

FORECASTING OF RECESSIONS VIA DYNAMIC PROBIT FOR TIME SERIES: REPLICATION AND EXTENSION OF KAUPPI AND SAIKKONEN (2008)

Park, Byeong U.; Simar, Léopold and
Valentin Zelenyuk

REPRINT | 2019 / 14



Forecasting of recessions via dynamic probit for time series: replication and extension of Kauppi and Saikkonen (2008)

Byeong U. Park¹ · Léopold Simar² · Valentin Zelenyuk³

Received: 12 June 2017 / Accepted: 25 March 2019
© Springer-Verlag GmbH Germany, part of Springer Nature 2019

Abstract

In this work, we first replicate the results of the fully parametric dynamic probit model for forecasting US recessions from Kauppi and Saikkonen (Rev Econ Stat 90(4):777–791, 2008) [which is in the spirit of Estrella and Mishkin (Rev Econ Stat 80(1):45–61, 1998) and Dueker (Rev Fed Reserve Bank St Louis 79(2):41–51, 1997)] and then contrast them to results from a nonparametric local-likelihood dynamic choice model for the same data. We then use expanded data to gain insights on whether these models could have warned the public about approach of the latest recession, associated with the Global Financial Crisis. Finally, we also apply both approaches to gain insights for 2018.

Keywords Forecasting of recessions · Nonparametric quasi-likelihood · Local-likelihood · Dynamic discrete choice

JEL Classification C14 · C22 · C25 · E37

1 Introduction

The goal of this article is threefold. First, we start by replicating some of the results from the parametric (linear) dynamic probit of Kauppi and Saikkonen (2008), who in turn followed and refined the approach originated in the seminal works of Estrella and Mishkin (1995, 1998) and Dueker (1997), for forecasting the probability of US recessions.

✉ Valentin Zelenyuk
v.zelenyuk@uq.edu.au

¹ Department of Statistics, Seoul National University, Seoul, Korea

² Institut de Statistique, Biostatistique et Sciences Actuarielles, Université Catholique de Louvain, Louvain-la-Neuve, Belgium

³ School of Economics and Centre for Efficiency and Productivity Analysis (CEPA), The University of Queensland, 530, Colin Clark Building (39), St Lucia, Brisbane, QLD 4072, Australia

Second, and more importantly, we want to do a nonparametric validation of their interesting results—by checking how sensitive they are to the parametric assumption about the linearity of the index function in their parametric probit model. Indeed, the parametric methods in general, and probit or logit approaches in particular (whether in a cross section, panel or time series context as ours), yield inconsistent estimates when the parametric assumptions are misspecified.

To perform the nonparametric validation, one may use many alternatives proposed in the literature and all have their own merits and limitations [e.g., see Henderson and Parmeter (2015) for related discussions and references therein for key works in nonparametric estimation]. Here, we use (with some modifications) the approach recently developed by Park et al. (2017), who in turn generalized the nonparametric quasi-likelihood approach of Fan et al. (1995) to a dynamic time series context. An important feature of this approach is that it embraces the parametric dynamic probit approach of Kauppi and Saikkonen (2008) as a special case, while requiring no parametric assumptions on the index function. Our modification of Park et al. (2017) is related to the use of the discrete kernel to handle the lagged dependent variable—they used a kernel in the spirit of Aitchison and Aitken (1976), while we used a generalized version of it where we combine a discrete kernel with an adaptive continuous kernel of Li et al. (2016), which allows for the continuous bandwidth to vary with the discrete variable.¹

Third, we use these two approaches for an expanded data set that includes more recent periods. Specifically, we try to get some insights on whether these models could have warned about the latest recession, associated with the Global Financial Crisis (GFC), if they were applied in 2007 and early 2008. Finally, we also apply both approaches to gain insights about 2018.

2 Some background

In their influential paper, Kauppi and Saikkonen (2008) asked a very important question that worries many people in the world: ‘What is the probability of a recession a month ahead or a year ahead?’ To tackle it, they revisited the approach originated by Estrella and Mishkin (1995, 1998) and Dueker (1997) for forecasting recessions using parametric *probit* model with only one continuous regressor (the spread) and one discrete regressor (the lagged dependent variable). Specifically, in their simple yet quite powerful model, the random variable Y_t takes on two values, 1 and 0, representing such events as ‘recession’ and ‘no recession’ at time t , and is assumed to follow the Bernoulli distribution with probability p_t , which in turn depends on past information on realizations of some explanatory variables and the realization(s) of Y_t available in a period $s < t$.

Formally, the focus of interest here is to estimate the conditional mean (which here is also a conditional probability) of Y_t to predict its realization in time period t , conditional on some explanatory factors known in a period $s < t$, some continuous

¹ One could also use other kernels and bandwidths, including the non-geometric discrete kernels like the discrete Epanechnikov kernel (see Chu et al. 2015). The asymptotic theory from Park et al. (2017), under their fairly mild assumptions, can be extended to encompass these and other cases.

variables $\mathbf{X}_s$ and some discrete variables $\mathbf{Z}_s$, i.e.,

$$\begin{aligned} m_s(\mathbf{x}, \mathbf{z}) &:= E(Y_t | \mathbf{X}_s = \mathbf{x}, \mathbf{Z}_s = \mathbf{z}) \\ &= P(Y_t = 1 | \mathbf{X}_s = \mathbf{x}, \mathbf{Z}_s = \mathbf{z}) = g^{-1}(f(\mathbf{x}, \mathbf{z})), \end{aligned} \quad (2.1)$$

where g is usually referred to as a link function and f as an index function. Most studies, and Kauppi and Saikkonen (2008) in particular, take the probit link for g and assume affine functional form for f and relaxing the latter is the main focus of this paper.

