPE = 0.0783. It was the only one of its type that improved upon the base linear regression model that had been proposed.

This model also improved upon the MSE/RMSE obtained with a persistence model (MSE = 0.167, RMSE = 0.4086) that predicted the price of the 10-year US treasury bond at $t+1$ through the value of t (Figure 3). This particular model therefore contributed value with respect to the two naïve models, thereby ensuring that work with temporal series was not merely a random walk.

The results of the *Voting Model* were very similar. They yielded slightly lower results: MAE = 0.2781, RMSE = 0.3386 and MAPE = 0.0782 (Figure 4). Nonetheless, the overfitting of these models is easily appreciated.

5.2 Models with stationary variables

The results of the base reference model were improved through the use of various algorithms. The metrics of the stationary models are shown in detail in Table 6.

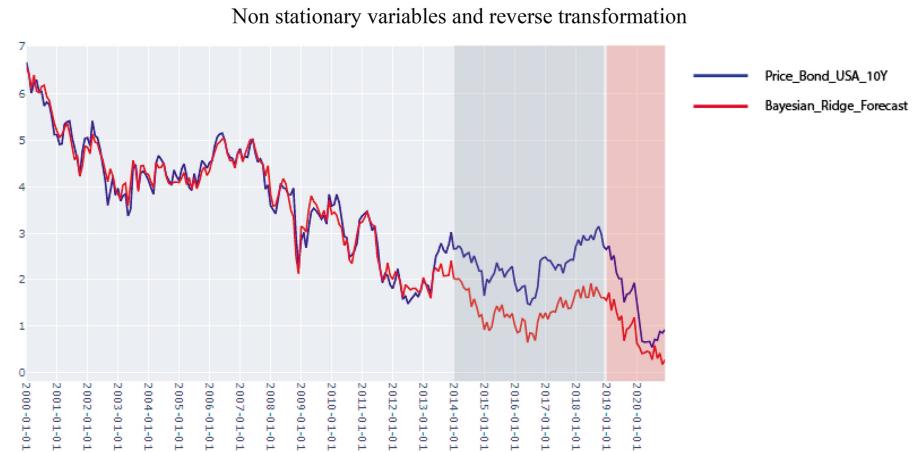
The *OMP Model* is among those that yielded the best results, improving on the base Linear Regression model and most of the other models (Figure 5). After the changes

Table 5.
Characteristics of the
linear regression
models

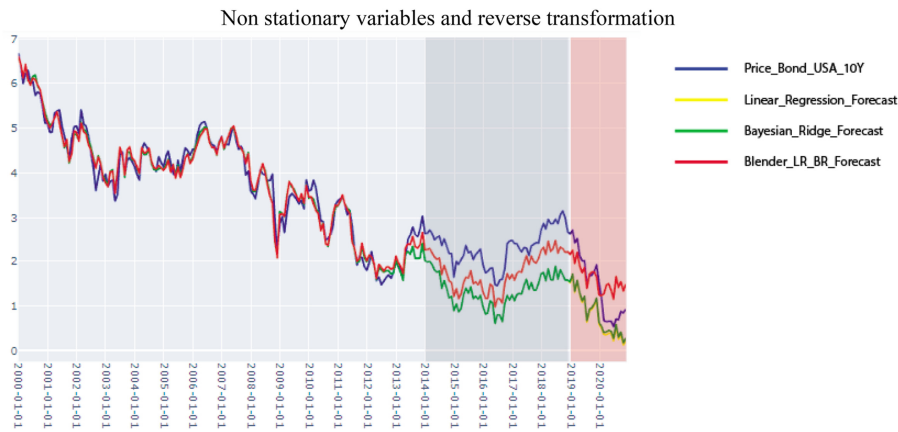
Variables	MAE		RMSE		MAPE	
	Mean	SD	Mean	SD	Mean	SD
Nonstationary	0.3451	0.1482	0.4018	0.1615	0.0866	0.0363
Stationary	0.3116	0.2323	0.3816	0.3816	5.3768	9.4791

Source(s): Table by authors

Figure 3.
Bayesian Ridge model
applying
nonstationary and
reverse transformation
variables



Source(s): Figure by authors



Source(s): Figure by authors

Figure 4.
Comparison of models
applying
nonstationary
variables and reverse
transformation

Model	RMSE	MAE	MAPE	Variables	Overfitting
Second Voting (OMP/RF/ET)	0.1715	0.1312	1.4816	Stationary	No
First Voting (OMP/CatBoost/AdaBoost/ XGBoost)	0.1759	0.1360	1.3781	Stationary	Yes
ET	0.1782	0.1358	1.3479	Stationary	No
OMP	0.1786	0.1418	1.558	Stationary	No
Random Forest	0.1858	0.1464	1.5869	Stationary	No
CatBoost	0.1860	0.1456	1.5149	Stationary	Yes
Adaboost	0.1889	0.1499	1.6574	Stationary	Yes
XGBoost	0.1899	0.1509	1.4860	Stationary	Yes
Voting (BR/RL)	0.3336	0.2781	0.0782	Nonstationary	Yes
Bayesian Ridge	0.3338	0.2751	0.0783	Nonstationary	Yes
Stationary Base Linear Regression	0.3816	0.3116	5.3768	Stationary	No
Non-Stationary Base Linear Regression	0.4018	0.3451	0.0866	Nonstationary	No

Source(s): Table by authors

Table 6.
Metrics obtained by
each model ordered
from lowest to
highest RMSE

introduced to reduce overfitting, the most relevant predictors for the dependent variable were the interbank interest rate of the US Federal Reserve (FedFunds), the reference interest rate of the BCE, the price of the German government bond and, finally, the variation of the Euro with respect to the Dollar.

