

PREDICTING RECESSIONS: A NEW MEASURE OF OUTPUT GAP AS PREDICTOR

Mastromarco, C., Simar, L. and V. Zelenyuk

DISCUSSION PAPER | 2019 / 23

Predicting Recessions: A New Measure of Output Gap as Predictor

Camilla Mastromarco* Léopold Simar[§] Valentin Zelenyuk^{††}

October 29, 2019

Abstract

In this paper, we merge two streams of literature: nonparametric methods to estimate frontier efficiency of an economy, which allows us to develop a new measure of output gap, and nonparametric methods to estimate probability of an economic recession. To illustrate the new framework we use quarterly data for Italy from 1995 to 2019, and find that our model, using either nonparametric or the linear probit model is able to provide useful insights.

JEL: C5, C14, C13, C32, D24, E37, O4,

Keywords: Output Gap, Robust Nonparametric Frontier, Generalized Nonparametric Quasi-Likelihood Method, Italian recession.

*Dipartimento di Scienze dell'Economia, Università degli Studi del Salento. Centro ECOTEKNE, Via per Monteroni - 73100 LECCE - Italy. Phone +39832298779, fax +390832298757, E-mail: camilla.mastromarco@unisalento.it.

[§]Institut de Statistique, Biostatistique et Sciences Actuarielles, Université Catholique de Louvain. Voie du Roman Pays 20 B - 1348 Louvain-La-Neuve - Belgium. Email: leopold.simar@uclouvain.be.

^{††}School of Economics and Centre for Efficiency and Productivity Analysis (CEPA) at The University of Queensland, Colin Clark Building (39), St Lucia, Brisbane, Qld 4072, Australia. Research support from his employer and from ARC grants (DP130101022 and FT170100401) is gratefully acknowledged.

1 Introduction

How to predict economic recessions of a country? This is a very important and challenging question interesting to a fairly wide audience. Many papers in the empirical macroeconomic literature proposed various methods to predict economic recessions, mainly focusing on the US. Here we follow one of the paradigms, started by Estrella and Mishkin (1995, 1998) further elaborated in various works (e.g., see Duecker 1997, Kauppi and Saikkonen 2008, and references cited therein), and we try to elaborate further by adapting some newly advanced methods in nonparametric statistics and in productivity and efficiency analysis.

Semiparametric and nonparametric methods are increasingly popular to analyze data in economics, business and other fields (e.g., see Horowitz 2009, Henderson and Parmeter 2015). Specifically, we use a nonparametric version of the dynamic probit for time series (Park et al. 2017) to model the dependent variable (recession vs. non-recession). Meanwhile, for the explanatory variables, besides the standard predictor such as the spread, we try to develop a method to incorporate the estimates of the efficiency scores of a country. For this purpose, we use the method of frontier estimation in nonparametric location-scale models (Florens et al. 2014) and robust conditional frontier methods (Cazals et al. 2002, Daraio and Simar 2005, Daouia and Gijbels 2011, Mastromarco and Simar 2018, etc.). We illustrate our approach on the case of the Italian economy.

In some sense, our paper is also related to and in the spirit of the work of Wheelock and Wilson (1995), who pioneered the use of efficiency estimates among predictors in the parametric probability models, in their case for predicting bank failures. Besides the focus on macroeconomic recessions rather than banks, the major distinctive features of our paper relative to theirs include (i) the use of recent nonparametric estimation methods for the discrete choice model (rather than a parametric one), (ii) the use of time-series data, with a dynamic component modeled explicitly and (iii) the use of more advanced methods for efficiency estimation that have become available very recently.

1.1 Predicting Recessions

Among the variety of different approaches attempting to model and forecast economic recessions, we will focus on those that employed the parametric binary choice approach and find that a good model for the prediction of the US recessions is a parsimonious model with only one of a few predictors, the most important of which is the interest rate spread and one discrete variable, the lagged dependent variable. The roots of this approach go back to at least

the seminal work of Estrella and Mishkin (1995, 1998), who thoroughly investigated various parametric models with many variables and concluded that the best forecasts resulted from a parsimonious probit model involving only one explanatory variable, the lagged spread. Duecker (1997) confirmed this result, yet also found that including the lagged dependent variable among regressors substantially improved the predicting power of the Estrella and Mishkin (1995, 1998) approach, especially for the recessions of the 1970s and 1990s that were missed by various other forecasting methods. Overall, the analyses in Estrella and Mishkin (1995, 1998) and Duecker (1997) suggest that their parsimonious model outperforms many alternative models that included many variables to gain a high in-sample fit, yet happened to be poorly forecasting the future. Also see Kauppi and Saikkonen (2008) for further refinements and more references and discussions.

This paper contributes to the empirical literature on predicting recessions by adding two novelties: i) we apply a nonparametric dynamic time series discrete response model suggested by Park et al. (2017); ii) we use a new measure of output gap as one of recession predictors. In particular, we employ a robust nonparametric frontier panel data model proposed by Mastromarco and Simar (2015) to estimate time-dependent conditional efficiency of countries and use this as a measure of output gap.¹ In a macroeconomics context, where countries are producers of output (i.e., GDP) given inputs (e.g., capital, labor) and technology, inefficiency can be identified as the distance of the individual production from the frontier. This frontier can be estimated by the maximum output of the reference country regarded as the empirical counterpart of an optimal boundary of the production set. Hence we might interpret the inefficiency as a measure of output gap with respect to the potential output of the technological frontier.

1.2 Existing Measures of Output Gap

Output gap is traditionally obtained as a deviation from a statistical measure of trend. One of the earliest and currently widely used statistical methods for measuring the output gap is based on measuring of output trend calculated by fitting a polynomial in time to output, the residual being the estimated cycle. This method imposes a strong prior on the smoothness of the trend. Another popular statistical approach uses a filter, Hodrick and Prescott (1997), to identify the trend and the cycle. The trend measure in this case is smooth but not deterministic. Baxter and King (1999) filter defines the cycle as having spectral power

¹Also see Cazals et al. (2002), Daraio and Simar (2005), Daouia and Gijbels (2011) for related discussions on robust nonparametric frontier.

in pre-specified frequencies. However, Murray (2003) stresses that this filter extracts an estimate of the cycle which includes some trend shock. Other statistical approaches need a model to identify the stochastic trend component. These statistical methods do not require smoothness but impose the restriction of no correlation between the cycle and the trend, which may lack theoretical support. Beveridge and Nelson (1981) suggest a measure of trend as a long run forecast of an ARMA model. The unobserved components model extracts an estimate of the trend and cycle using Kalman filter (Harvey 1985, Watson 1986, Clark 1987).

Differently from the statistical methods, the economic approaches estimate the output gap in the framework of production function (for example Galí and Gertler 1999). Recently, various studies (Kuttner 1994, Gerlach and Smets 1999, Apel and Jansson 1999, Roberts 2001, Basistha and Nelson 2007, Basu and Fernald 2009) tried to combine the statistical approach with the economic approach by estimating unobserved components multivariate model. These approaches do not impose smoothness or restrictive correlation structure, but estimate the output gap based on the empirical implications of the forward-looking Phillips curve.

1.3 Inefficiency as an Alternative Measure of Output Gap

Often, potential output is referred to as the production capacity of the economy. In our framework of the frontier model, potential output refers to the maximum level of output that can be produced for a given level of inputs, using full employment and capital utilization. The gap between the potential and actual outputs is interpreted as a measure of inefficiency which in our paper captures also the varying factor utilization over the cycle. The approach is closely linked to the production theory based approach in measuring output gap. We cast our empirical model in frontier form, treating the gap as an unobserved variable - efficiency scores - estimated using nonparametric frontier methods. In pursuing an economic based approach, we avoid imposing strong priors on the smoothness of the trend or cycle, and restrictive correlation structure between the trend and the cycle shocks.

Furthermore, parametric modelling may suffer from misspecification problems when the data generating process is unknown, as is usual in the applied studies. We propose a unified nonparametric framework for accommodating simultaneously the problem of model specification uncertainty and time dependence in panel data frontier model. Specifically, we estimate panel data frontier model using a flexible nonparametric two step approach to take into account the time dependence. Following recent development in nonparametric conditional frontier literature (Florens et al. 2014, Mastromarco and Simar 2015, 2018), we adapt the

nonparametric location-scale frontier model, where we link production inputs and output to time. In the first step we clean the dependence of inputs and outputs on time factors. These time factors capture the correlation among units. By eliminating the effect of these factors on production process we mitigate the problem of dependence across our time units and we are able to estimate a nonparametric frontier model from panel data. (In the application illustrate this approach for the data on 16 OECD countries.) In the second step we estimate the frontier and the efficiency scores using inputs and outputs whitened from the influence of time.