While Kauppi and Saikkonen (2008) considered many variations, the best performing models had only one continuous explanatory variable (lagged interest rate spread), similar as in Estrella and Mishkin (1995, 1998), and one discrete variable representing realizations of Y_t with one lag (denoted by y_{t-1}), which modeled the dynamics, as in Dueker (1997). They also tried this model with different lags for the spread, which performed similarly well with slight lead by the model where the spread is at lag two. The model with the fourth lag was among the best, and so, they chose it as the benchmark case because of its ability to produce longer forecasts than models with lower lags, and so, we will focus on this particular version here. To be precise, this benchmark parametric dynamic linear probit model was given by²

$$P(Y_t = 1 | X_{t-4}, Y_{t-1}) = \Phi(f(X_{t-4}, Y_{t-1})) = \Phi(b_0 + b_1 X_{t-4} + b_2 Y_{t-1}), \quad (2.2)$$

where ϕ denotes the cumulative distribution function of the standard normal distribution. We therefore will focus on this parsimonious model. Besides replicating results for this model, our main goal is to compare them to results from the nonparametric quasi-likelihood approach, which was originated by Fan et al. (1995) and recently generalized further by Park et al. (2017) to allow for the dynamic time series context where the lagged values of the discrete dependent variable models the dynamics. Specifically, the nonparametric quasi-likelihood function from Park et al. (2017) adapted to the model (2.2), and evaluated at a point of interest (x_o, y_o) , is given by

$$\begin{aligned} \mathcal{L}_T(\beta_0, \beta_1 | x_o, y_o) &= \frac{1}{(T-4)} \sum_{t=5}^T Q\left(g^{-1}(\beta_0 + \beta_1(x_{t-4} - x_o)), y_t\right) \\ &\quad w_c(x_{t-4}, x_o) w_d(y_{t-1}, y_o), \end{aligned} \quad (2.3)$$

where $Q(\cdot, y_t)$ is quasi-likelihood function that takes the role of the likelihood of the mean when $Y_t = y_t$ is observed and g is a known link function (strictly increasing) and w_c^t and w_d^t are kernel weights for the continuous and discrete variables, respectively,

² We use the following convention for the notation: capital letter (e.g., Y_t) denotes an unobserved random variable in the time period t , while the fixed values it can or do take or points at which one want to estimate are denoted with small letters (e.g., y or y_t if the timing of the observation is to be emphasized).

where for the latter we use the ‘complete smoothing’ approach as described in Li et al. (2016).³

Because Y is binary, we have $P(Y_t = y | x_{t-4}, y_{t-1}) = m(x_{t-4}, y_{t-1})^y [1 - m(x_{t-4}, y_{t-1})]^{1-y}$, for $y = 0, 1$, and so, we can replace $Q(\mu, y)$ by

$$\ell(\mu, y) = y \log \left(\frac{\mu}{1 - \mu} \right) + \log(1 - \mu). \quad (2.4)$$

Here and below, we write $m = m_{t-1}$ for simplicity. Furthermore, taking *probit* link $g(t) = \Phi^{-1}(t)$, we get

$$\begin{aligned} \mathcal{L}_T(\beta_0, \beta_1 | x_o, y_o) = & \frac{1}{(T-4)} \sum_{t=5}^T \left[y_t \log \left(\frac{\Phi(\beta_0 + \beta_1(x_{t-4} - x_o))}{1 - \Phi(\beta_0 + \beta_1(x_{t-4} - x_o))} \right) \right. \\ & \left. + \log(1 - \Phi(\beta_0 + \beta_1(x_{t-4} - x_o))) \right] w_c(x_{t-4}, x_o) w_d(y_{t-1}, y_o). \end{aligned} \quad (2.5)$$

Maximization of (2.5) at a given point of interest (x_o, y_o) gives $\hat{\beta}_0(x_o, y_o)$ that serves as the estimator of f at (x_o, y_o) and also gives $\hat{\beta}_1(x_o, y_o)$ that serves as the estimator of the first partial derivatives of f at (x_o, y_o) . The conditional mean function $m(x_o, y_o)$ can then be estimated by inverting the link function evaluated at $\hat{\beta}_0(x_o, y_o)$, i.e., $g^{-1}(\hat{\beta}_0(x_o, y_o))$.⁴