The *CatBoost Model* performance was not poor overall, although it was seriously affected by overfitting, as often happens with “greedy algorithm” based models. Despite having taken this factor into account and having taken the necessary measures, the R^2 values of the training and the test dataset were 0.934 and 0.501, respectively, which points to a large gap; the sign of an overfitted model (Figures 6 and 7).

The most important features of the CatBoost model were as follows: M_1 monetary aggregate of both the European Union and the US, the price of the German government bond, the closing price of VEIEX, the variation of the Euro with respect to the Dollar, the total of the monetary base, the change of consumer credit, the S&P 500 price and private European credit.

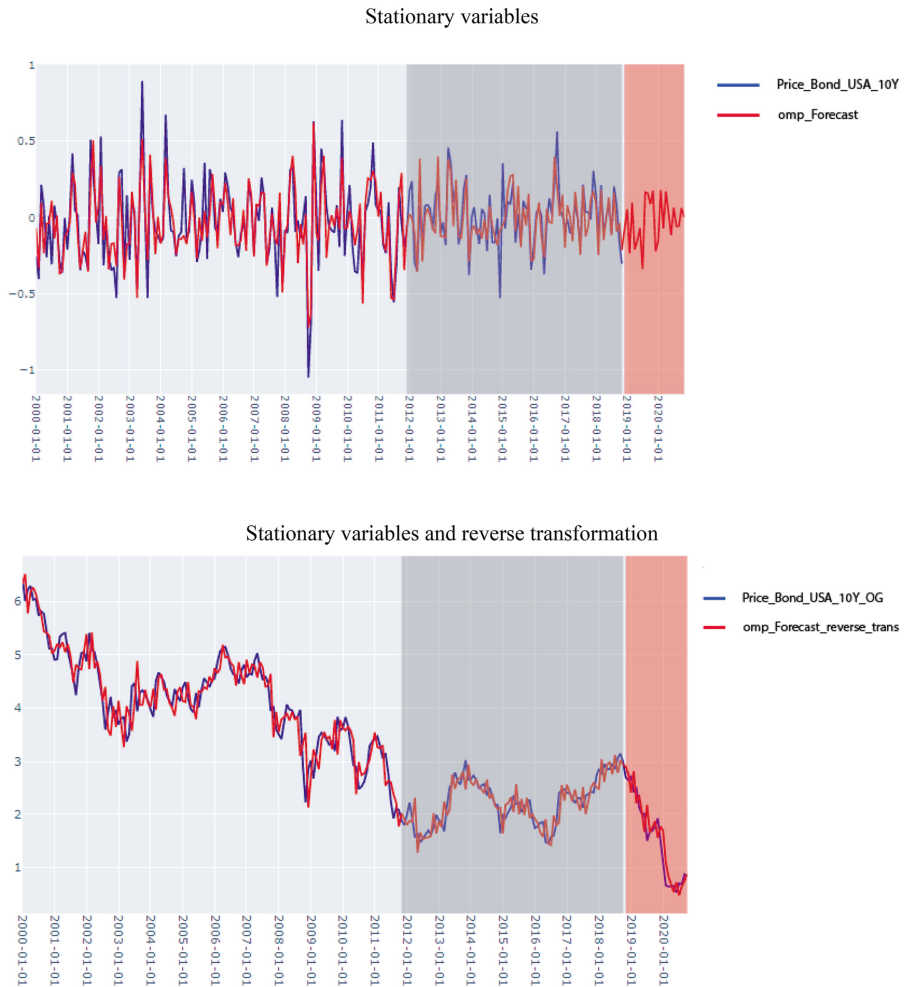


Figure 5.
OMP model

Source(s): Figure by authors

The most influential independent variables of *AdaBoost Model* were as follows: the price of the German government bond (which is the norm for the majority of models), followed by the M_1 monetary aggregate of the European Union, the variable year, the variation of the Euro with regard to the Dollar, the closing price of the VEIEX, private US and European credit and, finally, the total monetary base.

When we look at the distribution of the residuals of this model, a certain degree of overfitting can be seen. The R^2 squared of the training dataset was 0.862 while that of the test dataset was 0.468 (Figure 8).

The metrics obtained with the *XGBoost* model were as follows: MAE = 0.1509, RMSE = 0.1899 and MAPE = 1.4860 (Figure 9).

The results of the Feature importance model measured the degree to which one variable influenced the results based on the predictions of a particular model (Oh, 2019). We can see



Source(s): Figure by authors

Figure 6.
CatBoost model

that the most relevant variables were the credit of all commercial banks, the price of the 10-year German bond and the closing price of the VEIEX. Followed to a lesser degree, but even so with a notable influence, the variables M1 EU and Total Monetary Base.

The overfitting was high in this model, obtaining an R^2 in the training and in the test datasets of 0.869 and 0.485, respectively.

The *ET Model* presented very good results, better than those of the base Linear Regression model and the majority of the other models that were prepared (Figure 10).