1.4 The Contribution in a Nutshell and a Roadmap

The main idea of this paper is to merge the interesting streams of literature described above: the novel nonparametric methods to estimate frontier efficiency of an economy, as a new measure of output gap, and the novel nonparametric method to estimate probability of economic recession. We do this by deploying a generalised nonparametric quasi-likelihood method in the context of dynamic discrete choice models for time series data (Park et al. 2017). To illustrate the new framework we use data from 1995 to 2019, with quarterly frequency, and find that our model using either nonparametric or the linear probit model, applied frequently in this context, is able to offer additional insights into the literature.

The paper is organized as follows. Section 2 presents the methodology. Specifically, Section 2.1. explains nonparametric discrete choice models for time series to predict recessions. Section 2.2. introduces our proposed measure of output gap and explains time-dependent conditional efficiency scores and the nonparametric estimation. This section elucidates the location-scale models to eliminate the influence of common time factors and external variables. Section 3 illustrates the empirical application and summarizes the main findings of the paper. Section 4 gives concluding remarks.

2 Methodology

2.1 Forecasting Model

In this section, we summarize the elements from Park et al. (2017) (hereafter PSZ) that are needed in our setup to forecast economic recessions. The model should provide the elements for analyzing the behavior of a discrete variable in a time series setup. The approach is nonparametric.

Suppose we observe $(\mathbf{X}^t, \mathbf{Z}^t, Y^t)$, $t = 1, \dots, T$, where $\{(\mathbf{X}^t, \mathbf{Z}^t, Y^t)\}_{t=-\infty}^{\infty}$ is a stationary random process. We assume, as in PSZ that the process satisfies strong mixing conditions that typically allows time dependence which disappear at a geometrical rate when the time lags are too large.²

The response variable is binary taking the values 0 and 1; in our set-up, $Y = 1$ corresponds to a recession and $Y = 0$ is otherwise. The vector of covariates $\mathbf{X}^t$ is of dimension r and of continuous type, whereas $\mathbf{Z}^t$ is a discrete vector of dimension k . The components of $\mathbf{Z}^t$ may be lagged values of the response Y , e.g., Y^{t-1} , Y^{t-2} . The idea is to estimate the mean function

$$m(\mathbf{x}, \mathbf{z}) = \mathbb{E}(Y | \mathbf{X} = \mathbf{x}, \mathbf{Z} = \mathbf{z}). \quad (2.1)$$

Since Y is binary we have

$$\mathbb{P}(Y = y | \mathbf{X} = \mathbf{x}, \mathbf{Z} = \mathbf{z}) = m(\mathbf{x}, \mathbf{z})^y [1 - m(\mathbf{x}, \mathbf{z})]^{1-y}, \text{ for } y \in \{0, 1\}. \quad (2.2)$$

A key ingredient in these discrete choice models is the link function g , which is a strictly increasing function, defining the function f as

$$f(\mathbf{x}, \mathbf{z}) = g(m(\mathbf{x}, \mathbf{z})). \quad (2.3)$$

In parametric models, it is assumed that $f(\mathbf{x}, \mathbf{z})$ takes a parametric form, and then $m(\mathbf{x}, \mathbf{z}) = g^{-1}(f(\mathbf{x}, \mathbf{z}))$. Thus, a wrong choice may jeopardize the estimation of m . In nonparametric settings, $f(\mathbf{x}, \mathbf{z})$ will be locally approximated by some local polynomial around $(\mathbf{x}, \mathbf{z})$, so the choice of g is much less important. Approximating locally the functions $g_1(m(\mathbf{x}, \mathbf{z}))$ or $g_2(m(\mathbf{x}, \mathbf{z}))$ for two different link functions g_1 and g_2 does not make much difference. One may simply take the identity function, though since the range of the target m is $[0, 1]$, we will choose a link that guarantees the correct range (like Probit or Logit). Now, given the link g and the sample $\{(\mathbf{X}^t, \mathbf{Z}^t, Y^t)\}_{t=1}^T$, we see from (2.2) that the log-likelihood of f is given by $\sum_{t=1}^T \ell(g^{-1}(f(\mathbf{X}^t, \mathbf{Z}^t)), Y^t)$ where $\ell(\mu, y) = y \log\left(\frac{\mu}{1-\mu}\right) + \log(1 - \mu)$.

Let $(\mathbf{x}, \mathbf{z})$ be a fixed point of interest at which we want to estimate the value of the mean function m , or equivalently of its transformed function f . In a nonparametric approach, we will apply local smoothing techniques to the observations $(\mathbf{X}^t, \mathbf{Z}^t)$, which are in the neighborhood of $(\mathbf{x}, \mathbf{z})$. As explained in PSZ, this leads to weighting the observation $(\mathbf{X}^t, \mathbf{Z}^t)$ near $(\mathbf{x}, \mathbf{z})$ by some kernel. For the continuous variables (X) , usual continuous kernels

²See PSZ, Section 3.1, for mathematical details and additional references.

(Gaussian, Epanechnikov, etc.) can be used, while for the discrete variables (Z), some appropriate discrete kernels have to be used. Here we use the product kernel $w_c^t(\mathbf{x}, \mathbf{z}) \times w_d^t(\mathbf{z})$ defined as

$$w_c^t(\mathbf{x}, \mathbf{z}) = \prod_{j=1}^r K_{h_j}(x_j, X_j^t, \mathbf{z}). \quad (2.4)$$

$$w_d^t(\mathbf{z}) = \prod_{l=1}^k \gamma_l \mathbb{I}(Z_l^t \neq z_l) \quad (2.5)$$

where $\mathbb{I}(A)$ denotes the indicator function such that $\mathbb{I}(A) = 1$ if A hold and zero otherwise and $\gamma_l \in [0, 1]$ is the bandwidth for the l^{th} discrete variable, while for the continuous kernels, we have

$$\begin{aligned} K_{h_j}(x_j, X_j^t, \mathbf{z}) &= \frac{1}{h_j(1)} K\left(\frac{X_j^t - x_j}{h_j(1)}\right) \times \mathbb{I}(Z^t = z(1)) \\ &+ \frac{1}{h_j(2)} K\left(\frac{X_j^t - x_j}{h_j(2)}\right) \times \mathbb{I}(Z^t = z(2)) \end{aligned}$$

for a symmetric kernel function K and two bandwidth, $h_j(1) > 0$ and $h_j(2) > 0$, corresponding to the two groups denoted as $z(1)$ and $z(2)$, for each j^{th} continuous variable. The discrete kernel is in the spirit of Aitchison and Aitken (1976), except that it is standardized to be between 0 and 1. The continuous kernel is a generalized kernel proposed by Li et al. (2016), which allows different bandwidths for the continuous variables across various groups defined by the values of $\mathbf{Z}$, and thus allowing for more flexibility in terms of the fitted curvatures in the two groups. It is worth noting that when $\gamma_j = 0$ one performs separate estimation for each group identified by the values of Z_j . When $\gamma_j = 1$, one considers that Z_j is irrelevant and so all the groups are pooled together, although different bandwidths for continuous variables may still imply different curvatures in the two groups. In practice, these bandwidths will be determined by likelihood-based cross-validation adapting the ideas from PSZ, Section 4.

For approximating $f(\cdot, \cdot)$ locally near the point $(\mathbf{x}, \mathbf{z})$ we will not make use of the link function, nor of the likelihood function. The local approximation is linear in the direction of the continuous variable and constant in the direction of the discrete variables. To be specific, we have

$$f(\mathbf{u}, \mathbf{v}) \approx f(\mathbf{x}, \mathbf{z}) + \sum_{j=1}^r f_j(\mathbf{x}, \mathbf{z})(u_j - x_j), \quad (2.6)$$

where $f_j(\mathbf{x}, \mathbf{z}) = \partial f(\mathbf{x}, \mathbf{z}) / \partial x_j$. So the local approximation can be viewed as a first order

Taylor's expansion of f in $\mathbf{x}$, near $(\mathbf{x}, \mathbf{z})$.