In a nutshell, the main benefit of the nonparametric approach is that it allows much greater flexibility because instead of assuming a linear or any other particular parametric form for f , we approximate f locally in the direction of the continuous covariates and constant in the direction of the discrete covariates. As a result, we note that unlike in the parametric linear probit, $\hat{\beta}_1(x_o, y_o)$ generally can vary with (x_o, y_o) and thus is able to identify different impacts of explanatory variables onto the response variable. Why is this important in the context of forecasting recessions with the spread? As was also pointed out by Kauppi and Saikkonen (2008): ‘the interest rate spread can be regarded as a measure of the stance of monetary policy, while there is evidence that the impact of monetary policy on the real economy may be different during recessions and expansions.’ Continuing this line of thought, we note that the impact does not have to be the same (as in the linear model) even within one type of regime, be it ‘recession’ or ‘expansion’. Indeed, in the last few decades, society has been witnessing very different recessions and very different expansions and these differences may translate into different impacts of monetary policy, i.e., different values of the coefficient of x for different values of x .⁵ The nonparametric approach

³ Specifically, we note that $Q(\cdot, y_t)$ is defined by $\partial Q(\mu, y_t) / \partial \mu = (y_t - \mu) / V(\mu)$, where V is a chosen function for the working conditional variance model; see Park et al. (2017) and Li et al. (2016) for more of the technical details for the methodology used in this paper.

⁴ Here, it is also worth noting that although the implementation requires the choice of the link function (e.g., probit), the theory in Park et al. (2017) suggests that the asymptotic properties of the estimators are largely independent on the choice of g as long as it is sufficiently smooth and strictly increasing, because the estimation is performed locally.

⁵ For example, see a nice survey on related discussion by Florio (2004).

allows these variations. Indeed, the main value of the nonparametric approach is that it allows for the variation in the degree of influence of the explanatory variable onto the response variable, not just at different regimes ('recession' vs. 'expansion'), but at any point. Through the kernels and optimally selected bandwidths, it gives greater weight to the nearest observations, fitting the data in a way that minimizes the trade-off between the bias and variance, and without imposing a particular parametric form on the impact. The primary goal of this article is therefore to see if this method can improve upon the existing parametric approach.

3 Replication and nonparametric validation

We collected the data from the sources cited in Kauppi and Saikkonen (2008) for the same period as in their study: from 1955:Q4 to 2005:Q4. The dates of the peaks and troughs are according to NBER's Business Cycle Dating Committee, hereafter NBER.⁶ The information about the interest rates was sourced from the Federal Reserve Bank.⁷

It is important to note that different definitions of what is recession and what is the spread have been used in the literature to construct variables from this same data set, and some of these differences produce difference in results. Here, we try to follow exactly the definitions of Kauppi and Saikkonen (2008). Specifically, the dependent variable is constructed by setting $y_t = 1$ if 'US economy is considered as in recession' in the quarter t and 0 otherwise. Specifically, a given quarter is defined as the first quarter of a recession period if its first month or the preceding quarter's second or third month is classified by the NBER as the 'business cycle peak'; a given quarter is defined as the last quarter of a recession period if its second or third month or the subsequent quarter's first month is classified by the NBER as the 'business cycle trough'. The first and the last quarters define a recession period, during which all quarters are recession quarters (where $y_t = 1$). All the quarters that are not included in a recession period are called expansion quarters (where $y_t = 0$). Meanwhile, the spread variable is constructed as the difference between the 10-year US Treasury bond rate and the 3-month US Treasury bill rate.

The full sample here consists of 201 quarterly observations on US recessions and on the spread. Using these data, the parametric linear dynamic probit gave: $\hat{b}_0 = -1.15$, $\hat{b}_1 = -0.46$, $\hat{b}_2 = 1.88$ and $\text{Pseudo-}R^2 = 0.42$ —these are very similar as the estimates in Kauppi and Saikkonen (2008), $\hat{b}_1 = -0.55$ and $\hat{b}_2 = 1.79$ and $\text{Pseudo-}R^2 = 0.41$. (The small difference could be due to different optimization routines or/and slight differences in the data on recessions or on the spread.)

To perform the nonparametric estimation, we first need to obtain optimal bandwidths. For this, we use the leave-one-out maximum likelihood cross-validation criterion, maximizing it numerically over three dimensions. Here, recall that we use an adaptive bandwidth for the continuous variable which allows the bandwidth to be different according to the values of y_{t-1} . Intuitively, this means that we allow more

⁶ See <http://www.nber.org/cycles/>.

⁷ See <http://www.federalreserve.gov/releases/h15/data.htm>.

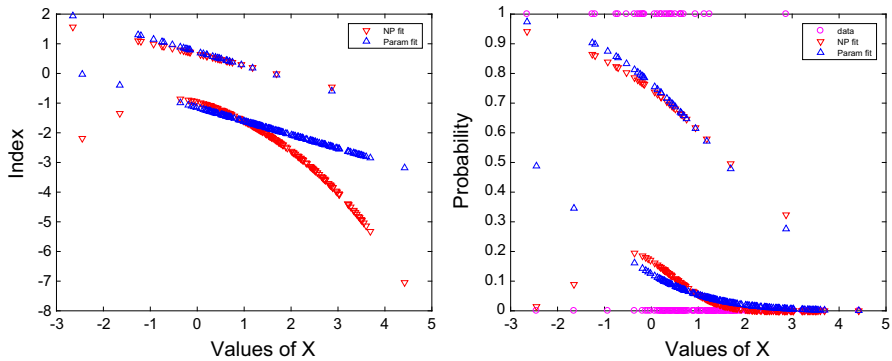


Fig. 1 Recessions in US example, from 1955:Q4 to 2005:Q4. Left panel estimates the index function $f(x, 0)$ and $f(x, 1)$, and right panel estimates the probabilities evaluated at observed points, as function of x_{t-4} . The two levels correspond to the realizations of either $y_{t-1} = 1$ (recession years, higher level) or $y_{t-1} = 0$ (no recessions years, lower level)