With regard to overfitting, the principal reason for the selection of this algorithm was understood to be somewhat less than in other models that were prepared. The R^2 of the training dataset was 0.613, while the test dataset had an R^2 of 0.543, which is an acceptable difference.

The predictions of this model were principally influenced by the price of the 10-year German government bond and private credit, both in the EU and in the USA.

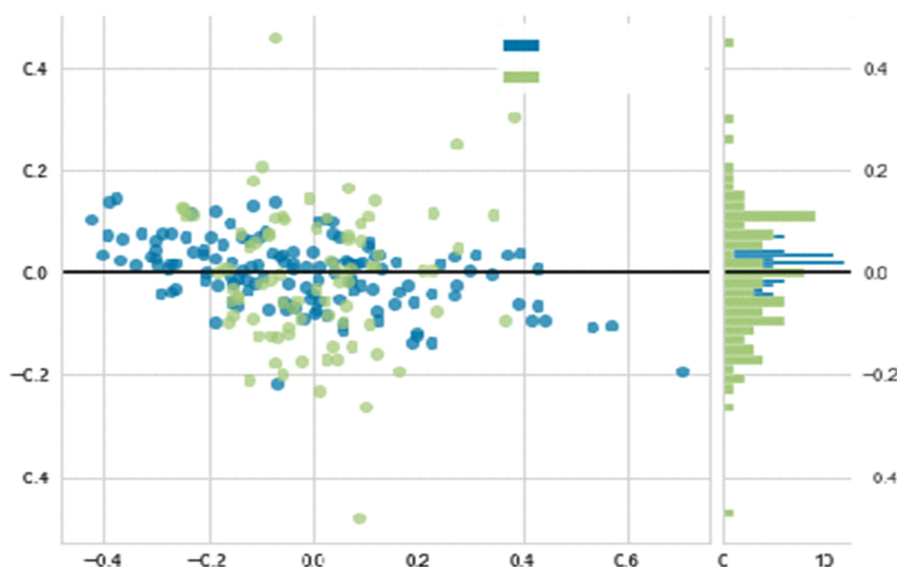


Figure 7.
Distribution of
residuals in the
CatBoost model with
stationary variables

Source(s): Figure by authors

The results of the *Random Forest Model* were somewhat less accurate than the results of ET. Likewise, the computing time was shorter (Figure 11).

Random Forest presented no overlearning problems, with R^2 values for the training and for the test datasets of 0.583 and 0.517, respectively.

The independent variable which had markedly greater importance when making predictions for this model was once again the 10-year German bond, followed to a lesser extent by the total monetary base, the variable FedFunds and private credit both in Europe and the USA.

Two voting models were developed that yielded some of the best results this time with the stationary variables.

The first model in which OMP–CatBoost–AdaBoost–XGBoost were combined and that generated new predictions through consensus between estimators (majority voting average probabilities). The metrics resulting from this combination were very encouraging.

The voting model yielded: MAE = 0.1360, RMSE = 0.1759 and MAPE = 1.3781 (Figure 12).

The second voting model included those models that not only improved the base model but also presented less overfitting: OMP–Random Forest–Extra Trees.

The metrics of this new model yielded the best results of the study: MAE = 0.1312, RMSE = 0.1715 and MAPE = 1.4816.

The principal objective of this latter model was to improve the metrics that had previously been obtained without committing the error of increasing the overfitting. This objective was satisfactorily achieved through the design of the voting model. The R^2 values of the training and the test datasets were 0.647 and 0.552, respectively, revealing a difference of only 0.095, which can be attributed to what is known as the generalization gap.

6. Analysis and discussion of the results

According to the results of each algorithm and the metrics with which they are comparable, we now focus on Table 6, which summarizes the indicators of the models, among which the



Figure 8.
AdaBoost model

Source(s): Figure by authors

RMSE may be highlighted to consider the goodness of fit of these models and whether there is a considerable presence of overfitting in the model.

The low performance of the models with nonstationary variables is evident (in accordance with the literature), showing problems of overfitting and the worst metrics among all the models. Only the Bayesian Ridge and the Voting model, prepared with a linear regression model and the earlier Bayesian Ridge model, managed to overcome the benchmark model with stationary variables among all the models that were tested.

The machine learning models based on stationary variables presented better RMSE values, as well as the other metrics under observation.

The CatBoost, AdaBoost and XGBoost models and the first voting model, prepared with the OMP, CatBoost, AdaBoost and XGBoost models, presented strong overfitting that represents a major limitation, despite their exceptional results.

On the other hand, the OMP, RF and ET models, and the second voting model (prepared with those three models) yielded exceptional RMSE values of 0.1786, 0.1858, 0.1782 and