To estimate $f(\mathbf{x}, \mathbf{z})$ and its partial derivatives $f_j(\mathbf{x}, \mathbf{z})$ we thus maximize

$$T^{-1} \sum_{t=1}^T w_c^t(\mathbf{x}, \mathbf{z}) w_d^t(\mathbf{z}) \ell \left(g^{-1} \left(\beta_0 + \sum_{j=1}^r \beta_j (X_j^t - x_j) \right), Y^t \right) \quad (2.7)$$

with respect to β_0 and β_j , $j = 1, \dots, r$. The solutions $\hat{\beta}_0 = \hat{f}(\mathbf{x}, \mathbf{z})$ and $\hat{\beta}_j = \hat{f}_j(\mathbf{x}, \mathbf{z})$ for $j = 1, \dots, r$. Then an estimator of the mean function $m(\mathbf{x}, \mathbf{z})$ is obtained by inverting the link function: $\hat{m}(\mathbf{x}, \mathbf{z}) = g^{-1}(\hat{\beta}_0)$.

The theory in PSZ shows that the asymptotic properties of the estimators does not much depend on the choice of the link function, as long it is smooth enough and strictly increasing, because the estimation is performed locally. We will choose below the probit link, i.e. $g(s) = \Phi^{-1}(s)$, where Φ is the cumulative distribution function of the standard normal distribution. So we have to maximize in (β_0, β_j) , $j = 1, \dots, r$

$$T^{-1} \sum_{t=1}^T w_c^t(\mathbf{x}, \mathbf{z}) w_d^t(\mathbf{z}) \left[Y^t \log \left(\frac{\Phi \left(\beta_0 + \sum_{j=1}^r \beta_j (X_j^t - x_j) \right)}{1 - \Phi \left(\beta_0 + \sum_{j=1}^r \beta_j (X_j^t - x_j) \right)} \right) + \log \left(1 - \Phi \left(\beta_0 + \sum_{j=1}^r \beta_j (X_j^t - x_j) \right) \right) \right]. \quad (2.8)$$

The properties of the resulting estimators follow from PSZ. In summary, under certain regularity assumptions and with the optimal order of the bandwidths, $h_{c,j} := (h_j(1) + h_j(2))/2 \propto T^{-1/(r+4)}$ and $\gamma_j \propto T^{-2/(k+4)}$, Theorem 3.1 in PSZ establishes

$$\sqrt{T \bar{h}_c} \left(\hat{f}(\mathbf{x}, \mathbf{z}) - f(\mathbf{x}, \mathbf{z}) + \sum_{j=1}^r O(h_{c,j}^2) + \sum_{j=1}^k O(\gamma_j) \right) \xrightarrow{\mathcal{L}} N(0, V(\mathbf{x}, \mathbf{z})), \quad (2.9)$$

where $\bar{h}_c = \prod_{j=1}^r h_{c,j}$ and the variance V has a complicated expression which depends on properties of the data generation process (DGP) (see PSZ for details). We see from (2.9) that the optimal bandwidths balance, as often the case, between the square of the bias terms and the variance.

Remark 1: It is worth noting that if the bandwidths for continuous variables increase such that they cover all the observations on those variables, the nonparametric approach yields very similar estimates as the parametric approach that assumes (2.6) holds exactly. In this sense, the parametric approach can be viewed as a special case of the nonparametric

approach, in the sense that the latter allows for much more flexibility and can be ‘reduced’ to the former by removing the flexibility through tuning the bandwidths to be large enough. Importantly, if the nonparametric approach employs a suitable data-driven procedure to select all the bandwidths that fit the data in an optimal way (e.g., by optimally balancing the bias and the variance), then, in a sense, it lets the data speak for itself, fitting the curvature that appears as the most appropriate for the given data. Here, for the bandwidths selection we use leave-one-out maximum likelihood cross-validation (MLCV) method as explained in PSZ, although other suitable methods can also be used here.

Remark 2: the nonparametric approach can also be viewed as a tool for validation of a suitable parametric approach. Indeed, when a parametric approach that assumes a particular (and perhaps very restrictive) functional form yields very similar results or conclusions as the nonparametric approach that allows for much more flexibility, this should give more confidence in the results or conclusions from the parametric approach, despite its restrictive assumptions. We will find this consideration very useful in our empirical application section for the particular data we use there.

2.2 Efficiency and Estimation of the Output Gaps

We propose as output gap our measure of inefficiency. The output gap is an economic measure of the difference between the actual output of an economy and its potential output. Potential output is the maximum amount of goods and services an economy can turn out when it is most efficient—that is, at full capacity. Often, potential output is referred to as the production capacity of the economy. In the context of this paper, we assume that a country is producer of output (i.e., GDP) given inputs (e.g., capital, labor) and available technology. The inefficiency is defined as the distance between the actual production and its maximum or frontier potential, given the inputs and technology.

As explained above, we would like to use the level of inefficiency of the country for a particular year by considering the so-called conditional inefficiency (Cazals et al. 2002, Daraio and Simar 2005, Mastromarco and Simar 2015). Inputs here are Capital (K) and Labor (L) and the output is the GDP (Q), and we have quarterly data $t = 1, \dots, T$ for 16 OECD countries. Evaluating marginal efficiency measures by considering the so-called meta-frontier of the 3-dimensional cloud of T points $\{(K_t, L_t, Q_t)\}_{t=1}^T$ would not make too much sense since the technology certainly varies over years. We will rather consider the conditional efficiency measure where we condition on the time period. As suggested in Mastromarco and Simar (2015), to introduce the time dimension we consider indeed, with

some abuse of notation, time as a conditioning variable W and we define the attainable production set at time t as the support of the conditional probability

$$H_{K,L,Q|W}(\xi, \zeta, \eta | W = t) = \text{Prob}(K \leq \xi, L \leq \zeta, Q \geq \eta | W = t), \quad (2.10)$$

which can be interpreted as the probability of observing, at time t , a production plan dominating a given point (ξ, ζ, η) . So, the feasible technology Ψ^t can be defined as

$$\Psi^t = \{(\xi, \zeta, \eta) \in \mathbb{R}_+^3 | H_{K,L,Q|W}(\xi, \zeta, \eta | W = t) > 0\}. \quad (2.11)$$

Finally, this leads to consider for the output orientation the conditional efficiency score

$$\lambda(\xi, \zeta, \eta | t) = \sup\{\lambda | (\xi, \zeta, \lambda \eta) \in \Psi^t\} \geq 1, \quad (2.12)$$

which is known as the Farrell-Debreu output oriented efficiency measure (see e.g. Kumar and Russell 2002, for its use in a related context but using a simpler estimator). Nonparametric estimators of these efficiency scores have been developed and their asymptotic properties are well-known (see e.g. Jeong et al. 2010). Here, we will follow the approach suggested by Florens et al. (2014) which has some advantages described below.

In the first step, a flexible nonparametric model is used to whiten the inputs (K, L) and the output Q from the effect of time W . We have the following model

$$\begin{aligned} K_{it} &= \mu_K(t) + \sigma_K(t)\varepsilon_{K,t} \\ L_{it} &= \mu_L(t) + \sigma_L(t)\varepsilon_{L,t} \\ Q_{it} &= \mu_Q(t) + \sigma_Q(t)\varepsilon_{Q,t}, \end{aligned} \quad (2.13)$$

where we assume that $(\varepsilon_K, \varepsilon_L, \varepsilon_Q)$ are ‘independent’ of time W , with $\mathbb{E}[\varepsilon_\ell] = 0$ and $\mathbb{V}[\varepsilon_\ell] = 1$ for $\ell = K, L, Q$. The estimation of the mean and variance functions are done by local polynomial smoothing as explained in detail in Florens et al. (2014). They suggest also a bootstrap test for testing the assumption of independence, but in our application below we will evaluate various correlations (Spearman, Pearson and Kendall) to check if this assumption is reasonable.