flexibility in fitting the curvatures under the two regimes, i.e., allowing for very different nonlinearities in the response function under recession and expansions. As we mentioned above, allowing for such flexibility is important for this context because the responsiveness of the real economy to the monetary policy (where the spread is the quintessence) was noticed to be usually different during recessions and expansions [see Florio (2004) and references therein]. The resulting optimal value of the bandwidth for the discrete regressor was $\lambda = 2.2/10^5$, confirming the importance of the lagged dependent variable y_{t-1} among explanatory variables in the index function. Importance of this variable can also be seen from the fact that the two optimal bandwidths for the continuous variable are very different: $h_0 = 1.092$ when $y_{t-1} = 0$, and $h_1 = 381.263$ when $y_{t-1} = 1$, suggesting about very different curvatures in the two regimes, as can be seen in Fig. 1.

More specifically, from Fig. 1, one can see that for the group of data where economy was in expansion in $t - 1$ (i.e., $y_{t-1} = 0$), the nonparametric estimator of the index function displays a clear curvature (left panel of Fig. 1).⁸ For the other case, the effect of the spread is roughly linear. The curvature appears again in the nonparametric probit probabilities (right panel of Fig. 1). This interesting finding confirms the claims and other evidence in the literature about the asymmetry of monetary policy effectiveness under different regimes (Florio 2004). Such a difference in the impact is not apparent from the dynamic linear probit approach results.

To compare the parametric and nonparametric fits, we follow Kauppi and Saikkonen (2008) and use the Pseudo- R^2 of Estrella (1998), which is a monotone transformation of the likelihood derived from some desired properties: We get 0.4187 for the para-

⁸ As correctly pointed out by a referee, the observed curvature could also be due to estimation uncertainty and so the statistical testing (yet to be developed in this framework) is needed to confirm or reject it with greater certainty.

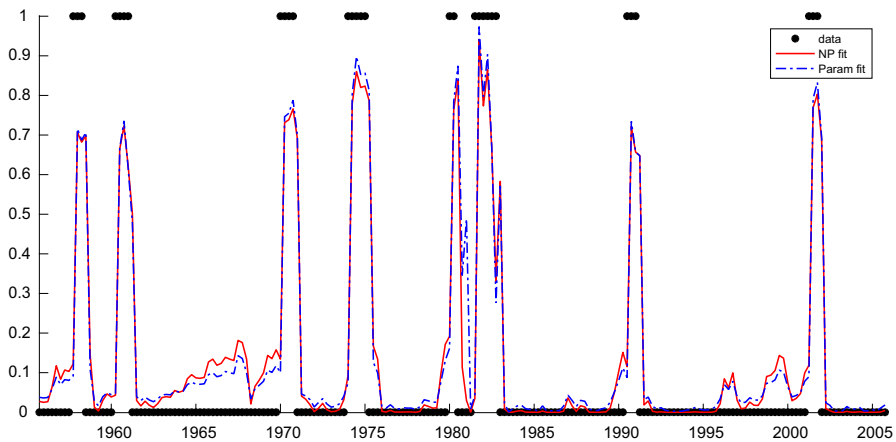


Fig. 2 Recessions in the USA for data as in Kauppi and Saikkonen (2008). In-sample forecasts from the linear dynamic probit (dot-dash line) and its nonparametric analogue (solid line). The ●'s are the realizations of Y_t (1 if recession and 0 otherwise)

metric and 0.4426 for the nonparametric cases.⁹ Figure 2 confirms the global quality of the in-sample fit for both parametric and nonparametric approaches.

We now proceed to the out-of-sample forecasts to make predictions of the recession one-period and two-periods ahead, where for the latter we follow the indirect iterative approach of Kauppi and Saikkonen (2008) and its nonparametric version.

We start with a subsample from 1955:Q4–2000:Q4 as the estimation sample and make one-period ahead and two-periods ahead forecasts, for 2001:Q1 and 2001:Q2, respectively. We then expand the sample by adding the next quarter (2001:Q1) to the estimation sample to produce new one-period ahead and two-periods ahead forecasts, for 2001:Q2 and 2001:Q3, respectively, and so on until we get the forecasts for the last observations in the sample (2005:Q3 and 2005:Q4).¹⁰ The results are then contrasted to the actual realizations and presented in Fig. 3.

Interestingly, note that the linear probit and the nonparametric approaches give fairly similar results for both types of out-of-sample predictions, especially with the one-period ahead forecasts. In particular, we note that the recession of early 2000s is almost identically warned by both models, as were most non-recessions.