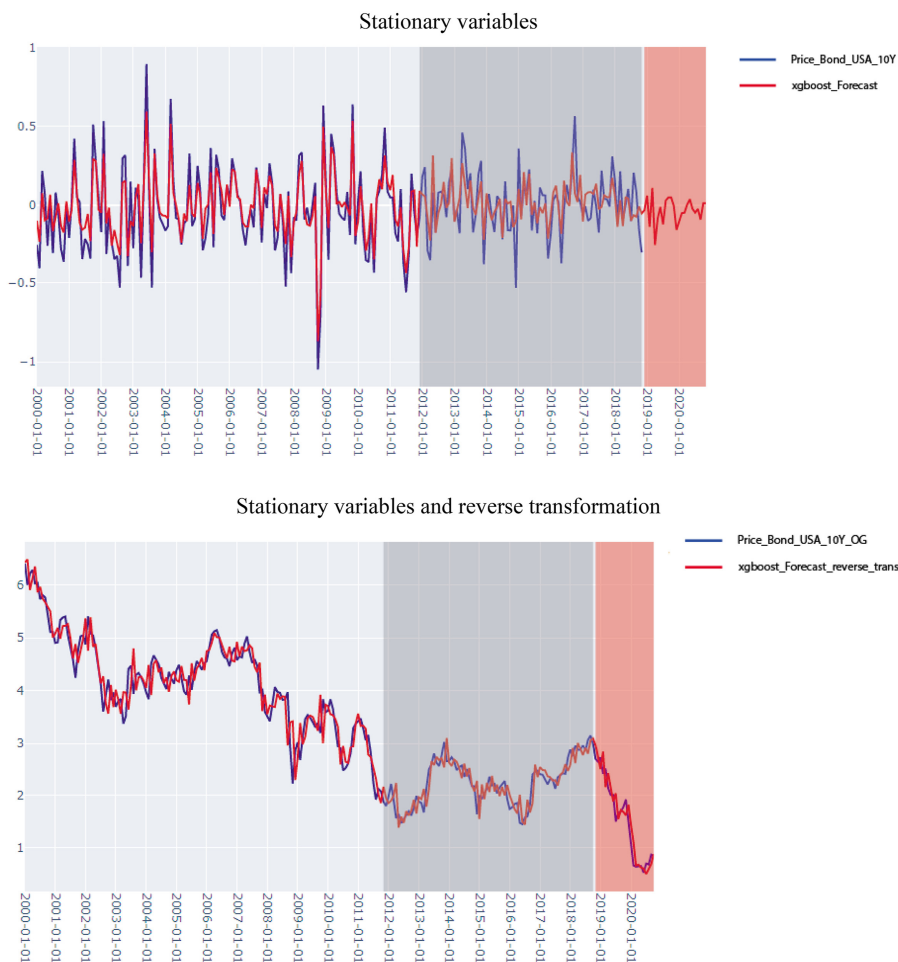


Figure 9.
XGBoost model with
stationary variables

Source(s): Figure by authors

0.1715, respectively. The presence of overfitting in each of these models was disregarded, so it was concluded that these four machine learning models with stationary variables yielded the best results. The second voting model with stationary variables was the one that yielded the best metrics, followed by the ET, the OMP and, finally, the RF models.

It must be noted that some variables were removed from some of the models, either because they only contributed to background noise or because their removal alleviated overfitting, helping to determine which variables had been the most important in the models.

The variable that had the most obvious relevance when predicting the dependent variable was the price of the 10-year German government bond, as its nature and behavior were very similar to the dependent variable, with which it showed a very high correlation (0.89). This correlation points to the presence of multicollinearity, although the value of 0.89 was below the threshold that is usually employed in the literature of 0.9.

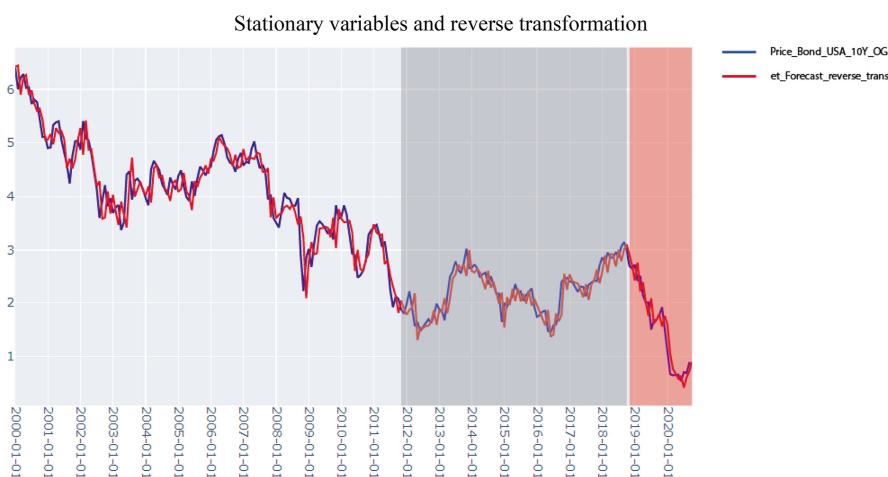
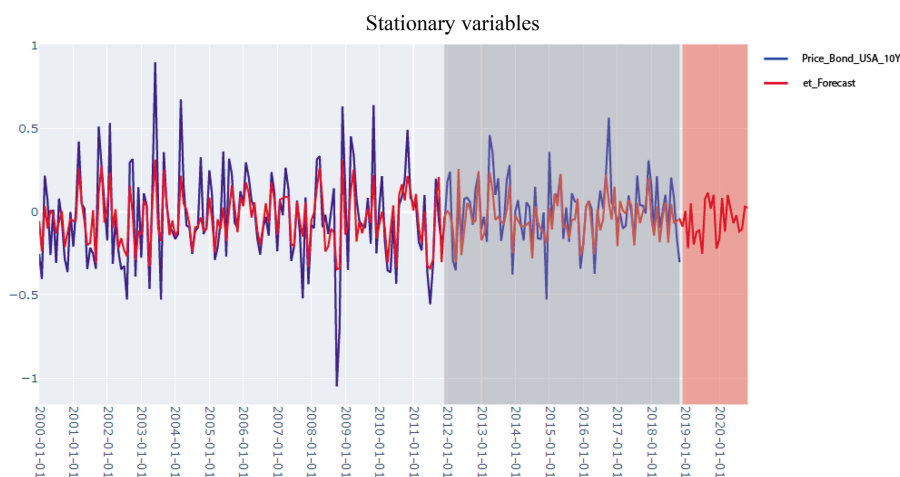