In our application, we first use the local-linear methods to estimate the mean functions $\mu_\ell(t)$, $\ell = K, L, Q$. From the squared residuals we estimate the variance functions $\sigma_\ell^2(t)$ by local constant methods (to avoid negative variances). Finally, Florens et al. (2014) define

the estimated ‘pure’ inputs and the estimated ‘pure’ outputs as

$$\begin{aligned}\widehat{\varepsilon}_{K,it} &= \frac{K_{it} - \widehat{\mu}_K(t)}{\widehat{\sigma}_K(t)}, \\ \widehat{\varepsilon}_{L,it} &= \frac{L_{it} - \widehat{\mu}_L(t)}{\widehat{\sigma}_L(t)}, \\ \widehat{\varepsilon}_{Q,it} &= \frac{Q_{it} - \widehat{\mu}_Q(t)}{\widehat{\sigma}_Q(t)},\end{aligned}\tag{2.14}$$

which are ‘pure’ in the sense of being filtered from time dependence. In this ‘pure units space’, we can compute the output directional distance to the efficient frontier.³ Since the output here is univariate, the efficient frontier in pure units is the function

$$\varphi(e_K, e_L) = \sup\{e_Q | \text{Prob}(\varepsilon_K \leq e_K, \varepsilon_L \leq e_L, \varepsilon_Q \geq e_Q) > 0\},\tag{2.15}$$

so that the directional distance of a point (e_K, e_L, e_Q) to the frontier is simply given by

$$\delta(e_K, e_L, e_Q) = \varphi(e_K, e_L) - e_Q \geq 0,\tag{2.16}$$

where the value zero indicates the point (e_K, e_L, e_Q) is on the efficient frontier. Under the location-scale assumptions, it can be proven that the conditional frontier in original units can be recovered as (see Florens et al., 2014, for details)

$$\tau(\xi, \zeta | t) = \mu_Q(t) + \sigma_Q(t) \varphi\left(\frac{\xi - \mu_K(t)}{\sigma_K(t)}, \frac{\zeta - \mu_L(t)}{\sigma_L(t)}\right),\tag{2.17}$$

so that the gap in the output to reach the frontier level is given by

$$G_Q(\xi, \zeta, \eta | t) = \sigma_Q(t) \delta(e_K, e_L, e_Q).\tag{2.18}$$

The nonparametric estimators of these various elements are obtained by plugging the estimators of the mean and variance functions derived above. One of the main advantages of this location-scale approach is that for estimating the functions $(\mu_\ell(t), \sigma_\ell(t))$ we require only smoothing in the center of the data in a standard regression setup. As pointed in Bădin et al. (2019), a direct estimation of $\lambda(\xi, \zeta, \eta | t)$ requires delicate problems of optimal bandwidths selection for estimating the support of the conditional $H_{K,L,Q|W}(\xi, \zeta, \eta | W = t)$.

³We need to use directional distance here since the ‘pure’ inputs and the ‘pure’ outputs may take negative values.

So at the end of this step of efficiency estimations we end up in practice with estimated efficiency scores in the pure units $\delta(e_{K_t}, e_{L_t}, e_{Q_t})$ and, if wanted, the measures of the gaps in original units of the DGP, i.e., $G_Q(K_t, L_t, Q_t|t)$ at each observation $t = 1, \dots, T$. These values (eventually lagged) will be used to improve the prediction of recession in our application below.

Real data samples contain in general some anomalous data and the estimated frontier obtained by these nonparametric techniques can be fully determined by these outliers or extreme data points, jeopardizing the measurement of inefficiencies, potentially leading to unrealistic results. Cazals et al. (2002), Daouia and Simar (2007), in the frontier literature, propose an approach which aims to keep all the observations in the sample but to replace the frontier of the empirical distribution by (conditional) quantiles or by the expectation of the minimum (or maximum) of a sub-sample of the data. This latter method defines the order- m frontier that we will use here.

To be short, the partial output-frontier of order- m is defined for any integer m and for input values e_{K_t}, e_{L_t} , as the expected value of the maximum of the output of m units drawn at random from the populations of units such that $\varepsilon_K \leq e_K, \varepsilon_L \leq e_L$. Formally,

$$\tau_m(\xi, \zeta|t) = E [\max(\varepsilon_{Q,1t}, \dots, \varepsilon_{Q,mt})], \quad (2.19)$$

where the $\varepsilon_{Q,it}$ are drawn from the empirical conditional survival function $\widehat{S}_{\varepsilon_Q|\varepsilon_x}(e_Q|\widehat{\varepsilon}_{x,it} \leq e_x)$. This can be computed by Monte-Carlo approximation or by solving a univariate numerical integral (for practical details see Simar and Vanhems 2012).

If m increases and converges to ∞ and $n \rightarrow \infty$, it has been shown (see Cazals et al. 2002) that the order- m frontier and its estimator converge to the full frontier, but for a finite m , the frontier will not envelop all the data points and so is much more robust than the Free Disposal Hull (FDH) to outliers and extreme data points (see e.g. Daouia and Gijbels 2011, for the analysis of these estimators from a theory of robustness perspective). Another advantage of these estimators is that besides the fact that their limiting distribution is normal, they achieve the parametric rate of convergence $(\sqrt{n})$.

3 Empirical Illustration: The Case of Modern Italy

3.1 Data in Brief

The dataset for which we try our approach to estimate the output gap consists of 99 quarterly observations from (1995 : Q1) till (2019 : Q2), on capital, labor and output of 16 OECD countries (Austria, Belgium, Denmark, Finland, France, Germany, Ireland, Israel, Italy, Korea, Netherlands, New Zealand, Norway, Spain, Sweden and United Kingdom). The data is from Organisation of Economic Development (OECD): Quarterly National Accounts (OECD stat.).⁴ As often suggested with macroeconomic data, all the variables are transformed in logarithms before estimation. The output, gross domestic product (GDP), is measured in million US dollar at 2015 constant price. For labour input, we use number of employed persons (in thousands) seasonally adjusted. Capital K is measured in million US dollar at 2015 constant price and constructed applying the perpetual inventory method (PIM) by using the real investment series (gross fixed capital formation).⁵

The spread variable is constructed as the difference between the 10-year Germany Treasury bond rate and the 10-year Italy Treasury bond rate in per cent per annum and is sourced from OECD.stat Monthly Monetary and Financial Statistics (MEI).⁶ We use this measure of spread because 10 year yield German bonds are the benchmark for the Euro area since they are considered by investors a risk-free market asset.⁷

The variables on recessions are constructed as following. We use the Composite Leading Indicators from OECD Reference Turning Points and Component Series data.⁸ The OECD identifies months of turning points. Our time series is composed of dummy variables that represent periods of expansion and recession. A value of 1 is a recessionary period, while a value of 0 is an expansionary period. For this time series, the recession begins on the quarter of the month of the peak and ends on the quarter of the month of the trough.

⁴The choice of the variables depends on data availability at quarterly frequency. See Appendix A for data description.

⁵PIM is necessitated by the lack of capital stock data across all the countries. The capital stock is constructed as $K_t = K_{t-1}(1 - \theta) + I_t$, where I_t is investment and θ the rate of depreciation assumed to be 6% (e.g., Hall and Jones, 1999; Iyer *et al.*, 2008). Repair and maintenance are assumed to keep the physical production capabilities of an asset constant during its lifetime. Initial capital stocks are constructed, assuming that capital and output grow at the same rate. Specifically, for country with investment data beginning in 1995, we set the initial stock, $K_{1995} = I_{1995} / (g + \theta)$, where g is output growth rate from 1995 to 2019. Estimated capital stock includes both residential and non-residential capital.

⁶The observation period is selected by the data availability.

⁷See *The Economist* <https://www.economist.com/blogs/buttonwood/2014/03/investing>.

⁸Which can be found at:

www.oecd.org/sdd/leading-indicators/oecdcompositeleadingindicatorsreferenceturningpointsandcomponentseries.htm

The dependent variable is constructed by setting $Y_t = 1$ if “recession” in the quarter t and 0 otherwise. We then fit the models to forecast the probability of being in recession in a given quarter using information from previous quarters.

3.2 Filtering the Inputs/Output and Efficiency Estimates

We have to run 3 location-scale models for K, L, Q respectively to clean the effect of time W .⁹ This provides the ‘pure’ inputs and ‘pure’ output, $\{(\hat{\varepsilon}_{K_t}, \hat{\varepsilon}_{L_t}, \hat{\varepsilon}_{Q_t})\}_{t=1}^T$ as explained above. The correlations of these ‘pure’ inputs/output with time are given in Table (1) (where $X1 = K, X2 = L, Y = Q$ and $Z = W$). Clearly these correlations are very small so we can infer that the assumption of independence between $(\varepsilon_K, \varepsilon_L, \varepsilon_Q)$ and W , which is part of our location-scale model, seems reasonable.

Robust measures of efficiency scores, providing the gaps in ‘pure’ units were computed with $m = 1500$. This choice was done for letting less than 5% of points above the order- m frontier, as shown in Figure 1. Note that from the values of $m = 1500$ and $m \rightarrow \infty$ (the full FDH frontier), all the results are quite similar.

The resulting efficiency scores $\hat{\delta}_{m,t}$ are shown in Figure 2, which illustrates that most of the time, the time effect has indeed been cleaned from the production process. We see also that most of the values of $\hat{\delta}_{m,t}$ are positive and only some take very small (near zero) negative values. Figure 3 exhibits the time path of output gaps in original units (in logs and re-scaled by their mean).

We give in the Appendix B the full table of results for all the time periods. The table also indicates the gaps G_t in original units of the DGP, Q_t , as defined above (in log scale and re-scaled by their mean).