The overall performance in the out-of-sample forecasts is also summarized in Table 1 via two goodness-of-fit measures, Pseudo- R^2 due to Estrella (1998) and Pseudo- R^2 due to Efron (1978). Both measures confirm very similar performance

⁹ Such a comparison appears reasonable at least heuristically, although rigorous theory needs to be developed to justify this or other comparisons of parametric and nonparametric fits in the context of dynamic discrete choice modeling and forecasting.

¹⁰ Here, the new optimal bandwidths were re-estimated with each addition to the sample, to account for the new information and for the new sample size.

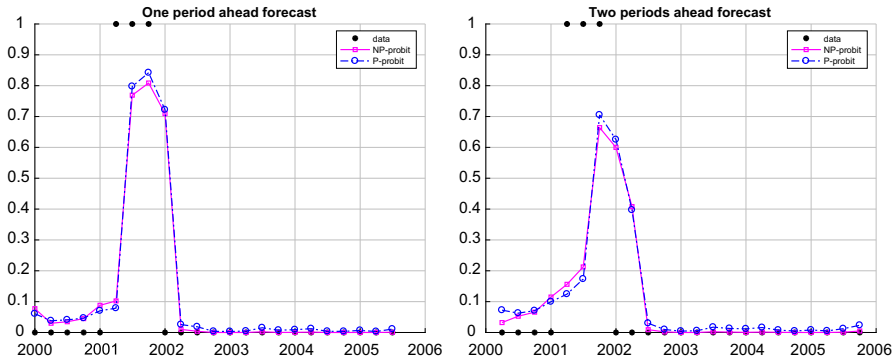


Fig. 3 Recessions in the USA for data as in Kauppi and Saikkonen (2008). One-period and two-periods ahead forecasts on left and right panels, respectively, were obtained for samples from 1955:Q4 to 1999:Q4 + t for $t = 0, 1, 2, \dots$. The $\bullet$'s are the actual realizations of Y_t (1 if recession and 0 otherwise)

Table 1 Goodness of fit for the out-of-sample forecasts

	Parametric probit	Nonparametric probit
Pseudo- R^2_{Estrella} for one-period ahead forecast	0.4075	0.4404
Pseudo- R^2_{Estrella} for two-periods ahead forecast	0.2624	0.3163
Pseudo- R^2_{Efron} for one-period ahead forecast	0.4488	0.4621
Pseudo- R^2_{Efron} for two-periods ahead forecast	0.2028	0.2464

The out-of-sample forecasts used in the goodness-of-fit measures were obtained for forecasts using samples from 1955:Q4 to 1999:Q4 + t for $t = 0, 1, 2, \dots$

of the two approaches, with the nonparametric approach giving slightly better forecasts, especially for the two-periods ahead forecasts (by a factor of about 1.2).¹¹

Overall, we can conclude that our nonparametric approach generally validates and supports (in terms of robustness of conclusions) the parametric approach used earlier by Kauppi and Saikkonen (2008), for this particular model and this data set.

4 Some insights with updated data

In this section, we summarize some insights that we are able to get with the expanded data set: We include the most recent periods as well as six earlier quarters, i.e., now we start from 1954:Q2 rather than 1955:Q4 and finish at 2017:Q4.

4.1 Could the GFC be warned earlier?

What would the two approaches have told us if we had tried them, for example, in mid-October 2007—would they have warned us about the looming GFC and the related

¹¹ Clearly, these values generally depend on the cutoff between the estimation subsample and the forecast validation, and so, we also varied the cutoff point and confirmed the general conclusion is robust, with the similarity increasing when the estimation sample is decreasing.

recession, often referred to as the ‘Great Recession’?¹² We try to answer this question in this subsection.

By mid-October 2007, the public already knew the spread for 2007:Q3, and since, this variable enters the model with 4 lags, one could use its actual values to get predictions for up to 2008:Q3 using the one-period ahead forecast. Here, recall that 2008:Q3 was the quarter when one of the biggest stock markets crashes in recent history, happened in September 2008, triggered by the collapse of one of the top five investment banks, Lehman Brothers, and rooted in the problems in the sub-prime mortgage market. Many people casually associate the Great Recession as starting with or near this particular quarter and this crash of the stock market and what followed after it. On the other hand, officially (according to the NBER approach), the peak of the US economic activity was reached in 2007:Q4 (in December), means that the recessionary period started already from 2008Q1. It is very important to clarify, however, that NBER announced that this was the case only as late as 1 December 2008, i.e., almost three months after the notorious market crash, and up to that point the public still wondered whether the US economy was in recession or not. While a few individuals have been warning about the approaching crash and the related recession since 2007 (and some did so for years prior to it), the mainstream economists at that time, before 2008:Q3, appear to have had no definite anticipation of them before they actually happened.¹³

Given this background, the forecasts for 2007:Q4–2008:Q3 are particularly interesting, answering the question whether any of the two approaches could have given any warning about the recession before the NBER’s *ex post* announcement (in December 2008), and especially before the market crash in September 2008.