Figure 10.
ET model

Source(s): Figure by authors

When testing other models in which this threshold was lower, there were significant losses of predictive capability. Maintaining the aforementioned variable, which to a great extent helped to predict the price of the US bond, when considering these circumstances, was therefore recommendable.

The representative variables of public liquidity had no high impact on the models. Among the variables of this group, the reference interest rate of the Central European Bank and the FedFunds (Federal Funds Rate) variables stood out most of all, having a relative relevance in models such as Random Forest, XGBoost and OMP.

The representative variables of private liquidity were the variables that more than any others helped the predictions of the different models (second only to the safe assets used as predictors). The variables of this group that may be highlighted because of their importance

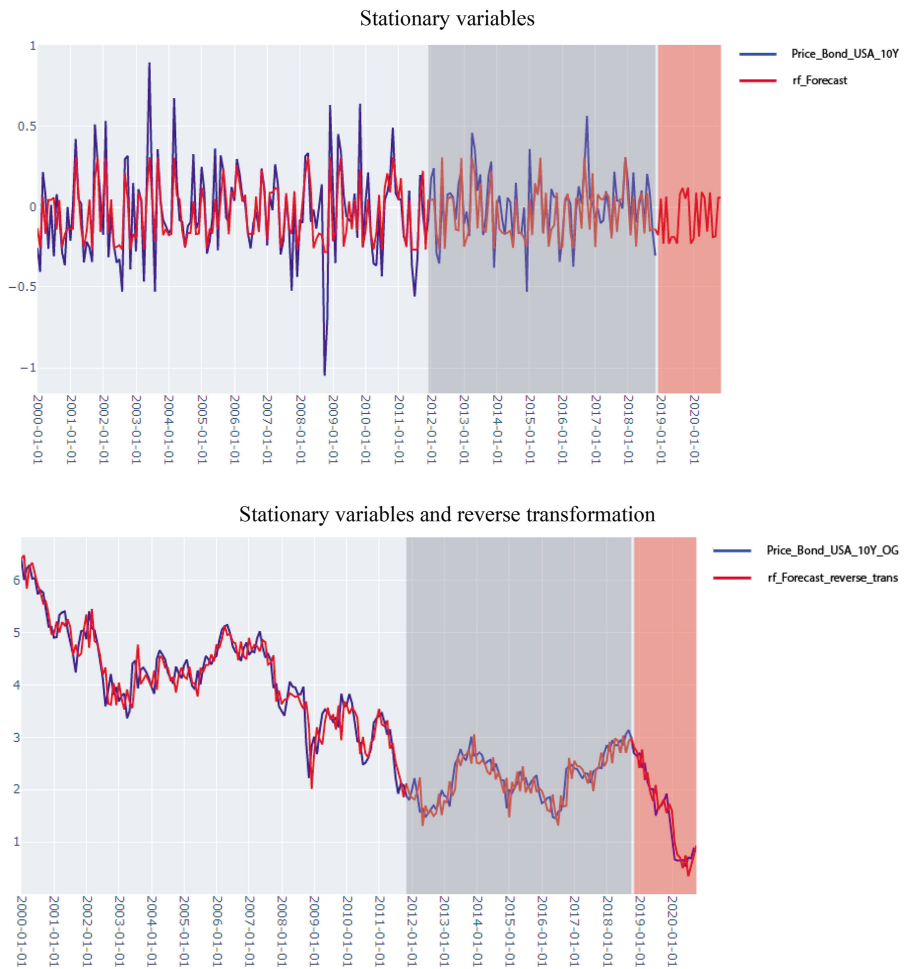


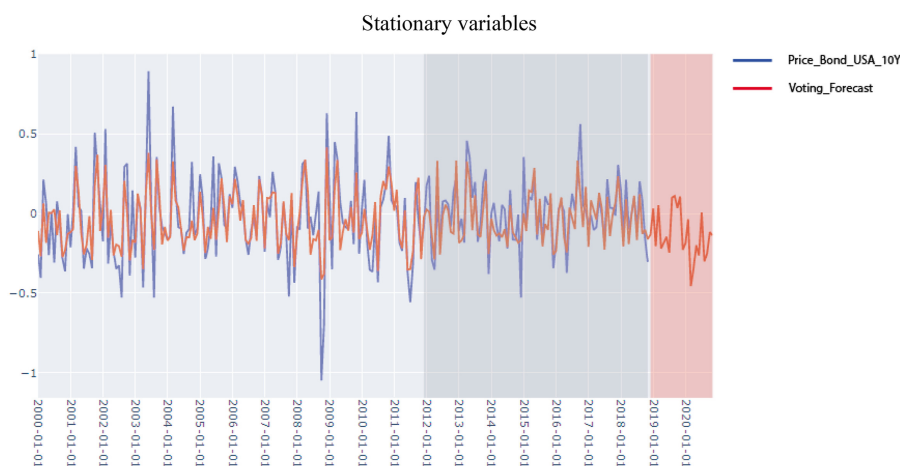
Figure 11.
Random forest model

Source(s): Figure by authors

were principally the M_1 monetary aggregate (for Europe and the USA), the M_0 monetary aggregate and the total monetary base (for the USA), closely followed by credit to the private nonfinancial sector (for Europe and the USA) and, to a lesser extent, the percentile change in consumer credit.