3.3 Fitting the Model

We apply our prediction model described above and estimate parametric linear probit model and our nonparametric model of PSZ. For the latter we use the complete smoothing technique suggested by Li et al. (2016), allowing different bandwidths for the continuous variables in the two groups determined by the values of Z .

As suggested in the literature for measuring the quality of the fit we indicate the value of the achieved Maximum Likelihood and of the Pseudo R^2 (see PSZ and the references

⁹For numerical convenience, all variables are scaled by their means, including the conditioning variable time, denoted here by W .

therein). The best fit we obtained for all the approaches were provided by the following model:

$$\mathbb{E}(Y|\mathbf{X} = \mathbf{x}, \mathbf{Z} = \mathbf{z}) = m(\mathbf{x}, \mathbf{z}) = m(X_1, X_2, Z), \quad (3.20)$$

where $X_{1,t} = Sp_{t-1}$ is the spread lagged by one period, $X_{2,t} = \Delta_{G,t-2}$ is the first difference of the estimates of output gaps (production efficiency) and lagged by two periods, i.e. $\Delta_{G,t} = G_{Q,t} - G_{Q,t-1}$. Finally, $Z_t = Y_{t-1}$. Our dependent variable is $Y_t = 1$ if “Italian economy is in recession” in the quarter t and 0 otherwise.

Our preferred model assumes that, the first difference of our measure of output gap, i.e. the variation, affects the probability of an economy to be in recession with some delays, due to market imperfections and frictions. Hence it acts as an indicator of recession two periods in advance, differently from the other indicator, the spread, usually used to forecast the recession, which in our case indicates only one period before a recession.

The estimation results are shown in Table 2. We see that the nonparametric complete smoothing approach offers similar results as the parametric probit on both the achieved maximum likelihood value and of the pseudo- R^2 . Indeed, the pseudo- R^2 is around 68% for the nonparametric approach, while it is 67% for the parametric approach. Furthermore, we find that the bandwidths are quite large for group 1 (recession) and for group 0 (non-recession):¹⁰ for group 1 we have $h_{spread} = 6.722$, $h_{\Delta G} = 3.335$ and for group 0 we have the values 3.580 and 11.129 respectively, while $\gamma = 0.026$. This evidence, with the previous one for the pseudo- R^2 , shows that the the linearity assumption is probably reasonable for both X_1 (the Spread, Sp_{t-1}) and X_2 (the output gap, $\Delta_{G,t-2}$). The small value of the bandwidth for the dummy variable ($\gamma = 0.061$) evidences the difference in the model between the recession and non-recession periods. The mean value of β_1 (coefficient for the Spread) is positive in nonparametric (0.2440) and parametric approach (0.2724), as well as the β_2 (coefficient for the output gap) which is positive in nonparametric and parametric approach (0.1371 and 0.2016 respectively).

Figure 4 exhibits the boxplots of the resulting local estimators of β_0, β_1 and β_2 over the T periods for group 1 and for group 0. It is interesting to see the different values of the local values of the coefficients β for the two groups of observations. For group 1 and for group 0, the estimates are similar, especially for the spread $X_{1,t} = Sp_{t-1}$ and different in the median values for the first difference of our measure of output gap $X_{2,t} = \Delta_{G,t-2}$, which appears to be a powerful indicator to forecast recession during recession periods. Moreover, the evidence

¹⁰The mean and standard deviation are 1.420 and 1.409 for the spread and 0.022 and 0.034 for the first difference in GAP.

shows that, the positive growth in inefficiency, our measure of output gap, increases the probability to be in recession more during downturn than during stable periods.

We now look at the in-sample fit for modeling the probability of recessions in Figure 5. We can indeed observe that both the nonparametric and parametric approaches fit the data well (as seen with the various measures described above). In particular, note that all recession periods, as established by the turning points of OECD.stat, i.e. of the Q1-1996 - Q1-1999, Q1-2001 - Q2-2002, Q2-2008 - Q2-2009, Q3-2011 - Q1-2013, Q1-2018 - Q2-2019 are successfully captured by our model both using the parametric and the nonparametric approaches.

Now we proceed to the out-of-sample forecasts to see if we can have a reasonably good prediction of the recession periods (one-period and two-periods ahead), using the data from the beginning till 2016:Q1.¹¹

The forecasts of the recessions are displayed in Figure 6. In most cases (and on average), we can observe a slightly better forecast value for the parametric estimator, both for the case of one-period ahead and for the two-periods ahead forecasts for the probability of observing a recession. In particular, note that, with one-period ahead forecasts, both approaches correctly and somewhat similarly warn about the recession in Q2-2018 - Q2-2019: the parametric approach slightly outperforms in two period ahead forecast, and the two are almost identical for one-period ahead forecast. Both approaches miss on Q1-2018, suggesting that the probability of recession is about 10%. Both approaches with one-period ahead forecasts also correctly alert about the non-recession (or recovery) in Q1-2016 through all Q1-2018. For the two-period ahead forecasts, the parametric approach also appears to be, to a certain extent, superior to the nonparametric approach.

It is worth recalling here that the parametric approach can be viewed as a special case of the nonparametric approach, in the sense that the latter allows for much more flexibility and can be restricted further to obtain the former through reducing this flexibility. Moreover, as discussed above, for the nonparametric approach by employing a data-driven procedure to select all the bandwidths to fit this data in an optimal way, we let the data speak for itself and fit the curvature that appears ‘best’ (in terms of MLCV) for this particular data. Interestingly, for this data set we see that despite assuming a naive (linear) and quite restrictive (e.g., constancy of the first derivative) functional form for the index function, the parametric approach still produced very similar conclusions and very similar or even slightly better forecasts than the nonparametric approach that allows for much more flexibility. This

¹¹Available sample size (past) at this time is 81.

suggests that we can have more confidence in the results and conclusions from the parametric approach, even though it imposes very restrictive assumptions. Of course for other data (e.g., for other countries or even the same country but for different time periods or with different variables) this similarity of parametric and nonparametric approaches may or may not hold *a priori* and so needs to be verified and validated on a case-by-case basis. (Indeed, it is very easy to construct an example when parametric and nonparametric approaches deliver very different results and conclusions, e.g., see Monte Carlo examples in PSZ).

4 Concluding Remarks

In this paper, we have attempted to merge two so far largely unrelated streams of literature. The first stream is about the non-parametric methods to estimate frontier efficiency of an economy. We considered various methods among the myriad of approaches, selecting and tailoring one that currently appears to be most suitable for a new measure of output gap to be used, *inter alia*, for estimating probability of economic recession. For the latter goal we have chosen the paradigm started by Estrella and Mishkin (1995, 1998), further refined by Duecker (1997) and Kauppi and Saikkonen (2008) as well as their nonparametric version recently developed by Park et al (2017). Naturally, endeavoring to merge the economic efficiency literature with other from the many paradigms for forecasting of economic recessions would be a natural direction for future research.

To illustrate our proposed framework that resulted from the merger of two different literatures, we illustrate it with data on Italy. In particular, we utilize most recent available data (from 1995 to 2019) and find that our approach (using both the linear probit model and its non-parametric version), is capable of giving useful insights, although of course is not crystal ball and more work is needed to refine and further improve the method. In particular, development of the asymptotic theory for the statistical inference in this approach (e.g., via bootstrap) would be the key direction for future research.