The practical difficulty of applying the dynamic approach here is the substantial delay with which NBER announces their decisions on dating the recessions and expansions. As a result, in 2007:Q3, the information on realizations of Y_t was known with certainty only till November 2001—the latest turning point (trough) happened before 2007:Q3, and announced by NBER as late as July 2003. Therefore, while in 2007:Q3 the public knew the expansion started after November 2001 (i.e., from 2002:Q1 we have Y_t switching to 0 from 1), there was still uncertainty for how long that expansionary period would last.

Note that Kauppi and Saikkonen (2008) faced the same problem when performing their forecasting exercise in September 2006. To resolve the problem, they essentially relied on the expert judgment, stating that ‘[they] are confident (and the public is) that the U.S. economy was still in expansion in the last quarter of 2005...’ (p.783.), and so, they proceeded by setting $Y_t = 1$ from 2002:Q2 till 2005:Q4.

Here, we use a different approach to deal with this practical problem: We use our nonparametric approach to forecast the unknown realizations of Y_t , iteratively, from actual information known up to 2002:Q1 (the latest non-recessionary period known to the public). Specifically, we use the actual data from 1954:Q2–2002:Q1 for estimating via our nonparametric approach and then use our *one-period ahead* forecasts about

¹² e.g., see Temin (2010) for discussions about this period.

¹³ For example, the warnings about the looming crisis from Nouriel Roubini are often referred to September 2006, yet it appears they were not taken seriously by many before the collapse of Lehman Brothers, and even less so before another top five investment bank, Bear Sterns, collapsed in March 2008. For example, see the article ‘Dr. Doom’ , by S. Mihm in New York Times Magazine (15 August 2008).

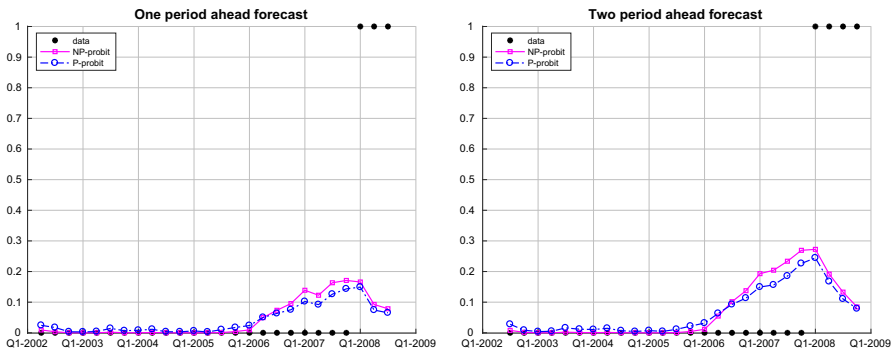


Fig. 4 One-period and two-periods ahead forecasts of US recessions on left and right panels, respectively, were obtained for samples from 1954:Q2 to 2001:Q1 + t for $t = 0, 1, 2, \dots$. The $\bullet$'s are the actual realizations of Y_t (1 if recession and 0 otherwise)

the probability of recession in 2002:Q2 to decide and impute the unknown value of the recession variable Y_t in 2002:Q2 to be 1 if the forecasted probability is greater or equal to 0.5, and set it to 0 otherwise.¹⁴ Such forecasting exercise is then repeated to forecast 2002:Q3 with the data from 1954:Q2 to 2002:Q2 where the unknown realization of Y_t in 2002:Q2 is replaced with our forecast obtained in the previous step. We roll forward such forecasting iterations till 2008:Q3, i.e., till the last actual observation on the most important predictor, the spread (with lag 4) for 2007:Q3 that was known to the public by mid-October 2007.¹⁵

The figures for the in-sample fit are very similar to those presented in the previous section (despite a larger sample), and so, we do not present them and rather focus on the out-of-sample forecasts, which are presented in Fig. 4. These results suggest that, as most of the mainstream economists at that time, both approaches could not definitely predict the Great Recession. However, they did give a useful warning that a recession might be approaching. Indeed, starting from 2006, the estimated probability of recession sharply risen from nearly zero—the level sustained for all quarters since the expansion started—to 0.05–0.1 about 1.5 years prior to the recession, with a further rise to 0.15–0.25 a few quarters before the recession.¹⁶ By the time, NBER announced the economy had been in the recession for about 1 year, and these estimated probabilities roughly halved, though still remained substantially above zero.

Looking at the previous recessions (e.g., see Fig. 2), one can see that a similar phenomenon (a jump in estimated probabilities from nearly zero to 0.10–0.20) also happened before several previous recessions (as well as at some other times). This suggests that the threshold for an alert of ‘substantial odds’ of a recession when

¹⁴ In principle, one can also use the two-periods ahead forecasts as those discussed in the previous section (and we did so and obtained the same conclusion) or construct k -periods ahead forecasts, or a combination of these approaches with various weights, or try different thresholds besides 0.5—the many possible exercises we leave for future explorations.

¹⁵ The new optimal bandwidths were re-estimated at each iteration to account for the new information and for the new sample size.

¹⁶ Interestingly, note that the two-periods ahead forecasts gave stronger warnings about the looming recession than the one-period ahead forecasts, as well as did so earlier.

considering these probability forecasting models perhaps should be set at much more conservative (lower) level than 0.5, which we used above. For example, one may search for an optimal (e.g., based on mean squared error criterion) threshold that historically was the best in correctly predicting recessions in out-of-sample forecasts—we leave this to be explored in future studies.