International liquidity is, outstandingly, the type of liquidity with the vaguest of definitions and it is especially difficult to measure with precision, consequently, the majority of its representative variables have been converted into background noise in the models. In the feature selection step, both the index of global liquidity prepared with BIS indicators and the DXY indicator of the Dollar versus the basket of currencies (the variables employed as proxies of international liquidity) were removed from the majority of the models. In general, it contributed to noise and could have generated overfitting in the models that maintained it.

The price of the 10-year German bond was not the only variable to have substantial influence on the results of the models. Other variables grouped as safe assets and used as



Stationary variables and reverse transformation

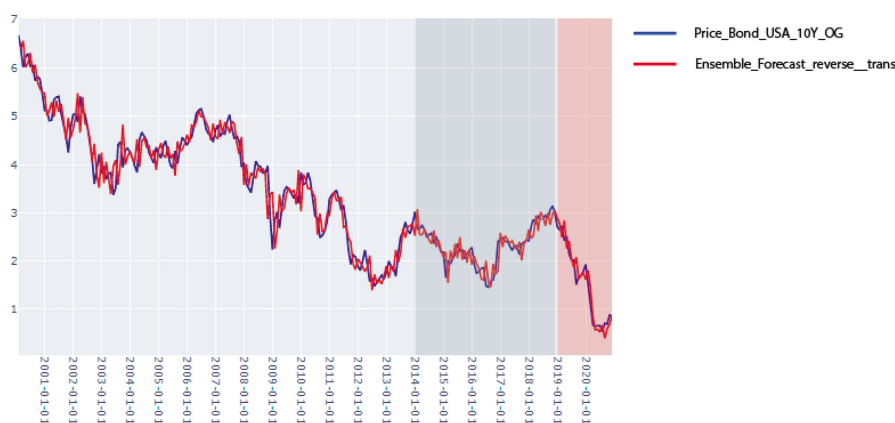


Figure 12.
First voting model

Source(s): Figure by authors

predictors also turned out to be useful for predicting the dependent variable, especially the variation of the Euro with regard to the Dollar, followed by the closing price of the VIEIX and to a lesser degree the variation of the S&P 500.

7. Conclusions, recommendations and limitations

On the basis of the above discussion, the following conclusions can now be presented.

The two models that yielded the best results, both in terms of their RMSE and their quality of fit, were based on decision tree algorithms: the Extra Trees model and the Random Forest model. There was also one model based on the OMP greedy algorithm. (All three models employed variables that had been converted into stationary variables following their transformation into the first difference in the natural log.) These models yielded better results

than both the traditional regression base models and the nonstationary variables, as we can see in Table 6, thereby responding to research questions Q₁ and Q₂.

A series of new predictions were generated using a combination of these three best-performing models in an ensemble (second voting model with stationary variables), predictions that were better than the other results obtained in this study, with an RMSE of 0.1715. On average, their predictions of 10-year US bond prices over 2019 and 2020 only deviated 0.1715 from the real price (a deviation expressed in the same units used for the bond); results which are responses to Q₃ and Q₄.

In line with the theoretical predictions, the models with nonstationary variables presented clear overfitting as previously mentioned. Regarding the stationary models, those based on boosting and greedy algorithms presented overlearning in each case, except for the one running the OMP algorithm. Tree-based models were the best performers in this regard, especially the ET model, which was specifically selected for this specific reason as supported by the theory presented in Section 2 (State of the Art).

The second ensemble comprising OMP, ET and Random Forest models showed no overfitting either, obtaining quality results in that regard. However, the first voting model, made up of the boosting and greedy algorithm models mentioned above, presented serious overlearning problems. An expected result, considering that the models with which it was configured also presented this problem, together with the fact that the voting approach makes the models more robust, although it can also contribute to increasing the difference between the proportion of variance explained between the training set and the test set, thereby contributing to overfitting (Q₅).

Likewise, the second voting model presented fewer overlearning problems, while the nonstationary variables suffered more from this problem (Q₅).

In response to Q₆, the variables that contributed most to the formulation of the bond price predictions were the price of the 10-year German government bond, the closing price of the VEIEX, and the M₀ and M₁ monetary aggregates. The variables that represent the value of other assets considered as safe and assets that represent private liquidity have been the most useful for the preparation of the models. While the public liquidity variables that were defined contributed to a lower number of models, they nevertheless did so in a significant manner (FedFunds and reference interest rate of the BCE). In general, they generated inconvenient levels of noise when training the models, although far less so than the variables of international liquidity.

Regarding the recommendations derived from the analysis of liquidity metrics, first, the importance of tracking private liquidity metrics such as banking credit and monetary aggregates was clear, given that when those indexes descended, clear upward trends of the price of the US bonds were observed, and *vice versa*.