Pearson correlations			
	ε_{X1}	ε_{X2}	ε_Y
W	0.000184	0.000102	0.000192
Spearman rank correlations			
	ε_{X1}	ε_{X2}	ε_Y
W	-0.003584	-0.018066	-0.009535
Kendall correlations			
	ε_{X1}	ε_{X2}	ε_Y
W	0.002918	-0.011346	-0.004158

Table 1: *Correlation between $W = \text{time}$ and pure inputs ε_{X1} , ε_{X2} and pure output ε_Y .*

	Parametric Probit	Nonparametric Probit
$\hat{\beta}_0$	-1.5799	-0.3746
$\hat{\beta}_1$	0.2440	0.2724
$\hat{\beta}_2$	0.1371	0.2016
$\hat{\beta}_3$	2.5734	-
Maximum Likelihood	-0.30017077	-0.29936822
Pseudo- R^2	0.6741	0.6753

Table 2: *Parametric and Nonparametric Probit*

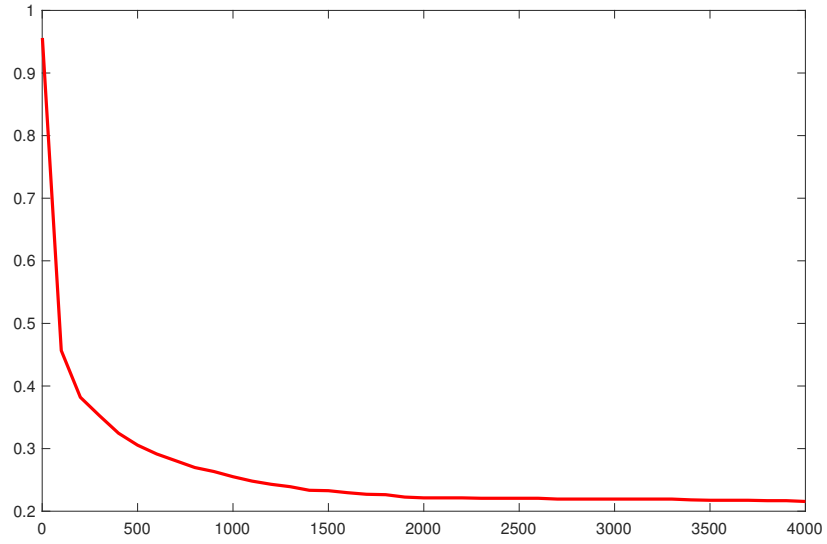


Figure 1: *The percentage of points left above the order- m frontier, as a function of m . We selected $m = 1500$ letting 5% of data points above the frontier.*

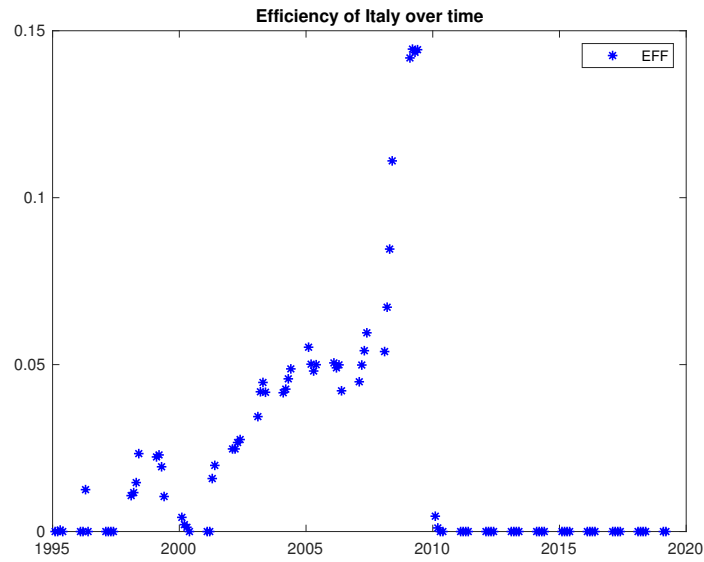


Figure 2: *Evolution of efficiency of Italy $\hat{\delta}_{m,t}$ over time.*

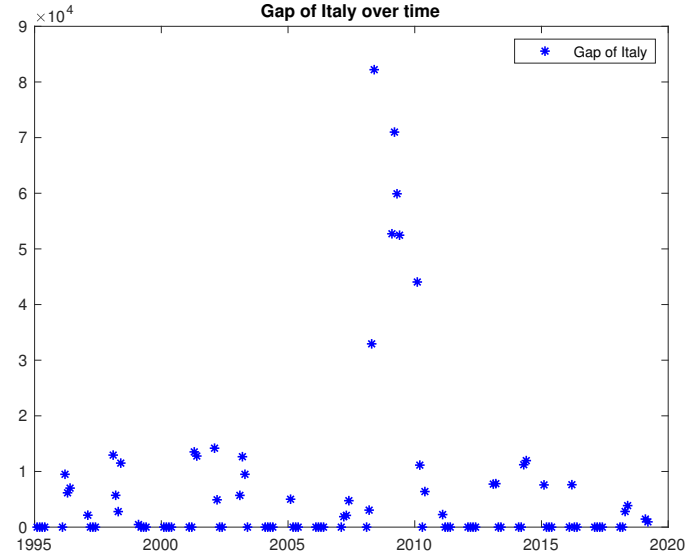


Figure 3: *Evolution of GAP of Italy over time.*

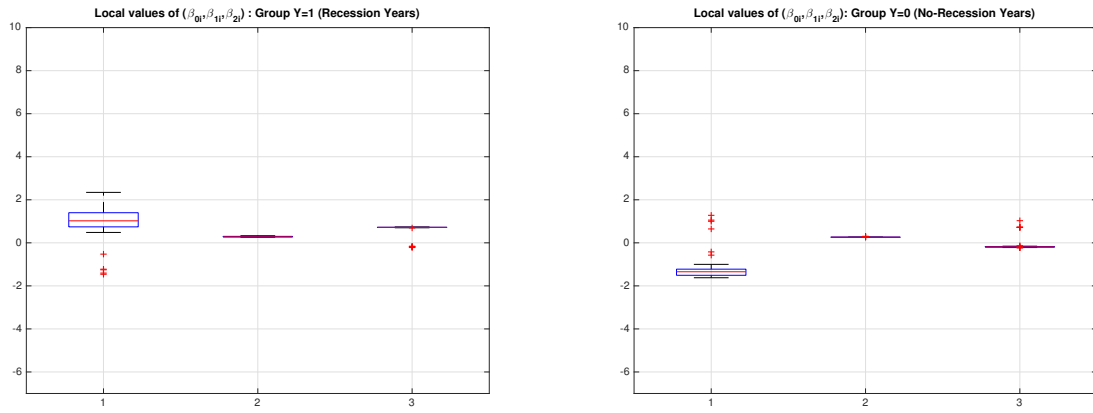


Figure 4: *Boxplots of the estimated local β 's, with the full sample of $n = 98$ data points.*

