



# Nonparametric long term prediction of stock returns with generated bond yields



Michael Scholz<sup>a,\*</sup>, Stefan Sperlich<sup>b</sup>, Jens Perch Nielsen<sup>c</sup>

<sup>a</sup> Department of Economics, University of Graz, Universitätsstraße 15/F4, 8010 Graz, Austria

<sup>b</sup> Geneva School of Economics and Management, Université de Genève, Bd du Pont d'Arve 40, 1211 Genève 4, Switzerland

<sup>c</sup> Faculty of Actuarial Science and Insurance, Cass Business School, 106 Bunhill Row, London, EC1Y8TZ, UK

## ARTICLE INFO

### Article history:

Received December 2014

Received in revised form

March 2016

Accepted 26 April 2016

Available online 6 May 2016

### JEL classification:

C14

C53

C58

G17

G22

### Keywords:

Prediction

Stock returns

Bond yield

Cross validation

Generated regressors

## ABSTRACT

Recent empirical approaches in forecasting equity returns or premiums found that dynamic interactions among the stock and bond are relevant for long term pension products. Automatic procedures to upgrade or downgrade risk exposure could potentially improve long term performance for such products. The risk and return of bonds is more easy to predict than the risk and return of stocks. This and the well known stock-bond correlation motivates the inclusion of the current bond yield in a model for the prediction of excess stock returns. Here, we take the actuarial long term view using yearly data, and focus on nonlinear relationships between a set of covariates. We employ fully nonparametric models and apply for estimation a local-linear kernel smoother. Since the current bond yield is not known, it is predicted in a prior step. The structure imposed this way in the final estimation process helps to circumvent the curse of dimensionality and reduces bias in the estimation of excess stock returns. Our validated stock prediction results show that predicted bond returns improve stock prediction significantly.

© 2016 Elsevier B.V. All rights reserved.

## 1. Introduction and motivation

For a long time predicting asset returns has been a main objective in the empirical finance literature. It started with predictive regressions of independent variables on stock market returns. Typically, valuation ratios are used that primarily characterise the stock, for example the dividend price ratio, the dividend yield, the earnings price ratio or the book-to-market ratio. Other variables related to the interest rate like treasury-bill rates and the long-term bond yield, or macroeconomic indicators like inflation and the consumption wealth ratio, are often incorporated to improve prediction. For a detailed overview we refer to the examples and discussion in [Rapach et al. \(2005\)](#) or [Campbell and Thompson \(2008\)](#).

In this paper, we take the actuarial long term view using yearly data, and focus on nonlinear relationships between a set

of covariates. There are not many historical years in our records and data sparsity is of great importance in our approach. One could also use data of higher frequency as weekly or daily data, but one has to remember that the logistics of prediction is then very different. In our approach using yearly data bias might be of big importance while variance becomes less of an issue. In other words, the usual variance-bias trade-off depends on the horizon. An adequate model for monthly data might perform worse for yearly data and vice versa. The reason for the use of yearly data is our interest in actuarial models of long term savings and their possible econometric improvement (see e.g. [Bikker et al., 2012](#), [Guillen et al., 2013a,b](#), [Owaddally et al., 2013](#), [Guillen et al., 2014](#), or [Gerrard et al., 2014](#)). Our favoured methodology of validating the fully nonparametric models that we employ for the long term yearly data also originates from the actuarial literature (see [Nielsen and Sperlich, 2003](#)).

The apparent predictability found by many authors was controversially discussed. As [Lettau and Nieuwerburgh \(2008\)](#) note, correct inference is problematic due to the high persistence of financial ratios, which have poor out-of sample forecasting power that moreover shows significant instability over time. Therefore,

\* Corresponding author. Tel.: +43 316 380 7112; fax: +43 316 380 697112.

E-mail addresses: [michael.scholz@uni-graz.at](mailto:michael.scholz@uni-graz.at) (M. Scholz), [stefan.sperlich@unige.ch](mailto:stefan.sperlich@unige.ch) (S. Sperlich), [jens.nielsen.1@city.ac.uk](mailto:jens.nielsen.1@city.ac.uk) (J.P. Nielsen).

the question of whether empirical models are really able to forecast the equity premium more accurately than the simple historical mean was intensively debated in the finance literature. Recently, Goyal and Welch (2008) fail to provide benefits of predictive variables compared to the historical mean. In contrast, Rapach et al. (2010) recommend a combination of individual forecasts. Their method includes the information provided from different variables and reduces this way the forecast volatility. Elliott et al. (2013) suggest a new method to combine linear forecasts based on subset regressions and show improved performance over the classical linear prediction methods. More recently, Scholz et al. (2015) propose a simple bootstrap test about the true functional form to evidence that the null of no predictability of returns can be rejected when using information such as earnings.