Likewise, the tracking of public liquidity, given that there was a strong increase in the prices of risk assets when they increased in a significant way and *vice versa* when they fell. One important fall of the FedFunds predicted future price rises of the S&P 500 index.

Theoretical implications with regard to international liquidity are focused on the need to develop variables that reflect this category of liquidity more precisely. It is one of the most difficult challenges for researchers, due to the vast breadth of this conceptual dimension, as well as its diffuse and mutable definition.

The monitoring of liquidity could facilitate the identification of systemic risks within the financial system and the consequences of economic activity, as it has already been of assistance to policymakers. In an earlier initiative, the Federal Reserve of the US (<https://www.boj.or.jp/en/research/brp/fsr/index.htm>) and the Bank of Japan (<https://www.boj.or.jp/en/research/brp/fsr/index.htm>) began to publish their respective Financial Stability Reports. The methodology followed by the Federal Reserve of the US for the monitoring of financial stability may be found in Adrian *et al.* (2015).

Our study differs, insofar as it approaches the same concept of liquidity/financial stability from a proactive viewpoint rather than from a reactive one. The predictability of the existing forecasting models could be heightened, by selecting liquidity and the corresponding key variables as the leading indicators for forecasting changes to economic cycles.

The limitations arise from the inclusion of a series of variables of little utility when advancing predictions on the dependent variable. This increase in the number of futures only caused overfitting of the model, in particular, in relation to some international and public liquidity variables.

The correlation between US bond and German bond gave rise to problems of multicollinearity, though it never passed the established threshold in normal use. Although, they never affected the predictions.

The singular events arise from the interruption of COVID-19 in 2020, adding a major difficulty to the production of accurate predictions for that year.

References

- Abellán, J. and Castellano, J.G. (2017), "A comparative study on base classifiers in ensemble methods for credit scoring", *Expert Systems with Applications*, Vol. 73, pp. 1-10, doi: [10.1016/j.eswa.2016.12.020](https://doi.org/10.1016/j.eswa.2016.12.020).
- Adrian, T., Covitz, D. and Liang, N. (2015), "Financial stability monitoring", *Annual Review of Financial Economics*, Vol. 7, pp. 357-395.
- Alessi, L. and Detken, C. (2011), "Quasi real time early warning indicators for costly asset price boom/bust cycles: a role for global liquidity", *European Journal of Political Economy*, Vol. 27 No. 3, pp. 520-533, doi:[10.1016/j.ejpoleco.2011.01.003](https://doi.org/10.1016/j.ejpoleco.2011.01.003).
- Barrow, A. (2017), *Stanley Druckenmiller on Liquidity, Macro, and Margins*, MACRO OPS, available at: <https://macro-ops.com/stanley-druckenmiller-on-liquidity-macro-margins/>
- Bernanke, B.S., Bertaut, C.C., Demarco, L. and Kamin, S.B. (2011), "International capital flows and the return to safe assets in the United States, 2003-2007", FRB International Finance Discussion Paper, (1014), doi: [10.2139/ssrn.1837780](https://doi.org/10.2139/ssrn.1837780), available at SSRN: <https://ssrn.com/abstract=1837780>
- Borio, C., McCauley, R. and McGuire, P. (2011), "Global credit and domestic credit booms. Bank for International Settlements", *Quarterly Review*, Part 4, Vol. 2011 September, available at: https://www.bis.org/publ/qtrpdf/r_qt1109f.pdf
- Brockwell, P.J. and Davis, R.A. (2009), *Time Series: Theory and Methods*, Springer Science & Business Media, New York.
- Brownlee, J. (2016a), "A gentle introduction to xgboost for applied machine learning", *Machine Learning Mastery*, available at: <https://machinelearningmastery.com/gentle-introduction-xgboost-applied-machine-learning/>
- Brownlee, J. (2016b), "Overfitting and underfitting with machine learning algorithms", *Machine Learning Mastery*, available at: <https://machinelearningmastery.com/overfitting-and-underfitting-with-machine-learning-algorithms/>
- Bruno, V. and Shin, H.S. (2015a), "Capital flows and the risk-taking channel of monetary policy", *Journal of Monetary Economics*, Vol. 71, pp. 119-132.
- Bruno, V. and Shin, H.S. (2015b), "Cross-border banking and global liquidity", *The Review of Economic Studies*, Vol. 82 No. 2, pp. 535-564.
- Caruana, J. (2013), "Global liquidity: where do we stand?", *In speech at the Bank of Korea International Conference*.
- Cesa-Bianchi, A., Cespedes, L.F. and Rebucci, A. (2015), "Global liquidity, house prices, and the macroeconomy: evidence from advanced and emerging economies", *Journal of Money, Credit and Banking*, Vol. 47 No. S1, pp. 301-335, doi: [10.1111/jmcb.12204](https://doi.org/10.1111/jmcb.12204).