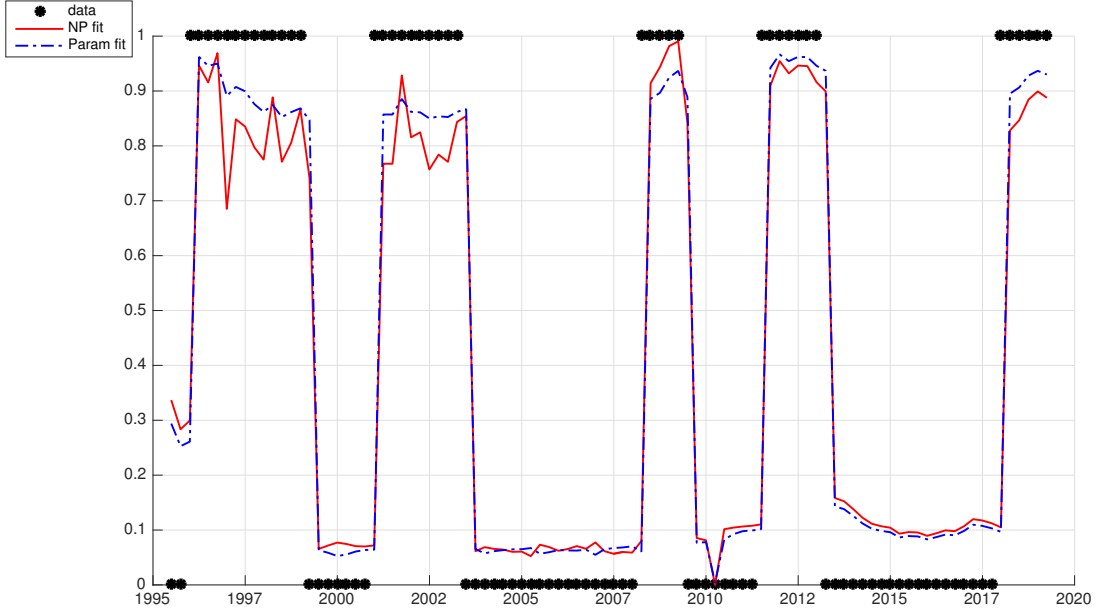


Figure 5: *In-sample fit for Recessions, with $n = 98$ data points.*

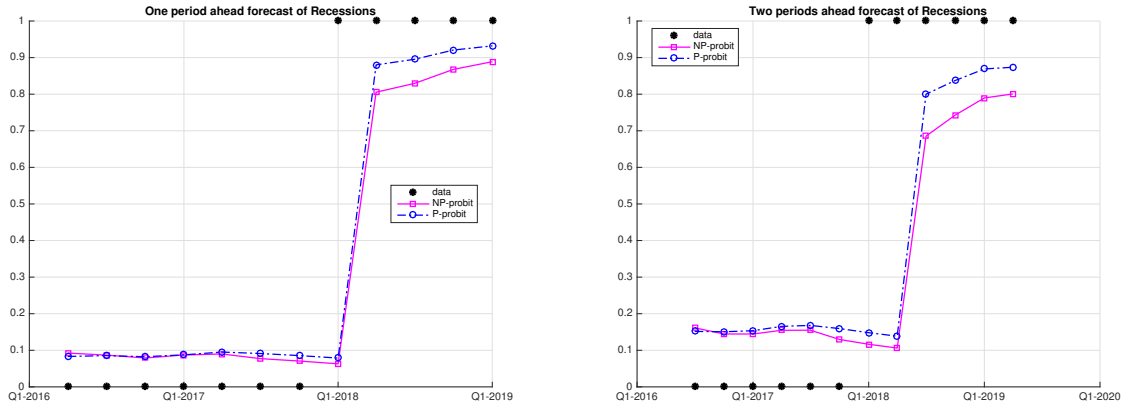


Figure 6: *Forecast of the Recessions, starting after $T = Q1 - 2016$ periods.*

A Data Description

<i>Country</i>		<i>Y</i>	<i>K</i>	<i>L</i>
Austria	Mean	3.35E+05	1.25E+06	3963.9
	SD	40359	1.13E+05	235.81
Belgium	Mean	4.11E+05	1.43E+06	4403.5
	SD	48008	2.16E+05	285.89
Denmark	Mean	2.34E+05	9.16E+05	2777.3
	SD	21562	2.10E+05	78.092
Finland	Mean	1.94E+05	7.41E+05	2405.5
	SD	25712	41161	156.8
France	Mean	2.24E+06	6.87E+06	26681
	SD	2.34E+05	2.15E+06	1356.2
Germany	Mean	3.14E+06	1.02E+07	40683
	SD	3.02E+05	6.55E+05	2052.3
Ireland	Mean	1.95E+05	2.48E+06	1854.3
	SD	63151	2.48E+06	250.37
Israel	Mean	2.01E+05	7.37E+05	3193.9
	SD	51294	1.37E+05	596.6
Italy	Mean	2.03E+06	5.78E+06	23710
	SD	92116	1.53E+06	1146
Korea	Mean	1.38E+06	6.84E+06	23429
	SD	3.92E+05	1.08E+06	2170.3
Netherlands	Mean	7.07E+05	2.31E+06	8412.5
	SD	84283	3.30E+05	485.19
New Zealand	Mean	1.27E+05	3.98E+05	2096.1
	SD	24578	1.06E+05	276.78
Norway	Mean	2.71E+05	8.50E+05	2498.4
	SD	34325	2.09E+05	221.85
Spain	Mean	1.36E+06	5.07E+06	18192
	SD	1.89E+05	9.93E+05	2143
Sweden	Mean	3.67E+05	1.24E+06	4480.6
	SD	59666	2.82E+05	336.44
United Kingdom	Mean	2.19E+06	5.69E+06	29036
	SD	2.93E+05	7.03E+05	1878.7

B Measures of Efficiencies

Date	δ	δ_m	λ	λ_m	GAP	GAP_m
1995.1	0	-0.21851	1	0.98365	0	-0.22212
1995.2	0	-0.26505	1	0.98018	0	-0.26932
1995.3	0.000446	-0.033551	1	0.99749	0.000453	-0.034077
1995.4	0	-0.035272	1	0.99737	0	-0.03581
1996.1	0	-0.21622	1	0.98387	0	-0.21942
1996.2	0	-0.215	1	0.98397	0	-0.21809
1996.3	0.012553	0.012265	1.0009	1.0009	0.012727	0.012435
1996.4	0	-0.21445	1	0.98402	0	-0.21732
1997.1	0	-0.21437	1	0.98404	0	-0.21714
1997.2	0	-0.21643	1	0.98391	0	-0.21913
1997.3	0	-0.21701	1	0.98388	0	-0.21961
1997.4	0	-0.11868	1	0.9912	0	-0.12004
1998.1	0.010762	-0.015923	1.0008	0.99882	0.01088	-0.016098
1998.2	0.011665	-0.01464	1.0009	0.99892	0.011787	-0.014793
1998.3	0.014662	0.011138	1.0011	1.0008	0.014808	0.011249
1998.4	0.023339	-0.003768	1.0017	0.99972	0.023559	-0.0038035
1999.1	0.022348	-0.004278	1.0017	0.99968	0.022546	-0.004316
1999.2	0.022951	0.020456	1.0017	1.0015	0.023142	0.020627
1999.3	0.019401	0.016882	1.0014	1.0012	0.019552	0.017014
1999.4	0.010497	0.00794	1.0008	1.0006	0.010573	0.0079975
2000.1	0.004246	0.002897	1.0003	1.0002	0.0042744	0.0029164
2000.2	0.00207	0.001371	1.0002	1.0001	0.0020827	0.0013794
2000.3	0.001275	0.000977	1.0001	1.0001	0.0012821	0.00098242
2000.4	0	-0.002169	1	0.99984	0	-0.0021798
2001.1	0	-0.00206	1	0.99985	0	-0.002069
2001.2	0	-0.000543	1	0.99996	0	-0.00054506
2001.3	0.015874	0.015159	1.0012	1.0011	0.015925	0.015208
2001.4	0.019816	0.01909	1.0015	1.0014	0.019868	0.01914
2002.1	0.024743	0.02401	1.0018	1.0018	0.024792	0.024058
2002.2	0.024744	0.023996	1.0018	1.0018	0.024778	0.024029
2002.3	0.026664	0.025895	1.0019	1.0019	0.026685	0.025915
2002.4	0.027563	0.026776	1.002	1.002	0.027567	0.02678

Continued on next page

Date	δ	δ_m	λ	λ_m	GAP	GAP_m
2003.1	0.034427	0.033625	1.0025	1.0025	0.034411	0.033609
2003.2	0.041866	0.041056	1.0031	1.003	0.04182	0.041011
2003.3	0.044656	0.043837	1.0033	1.0032	0.044579	0.043761
2003.4	0.041716	0.040892	1.003	1.003	0.041617	0.040795
2004.1	0.041579	0.040754	1.003	1.003	0.041454	0.040632
2004.2	0.042656	0.041826	1.0031	1.003	0.042501	0.041674
2004.3	0.045739	0.044903	1.0033	1.0033	0.045543	0.044711
2004.4	0.04872	0.047879	1.0035	1.0035	0.04848	0.047643
2005.1	0.055213	0.054366	1.004	1.0039	0.054905	0.054063
2005.2	0.050126	0.049274	1.0036	1.0036	0.049814	0.048967
2005.3	0.048081	0.047228	1.0035	1.0034	0.04775	0.046903
2005.4	0.049974	0.049115	1.0036	1.0035	0.049598	0.048745
2006.1	0.050511	0.04964	1.0036	1.0036	0.050098	0.049234
2006.2	0.049037	0.04815	1.0035	1.0035	0.048604	0.047725
2006.3	0.049926	0.049049	1.0036	1.0035	0.049452	0.048584
2006.4	0.042162	0.041261	1.003	1.003	0.041735	0.040843
2007.1	0.044862	0.043961	1.0032	1.0032	0.044378	0.043487
2007.2	0.049876	0.04896	1.0036	1.0035	0.049305	0.0484
2007.3	0.054173	0.053248	1.0039	1.0038	0.053518	0.052604
2007.4	0.059529	0.058596	1.0043	1.0042	0.058771	0.05785
2008.1	0.053911	0.052988	1.0039	1.0038	0.05319	0.052279
2008.2	0.067195	0.066277	1.0048	1.0047	0.066253	0.065348
2008.3	0.084608	0.083695	1.0061	1.006	0.083367	0.082468
2008.4	0.11102	0.11012	1.008	1.0079	0.10932	0.10843
2009.1	0.14186	0.14097	1.0102	1.0101	0.1396	0.13872
2009.2	0.1445	0.14278	1.0104	1.0103	0.14211	0.14042
2009.3	0.14355	0.14056	1.0103	1.0101	0.14108	0.13815
2009.4	0.14433	0.10546	1.0103	1.0076	0.14175	0.10358
2010.1	0.004596	0.003892	1.0003	1.0003	0.0045113	0.0038203
2010.2	0.001087	0.00084	1.0001	1.0001	0.0010663	0.000824
2010.3	0	-0.001903	1	0.99986	0	-0.0018656
2010.4	0	-0.001183	1	0.99992	0	-0.001159
2011.1	0	-0.00195	1	0.99986	0	-0.0019093
2011.2	0	-0.003906	1	0.99972	0	-0.0038222