A direct comparison of stocks and bonds, mostly used by practitioners, makes the so-called FED model. It relates yields on stocks, as ratios of dividends or earnings to stock prices, to yields on bonds. Asness (2003) shows the empirical descriptive power of the model, but notes also that it fails in predicting stock returns. One of his criticisms is the comparison of real numbers to nominal ones. Actually, most studies discuss separately the predictability in stock and bond markets. However, Shiller and Beltratti (1992) analyse the relation between stock prices and changes in long-term bond yields. Fama and French (1993) find that stock returns have shared variation due to the stock-market factors, and they are linked to bond returns through shared variation in the bond-market. Engsted and Tanggaard (2001) pose the interesting question of whether expected returns on stocks and bonds are driven by the same information, and to what extent they move together. In their empirical setting, they find that excess stock and bond returns are positively correlated. Aslanidis and Christiansen (2014) adopt quantile regressions to scrutinise the realised stock-bond correlation and the link to the macroeconomy. Tsai and Wu (2015) analyse the bond and stock market responses to changes in dividends. Lee et al. (2013) find dynamic interactions among the stock, bond, and insurance markets. For additional literature on the relation between stock and bond returns (especially co-movements, joint distributions, or correlations), see, for example, Lim et al. (1998), Ilmanen (2003), Guidolin and Timmermann (2006), Connolly et al. (2010), Baele et al. (2010), or Bekaert et al. (2010).

One overall idea of the this paper is to exploit the interrelationship of present values of stock returns and bond returns. They are after all both discounted cash flows. Our underlying assumption implies that expected returns are associated with variables related to longer-term aspects of business conditions, as mentioned in Campbell (1987). Consequently, we include in a nonparametric prediction model of excess stock returns the bond yield of the same year. This way, the bond captures a most important part of the stock return, namely the part related to the change in long-term interest rate. Nonlinear forecasting methods are a growing area of empirical research, see for example Guidolin and Timmermann (2006), McMillan (2007), or Guidolin et al. (2009). Nielsen and Sperlich (2003) find a significant improvement in the prediction power of excess stock returns due to the use of nonlinear smoothing techniques. Based on their findings, we focus on nonlinear relationships between a set of covariates and the bond yield of the same year. We apply for estimation a local-linear kernel smoother which nests the linear model without bias. For the purpose of bandwidth selection and to measure the quality of prediction we use a cross-validation measure of performance. It is a generalised version of the validated  $R^2$  of Nielsen and Sperlich (2003) and allows for a direct comparison of the proposed model with the historical mean.

An obvious problem is that the current bond yield is unknown. Thus, we have to predict it in a first step. Here, we also employ

fully nonparametric models and use a local-linear kernel smoother. This raises the question why it is necessary to use a two-step procedure. One could directly include the variables used for the bond prediction when forecasting stock returns. The problem is that such a model would suffer from the curse of dimensionality and complexity in several aspects: The dimension of the covariates, possible over-fitting, and the interpretability. In nonparametrics it is well known that the import of structure is an appropriate way to circumvent these problems.<sup>1</sup> Furthermore, Park et al. (1997) showed that an appropriate transformation of the predictors can significantly improve nonparametric prediction. In our approach, we utilise the additional knowledge about the structure that is inherent in the economic process that generates the data. We find that the inclusion of the generated variable shows notable improvement in the prediction of excess stock returns. Note that one does not achieve computational efficiency, but rather estimation efficiency from adding information. To our knowledge we are the first including nonparametrically generated regressors for nonparametric prediction of time series data. Therefore we also have to develop the theoretical justification for the use of constructed variables in nonparametric regression when the data are dependent.

For the empirical part we use annual Danish stock and bond market data (also used in Lund and Engsted, 1996, Engsted and Tanggaard, 2001, or Nielsen and Sperlich, 2003). We find that the inclusion of predicted bond yields greatly improves the prediction quality of stock returns in terms of the validated  $R^2$ . With our best prediction model for one-year stock returns we not only beat the simple historical mean but we also observe a large increase in validated  $R^2$  from 5.9% to 28.3%. To underline our findings, we also include in our empirical analysis the prediction of the ratio of stock returns and dividend yields getting similar results.

The paper proceeds as follows. Section 2 describes the prediction framework and the measure of validation. The mathematical justification is introduced in Section 3. Section 4 presents our findings from an empirical and a small simulation study. Section 5 concludes. Finally, Appendix contains proofs of our theoretical results.

## 2. The prediction framework

In the financial and actuarial literature traditional approaches like the classic  $R^2$ , the adjusted  $R^2$ , goodness-of-fit or testing methods are mainly used to measure in-sample forecasting power. More recently, out-of-sample statistics and tests are discussed, see for example Inoue and Kilian (2004), Clark and West (2006), Goyal and Welch (2008), or Campbell and Thompson (2008). In our study, we use a generalised version of the validated  $R^2$  ( $R_V^2$ ) of Nielsen and Sperlich (2003) based on leave- $k$ -out cross-validation. It measures how well a model predicts in the future compared to the historical mean. The classical  $R^2$  is often used, easy to calculate and has a straight forward interpretation. But it can hardly be used for prediction nor for comparison issues as it always prefers the most complex model. See also Valkanov (2003) or Dell'Aquila and Ronchetti (2006) for more relevant arguments for disregarding the classical  $R^2$  measure when selecting a model. For comparison often the adjusted  $R^2$  is applied, which penalises complexity via a degree of freedom adjustment. It is well known that this correction does not work in our case, see for example Sperlich et al. (1999).

The idea of the  $R_V^2$  is to replace total variation and not explained variation by their leave- $k$ -out cross-validated analogs. Note that cross-validation (cv) is a quite common in the nonparametric time

<sup>1</sup> An other possibility could be the optimal choice of regressors, see Vieu (1994).

Download English Version:

<https://daneshyari.com/en/article/5076362>

Download Persian Version:

<https://daneshyari.com/article/5076362>

[Daneshyari.com](https://daneshyari.com)