4.2 Some insights about 2018

Here, we do the same type of exercise as in the previous subsection with the aim to see what the two approaches forecast about 2018.

Again, note that when we started this exercise (mid-January 2018), the information on realizations of Y_t was known with certainty only until 2009:Q3—because the latest decision by NBER (announced in September 2010) was that the US economy reached a trough in June 2009, which in turn defined the next quarters as expansionary periods.¹⁷ Therefore, while we know the expansion started after June 2009 (and so 2009:Q3 is definitely counted as ‘non-recession’), there is still uncertainty for how long this expansionary period would last. Strictly speaking, until the NBER announces the time of the next peak, we cannot be 100% sure about the status of quarters after 2009:Q3, especially about the most recent quarters to date, though it is very likely that most of them until somewhere in 2016–2017 would be counted as expansionary.¹⁸

Again, one approach to resolve this practical challenge is to use expert judgment (similar as in Kauppi and Saikkonen (2008), as quoted above) and another is to use a model like ours to iteratively forecast the values from 2009:Q4, as we did in the previous subsection and also do here.

While we used substantially expanded sample, the figures for the in-sample fit are fairly similar to those we presented for the smaller sample in Sect. 3, and so, we do not present them here and focus our discussion on the out-of-sample forecast.

Figure 5 presents the one-period and two-periods ahead forecasts from both approaches, all showing very low estimated probability of recession (according to definition of NBER) up to 2018:Q4. Does this mean a recession in 2018 is very unlikely? Here, it is worth recalling remark of Kauppi and Saikkonen (2008), who pointed out that the recent US recessions were harder to predict than the previous ones, as we can also see from our results, and this prediction seems to have become even more challenging for the 2008–2009 recession, as we have seen in the previous section, and probably even more so for the next recession.

The lesson from the previous subsection also suggests that both approaches were able to warn about the looming latest recession only though a sharp yet still small rise in the estimated probability—from nearly zero (which was stable for a long time) to 0.05–0.1 about 1.5 years prior to the recession, with further rise to 0.15–0.25 a few

¹⁷ It might be also worth noting here that the NBER’s Business Cycle Dating Committee also met in April 2010 but without a decision, concluding: ‘Although most indicators have turned up, the committee decided that the determination of the trough date on the basis of current data would be premature...’.

¹⁸ On average, NBER’s decisions came with about 11 months delay (if counting since 1979 when they started the formal announcements of business cycle turning points). The longest delays, 20 and 19 months, were to announce the troughs of 1991 and 2001, yet there is no guarantee that the NBER does not take longer delay to announce the new peak that ends the latest expansion.

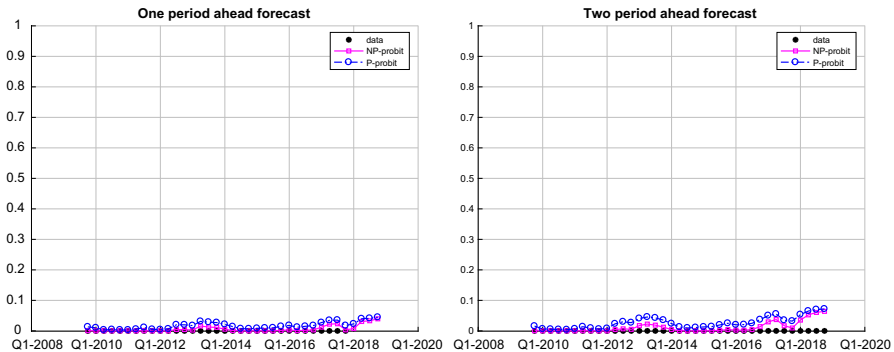


Fig. 5 One-period and two-periods ahead forecasts of US recessions on left and right panels, respectively, were obtained for samples from 1954:Q2 to 2009:Q3+ t for $t = 0, 1, 2, \dots$. The $\bullet$'s are the actual realizations of Y_t (1 if recession and 0 otherwise)

quarters before the recession. A similar phenomenon occurred before several other recessions too.

Thus, what Fig. 5 may hint us about is that we might be at the beginning of observing a similar pre-recession phenomenon now, at the beginning of 2018: a small yet sharp rise from nearly zero to about 0.05 in the estimated probabilities of observing a recession later in 2018. This is the largest jump since the end of the last recession (June 2009), with two other smaller jumps happening in 2013 and 2017. It might therefore be worthwhile to be on alert in the next quarters, monitoring if the estimated probability continues to rise as it did in 2008.