Continued on next page

Date	δ	δ_m	λ	λ_m	GAP	GAP_m
2011.3	0	-0.004637	1	0.99967	0	-0.0045348
2011.4	0	-0.004595	1	0.99967	0	-0.004491
2012.1	0	-0.006826	1	0.99951	0	-0.0066676
2012.2	0	-0.003307	1	0.99976	0	-0.0032284
2012.3	0	-0.003546	1	0.99975	0	-0.0034597
2012.4	0	-0.003994	1	0.99972	0	-0.0038945
2013.1	0	-0.001276	1	0.99991	0	-0.0012435
2013.2	0	-0.000628	1	0.99996	0	-0.00061167
2013.3	0	-0.004447	1	0.99968	0	-0.004329
2013.4	0	-0.036703	1	0.99739	0	-0.035709
2014.1	0	-0.00407	1	0.99971	0	-0.0039576
2014.2	0	-0.002331	1	0.99983	0	-0.0022654
2014.3	0	-0.001315	1	0.99991	0	-0.0012773
2014.4	0	-0.000639	1	0.99995	0	-0.00062037
2015.1	0	-0.00888	1	0.99937	0	-0.0086167
2015.2	0	-0.037797	1	0.99732	0	-0.036658
2015.3	0	-0.038495	1	0.99727	0	-0.037316
2015.4	0	-0.006379	1	0.99955	0	-0.0061805
2016.1	0	-0.001171	1	0.99992	0	-0.001134
2016.2	0	-0.000165	1	0.99999	0	-0.00015971
2016.3	0	-0.005117	1	0.99964	0	-0.0049506
2016.4	0	-0.000651	1	0.99995	0	-0.00062954
2017.1	0	-0.000332	1	0.99998	0	-0.00032091
2017.2	0	-0.000683	1	0.99995	0	-0.00065988
2017.3	0	-0.000485	1	0.99997	0	-0.00046837
2017.4	0	-0.001568	1	0.99989	0	-0.0015136
2018.1	0	-0.00429	1	0.9997	0	-0.0041393
2018.2	0	-0.001815	1	0.99987	0	-0.0017505
2018.3	0	-0.00188	1	0.99987	0	-0.0018124
2018.4	0	-0.009783	1	0.99931	0	-0.0094274
2019.1	0	-0.009008	1	0.99937	0	-0.0086771
2019.2	0	-0.034192	1	0.9976	0	-0.032923

Table 3: *Different measures of efficiency: Efficiency δ , Order-m Efficiency δ_m , Time Conditional Efficiency λ , Order-m Time Conditional Efficiency λ_m , Pure Time Conditional Efficiency GAP , Order-m Pure Time Conditional Efficiency GAP_m .*

References

- Aitchison, J. and Aitken, C. G. G.: 1976, Multivariate binary discrimination by the kernel method, *Biometrika* **63**, 413–420.
- Apel, M. and Jansson, P.: 1999, A theory-consistent system approach for estimating potential output and the NAIRU, *Economics Letters* **64**, 271–275.
- Basistha, A. and Nelson, C. R.: 2007, New measures of the output gap based on the forward-looking new keynesian phillips curve, *Journal of Monetary Economics* **54**, 498–511.
- Basu, S. and Fernald, J. G.: 2009, What do we know (and not know) about potential output?, *Federal Reserve Bank of St. Louis Review* **91**, 187–213.
- Baxter, M. and King, R.: 1999, Measuring business cycles: Approximate band-pass filters for economic time series, *Review of Economics and Statistics* **81**, 575–593.
- Beveridge, S. and Nelson, C.: 1981, A new approach to the decomposition of economic time series into permanent and transitory components with particular attention to the measurement of the business cycle., *Journal of Monetary Economics* **7**, 151–174.
- Bădin, L., Daraio, C. and Simar, L.: 2019, A bootstrap based approach for bandwidth selection in estimating conditional efficiency measures, *European Journal of Operational Research* **forthcoming**.
- Cazals, C., Florens, J. and Simar, L.: 2002, Nonparametric frontier estimation: A robust approach, *Journal of Econometrics* **106**, 1–25.
- Clark, P.: 1987, The cyclical component of US economic activity, *Quarterly Journal of Economics* **102**, 797–814.
- Daouia, A. and Gijbels, I.: 2011, Robustness and inference in nonparametric partial-frontier modeling, *Journal of Econometrics* **161**, 147–165.
- Daouia, A. and Simar, L.: 2007, Nonparametric efficiency analysis: A multivariate conditional quantile approach, *Journal of Econometrics* **140**, 375–400.
- Daraio, C. and Simar, L.: 2005, Introducing environmental variables in nonparametric frontier models: A probabilistic approach, *Journal of Productivity Analysis* **24**, 93–121.
- Dueker, M.: 1997, Strengthening the case for the yield curve as a predictor of u.s. recessions, *Review Federal Reserve Bank of St Louis* **79**, 41–51.
- Estrella, A. and Mishkin, F. S.: 1995, Predicting u.s. recessions: Financial variables as leading indicators. NBER Working Paper No. 5379.
- Estrella, A. and Mishkin, F. S.: 1998, Predicting u.s. recessions: Financial variables as leading indicators, *The Review of Economics and Statistics* **80**, 45–61.

- Florens, J., Simar, L. and van Keilegom, I.: 2014, Frontier estimation in nonparametric location-scale models, *Journal of Econometrics* **178**, 456–470.
- Gali, J. and Gertler, M.: 1999, Inflation dynamics: A structural econometric analysis, *Journal of Monetary Economics* **44**, 195–222.
- Gerlach, S. and Smets, F.: 1999, Output gaps and monetary policy in the EMU area, *European Economic Review* **43**, 801–812.
- Harvey, A.: 1985, Trends and cycles in macroeconomic time series, *Journal of Business and Economic Statistics* **3**, 216–227.
- Henderson, D. J. and Parmeter, C. F.: 2015, *Applied nonparametric econometrics*, Cambridge University Press.
- Hodrick, R. and Prescott, E. C.: 1997, Postwar u.s. business cycles: An empirical investigation, *Journal of Money, Credit and Banking* **29**, 1–16.
- Horowitz, J. L.: 2009, *Semiparametric and nonparametric methods in econometrics*, Vol. 12, Springer.
- Jeong, S., Park, B. and Simar, L.: 2010, Nonparametric conditional efficiency measures: asymptotic properties, *Annals of Operational Research* **173**, 105–122.
- Kauppi, H. and Saikkonen, P.: 2008, Predicting u.s. recessions with dynamic binary response models, *Review of Economics and Statistics* **90**, 777–791.
- Kumar, S. and Russell, R.: 2002, Technological change, technological catch-up, and capital deepening: Relative contributions to growth and convergence, *American Economic Review* **92**, 527–548.
- Kuttner, K.: 1994, Estimating potential output as a latent variable, *Journal of Business and Economic Statistics* **12**, 361–368.
- Li, D., Simar, L. and Zelenyuk, V.: 2016, Generalized nonparametric smoothing with mixed discrete and continuous data, *Computational Statistics & Data Analysis* **100**, 424–444.
- Mastromarco, C. and Simar, L.: 2015, Effect of fdi and time on catching-up: New insights from a conditional nonparametric frontier analysis, *Journal of Applied Econometrics* **30**, 826–847.
- Mastromarco, C. and Simar, L.: 2018, Globalization and productivity: A robust nonparametric world frontier analysis, *Economic Modelling* **69**, 134–149.
- Murray, C.: 2003, Cyclical properties of baxter king filtered time series, *The Review of Economics and Statistics* **85**, 472–476.
- Park, B., Simar, L. and Zelenyuk, V.: 2017, Nonparametric estimation of dynamic discrete choice models for time series data, *Computational Statistics & Data Analysis* **108**, 97–120.

- Roberts, J.: 2001, Estimates of the productivity trend using time-varying parameter techniques., *The B.E. Journal of Macroeconomics*. **Contributions to Macroeconomics 1 (Article 3)**.
- Simar, L. and Vanhems, A.: 2012, Probabilistic characterization of directional distances and their robust versions, *Journal of Econometrics* **166**, 342–354.
- Watson, M.: 1986, Univariate detrending methods with stochastic trends, *Journal of Monetary Economics* **18**, 49–75.
- Wheelock, D. C. and Wilson, P.: 1995, Explaining bank failures: Deposit insurance, regulation, and efficiency, *The Review of Economics and Statistics* **77**, 689–700.