- Chen, M.S.F.S.F., Liu, M.P., Maechler, A.M., Marsh, C., Saksonovs, M.S. and Shin, M.H.S. (2012), *Exploring the Dynamics of Global Liquidity*, International Monetary Fund, Washington, DC.
- Chhugani, R. (2020), *An Overview of Autocorrelation, Seasonality and Stationarity in Time Series Data*, Analytics India Magazine, available at: <https://analyticsindiamag.com/an-overview-of-autocorrelation-seasonality-and-stationarity-in-time-series-data/>
- Chung, K., Lee, J.E., Loukoianova, M.E., Park, M.H. and Shin, M.H.S. (2014), *Global Liquidity through the Lens of Monetary Aggregates*, International Monetary Fund, Washington, DC.
- Dunn, K. (2021), "Avoid R-squared to judge regression model performance", *Towards Data Science*, available at: <https://towardsdatascience.com/avoid-r-squared-to-judge-regression-model-performance-5c2bc53c8e2e>
- Eickmeier, S., Gambacorta, L. and Hofmann, B. (2014), "Understanding global liquidity", *European Economic Review*, Vol. 68, pp. 1-18, doi: [10.1016/j.eurocorev.2014.01.015](https://doi.org/10.1016/j.eurocorev.2014.01.015).
- Galindo, J. and Tamayo, P. (2000), "Credit risk assessment using statistical and machine learning: basic methodology and risk modeling applications", *Computational Economics*, Vol. 15 No. 1, pp. 107-143, doi: [10.1023/A:1008699112516](https://doi.org/10.1023/A:1008699112516).
- Geurts, P., Ernst, D. and Wehenkel, L. (2006), "Extremely randomized trees", *Machine Learning*, Vol. 63 No. 1, pp. 3-42.
- Goda, T., Lysandrou, P. and Stewart, C. (2013), "The contribution of US bond demand to the US bond yield conundrum of 2004-2007: an empirical investigation", *Journal of International Financial Markets, Institutions and Money*, Vol. 27, pp. 113-136, doi: [10.1016/j.jintfin.2013.07.012](https://doi.org/10.1016/j.jintfin.2013.07.012).
- Guerra, P., Castelli, M. and Côte-Real, N. (2022), "Machine learning for liquidity risk modelling: a supervisory perspective", *Economic Analysis and Policy*, Vol. 74, pp. 175-187, doi: [10.1016/j.eap.2022.02.001](https://doi.org/10.1016/j.eap.2022.02.001).
- Gutierrez, D. (2014), "Ask a data scientist: data leakage", *Inside BIGDATA*, available at: <https://insidebigdata.com/2014/11/26/ask-data-scientist-data-leakage/>
- Hatzius, J., Hooper, P., Mishkin, F.S., Schoenholtz, K.L. and Watson, M.W. (2010), *Financial Conditions Indexes: A Fresh Look after the Financial Crisis (No. W16150)*, National Bureau of Economic Research, Cambridge, MA.
- Hellwig, K.P. (2021), *Predicting Fiscal Crises: A Machine Learning Approach*, International Monetary Fund, Washington, DC.
- Howell, M. (2020), *Capital Wars: The Rise of Global Liquidity*, Palgrave Macmillan, New York.
- James, G., Witten, D., Hastie, T. and Tibshirani, R. (2013), *An Introduction to Statistical Learning*, Springer, New York, Vol. 112, p. 18.
- Jeanne, O. and Sandri, D. (2020), *Global Financial Cycle and Liquidity Management (No. W27901)*, National Bureau of Economic Research, Cambridge, MA.
- Khosravy, M., Dey, N. and Duque, C. (2020), *Compressive Sensing in Healthcare*, Academic Press, London.
- Landau, J.P. (2011), "Global liquidity-concept, measurement and policy implications", *CGFS Papers*, Vol. 45, pp. 1-33.
- Lane, P.R. and McQuade, P. (2014), "Domestic credit growth and international capital flows", *The Scandinavian Journal of Economics*, Vol. 116 No. 1, pp. 218-252.
- Miranda-Agrippino, S. and Rey, H. (2020), "The global financial cycle after Lehman", *AEA Papers and Proceedings*, Vol. 110, pp. 523-528, doi: [10.1257/pandp.20201096](https://doi.org/10.1257/pandp.20201096).
- Oh, S. (2019), "Feature interaction in terms of prediction performance", *Applied Sciences*, Vol. 9 No. 23, p. 5191.
- Sahin, E.K. (2020), "Assessing the predictive capability of ensemble tree methods for landslide susceptibility mapping using XGBoost, gradient boosting machine, and random forest", *SN Applied Sciences*, Vol. 2 No. 7, pp. 1-17.

-
- Shin, H.S. and Shin, K. (2011), *Procyclicality and Monetary Aggregates* (No. W16836), National Bureau of Economic Research, Cambridge, MA.
- Thiesen, S. (2020), "CatBoost regression in 6 minutes", *Towards Data Science*, available at: <https://towardsdatascience.com/catboost-regression-in-6-minutes-3487f3e5b329>
- Wang, R. (2012), "AdaBoost for feature selection, classification and its relation with SVM, a review", *Physics Procedia*, Vol. 25, pp. 800-807.
- Zucchi, K. (2021), "Why the 10-year U.S. Treasury yield matters", Investopedia, available at: <https://www.investopedia.com/articles/investing/100814/why-10-year-us-treasury-rates-matter.asp>

Further reading

- Borio, C. (2014), "The financial cycle and macroeconomics: what have we learnt?", *Journal of Banking and Finance*, Vol. 45, pp. 182-198.

Corresponding author

Ignacio Manuel Luque Raya can be contacted at: luquenacho0@gmail.com