5 Concluding remarks

The goal of this article was threefold. Firstly, we wanted to replicate some of the results from the parametric (linear) dynamic probit of Kauppi and Saikkonen (2008) for forecasting the probability of US recessions. And, secondly, we wanted to perform a nonparametric validation of their interesting results. By achieving these two goals, we conclude that the results of Kauppi and Saikkonen (2008) are successfully replicated and, more interestingly, the nonparametric approach gave a slightly better in-sample and out-of-sample fit. This is especially the case for the two-period ahead forecasts, where for both approaches the accuracy of the forecasts fell substantially relative to the one-period ahead forecasts, but fell more so for the parametric approach, with the difference in terms of the Pseudo- R^2 measures reaching a factor of about 1.2.¹⁹

On the other hand, it would be also fair to conclude that the parametric dynamic linear probit approach of Kauppi and Saikkonen (2008) gave fairly good approximation of what a more computationally demanding nonparametric version would give for the same data, especially for the in-sample fit and for the one-period ahead forecast (where

¹⁹ As correctly pointed out by the referee, rigorous comparisons of the two approaches should also consider the measures of random variation and accuracy of estimates, which are yet to be developed for this framework.

the difference was within 10%). This conclusion is particularly important since the asymptotics for the inference for the nonparametric approach of Park et al. (2017) is yet to be developed.²⁰

Our third goal was a task inspired by one of the referees—to use both approaches on expanded data to gain insights on whether these approaches could have warned about the recession associated with the global financial crisis, if they were applied in 2007 and early 2008, as well as any insights about 2018. Our conclusion from this task is that both approaches were able to warn the latest recession only through a sharp yet still relatively small rise in the estimated probabilities of recession—from nearly zero to 0.05–0.1 about 1.5 years prior to the recession and with later rise to 0.15–0.25 closer to the recession. A similar phenomenon also occurred before several earlier recessions and appears to happen now in 2018. Considering this together with the fact that the latest non-recessionary period, nearly nine years, is one of the longest expansions of the US economy since 1854,²¹ it seems it would not be unreasonable to expect the next recession occurring within the next 2 years, unless the government implements substantial measures to mitigate it. All in all, further monitoring and a more thorough study of the probability of a recession, well beyond the modest goals of this paper, seem to be highly warranted in the near future.

Acknowledgements We thank Adrian Pagan, colleagues and participants of various events where earlier versions of this work was presented for their valuable feedback. We acknowledge and thank for the financial support provided by the ARC Discovery Grant (DP130101022) and ARC FT170100401, from the ‘Interuniversity Attraction Pole’, Phase VII (No. P7/06) of the Belgian Science Policy, and the National Research Foundation of Korea (NRF) grant funded by the Korea government (MSIP) (No. NRF-2014R1A4A1007895). We also thank Ms. Ailin Leng for her technical assistance with some data collection at the early stage of the project. Finally, we thank the editor and two anonymous referees for the fruitful feedback that helped improving this paper substantially.

References

- Aitchison J, Aitken CGG (1976) Multivariate binary discrimination by the kernel method. *Biometrika* 63(3):413–420
- Chu CY, Henderson DJ, Parmeter CF (2015) Plug-in bandwidth selection for kernel density estimation with discrete data. *Econometrics* 3(2):199–214
- Dueker M (1997) Strengthening the case for the yield curve as a predictor of U.S. recessions. *Rev Fed Reserve Bank St Louis* 79(2):41–51
- Efron B (1978) Regression and ANOVA with zero-one data: measures of residual variation. *J Am Stat Assoc* 73:113–121
- Estrella A (1998) A new measure of fit for equations with dichotomous dependent variables. *J Bus Econ Stat* 16(2):198–205
- Estrella A, Mishkin FS (1995) Predicting U.S. recessions: financial variables as leading indicators. Working paper 5379, National Bureau of Economic Research

²⁰ In principle, one can use Theorem 3.1 of Park et al. (2017) to construct a point-wise confidence intervals. However, this requires the estimation of more complex functions such as the second derivatives of the index function. Instead, one may apply a bootstrap method to construct confidence bands, though one also need to check whether the bootstrap method works theoretically and how good it performs in small samples.

²¹ According to NBER, the only two longer expansions in the USA since 1854 were those between March 1991 and March 2001 (120 months) and between February 1961 and December 1969 (106 months), while the average was only about 39 months overall (over 33 business cycles, since 1854) and about 58 months since 1945 (over 11 business cycles).

- Estrella A, Mishkin FS (1998) Predicting U.S. recessions: financial variables as leading indicators. *Rev Econ Stat* 80(1):45–61
- Fan J, Heckman NE, Wand MP (1995) Local polynomial kernel regression for generalized linear models and quasi-likelihood functions. *J Am Stat Assoc* 90:141–150
- Florio A (2004) The Asymmetric Effects of Monetary Policy. *J Econ Surv* 18:409–426
- Henderson DJ, Parmeter CF (2015) *Applied nonparametric econometrics*. Cambridge University Press, Cambridge
- Kauppi H, Saikkonen P (2008) Predicting U.S. recessions with dynamic binary response models. *Rev Econ Stat* 90(4):777–791
- Li D, Simar L, Zelenyuk V (2016) Generalized nonparametric smoothing with mixed discrete and continuous data. *Comput Stat Data Anal* 100:424–444
- Park BU, Simar L, Zelenyuk V (2017) Nonparametric estimation of dynamic discrete choice models for time series data. *Comput Stat Data Anal* 108:97–120
- Temin P (2010) The Great Recession and the Great Depression. National Bureau of Economic Research (NBER), Working paper no. 15645

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.