



BANCA D'ITALIA
EUROSISTEMA

Temi di Discussione

(Working Papers)

The predictive power of Google searches
in forecasting unemployment

by Francesco D'Amuri and Juri Marcucci

November 2012

Number

891



BANCA D'ITALIA
EUROSISTEMA

Temi di discussione

(Working papers)

The predictive power of Google searches
in forecasting unemployment

by Francesco D'Amuri and Juri Marcucci

Number 891 - November 2012

The purpose of the Temi di discussione series is to promote the circulation of working papers prepared within the Bank of Italy or presented in Bank seminars by outside economists with the aim of stimulating comments and suggestions.

The views expressed in the articles are those of the authors and do not involve the responsibility of the Bank.

Editorial Board: MASSIMO SBRACIA, STEFANO NERI, LUISA CARPINELLI, EMANUELA CIAPANNA, FRANCESCO D'AMURI, ALESSANDRO NOTARPIETRO, PIETRO RIZZA, CONCETTA RONDINELLI, TIZIANO ROPELE, ANDREA SILVESTRINI, GIORDANO ZEVI.

Editorial Assistants: ROBERTO MARANO, NICOLETTA OLIVANTI.

ISSN 1594-7939 (print)

ISSN 2281-3950 (online)

Printed by the Printing and Publishing Division of the Bank of Italy

THE PREDICTIVE POWER OF GOOGLE SEARCHES IN FORECASTING UNEMPLOYMENT

by Francesco D'Amuri* and Juri Marcucci*

Abstract

We suggest the use of an index of Internet job-search intensity (the Google Index, GI) as the best leading indicator to predict the US monthly unemployment rate. We perform a deep out-of-sample forecasting comparison analyzing many models that adopt our preferred leading indicator (GI), the more standard initial claims or combinations of both. We find that models augmented with the GI outperform the traditional ones in predicting the unemployment rate for different out-of-sample intervals that start before, during and after the Great Recession. Google-based models also outperform standard ones in most state-level forecasts and in comparison with the Survey of Professional Forecasters. These results survive a falsification test and are also confirmed when employing different keywords. Based on our results for the unemployment rate, we believe that there will be an increasing number of applications using Google query data in other fields of economics.

JEL Classification: C22, C53, E27, E37, J60, J64.

Keywords: Google econometrics, forecast comparison, keyword search, US unemployment, time series models.

Contents

1. Introduction	5
2. Data.....	8
3. Forecasting models	14
4. Out-of-sample forecasting comparison	15
4.1 Main results	15
4.2 Formal tests of forecast accuracy	17
5. Robustness checks	20
5.1 Different in-sample/out-of-sample	20
5.2 State level forecasts	21
5.3 Nonlinear models.....	22
5.4 Alternative keywords	24
5.5 Falsification test	25
6. Comparison with the Survey of Professional Forecasters	26
7. Conclusions	28
References	29
Tables and figures	32

* Bank of Italy, Economic Research and International Relations.

1 Introduction¹

In this paper we suggest the use of the Google index (GI), based on internet job searches performed through Google, as the best leading indicator to predict the US monthly unemployment rate.

Quantitative data on internet use are becoming quickly available and will constitute an invaluable source for economic analysis in the near future. Following the growing popularity of the internet as a job-search tool and the increasing need for reliable and updated unemployment forecasts, especially during recessions, in this article we suggest the use of an indicator based on Google job-search-related query data (i.e., the Google Index, GI) as the best leading indicator to predict the US monthly unemployment rate.² We test the predictive power of this indicator by means of a deep out-of-sample comparison among more than five hundred forecasting models which differ along three dimensions: (i) The exogenous variables adopted as leading indicators; (ii) the econometric specification; and (iii) the length of the estimation sample. In particular, we estimate standard time series (ARMA) models and we augment them with the Initial Claims (IC, a widely accepted leading indicator for the US unemployment rate), the GI, or combinations of both. In carrying out our comparison, we include both linear and non-linear models, since the former typically capture short-run developments, while the latter can better approximate

¹We wish to thank F. Busetti, D. Depalo, O. Jorda, F. Lotti, R. Mosconi, C. Perna, A. Rosolia, P. Sestito, H. Varian and K. Zimmermann for their useful suggestions. We also thank seminar participants at the 32nd International Symposium on Forecasting, the 2nd International Conference in Memory of Carlo Giannini, the XVIII SNDE Symposium, the 45th Scientific Meeting of the Italian Statistical Society, the Fourth Italian Congress of Econometrics and Empirical Economics, Fondazione ENI Enrico Mattei and LUISS ‘Guido Carli’ University for their useful comments. The views expressed are those of the authors and do not necessarily reflect those of the Bank of Italy. Emails: francesco.damuri@bancaditalia.it (Francesco D’Amuri), juri.marcucci@bancaditalia.it (Juri Marcucci, corresponding author). Correspondence address: Bank of Italy, Economic Research Department, Via Nazionale 91, 00184, Rome, Italy.

²The time series of the US unemployment rate is certainly one of the most studied in the literature. Proietti (2003) defines this series as the ‘testbed’ or the ‘case study’ for many (if not most) non-linear time series models. In fact, many papers have documented its asymmetric behavior. Neftci (1984), DeLong and Summers (1986) and Rothman (1998) document the type of asymmetry called *steepness* for which unemployment rates rise faster than they decrease. Sichel (1993) finds evidence for another type of asymmetry called *deepness* in which contractions are deeper than expansions. McQueen and Thorley (1993) find *sharpness* for which peaks tend to be sharp while troughs are usually more rounded. In a recent paper, Barnichon and Nekarda (2012) develop a model based on labor market flows to forecast unemployment; according to their results, this approach can greatly improve the forecast accuracy of standard time series forecasts.

the dynamics of the unemployment rate during economic contractions. We also compare models estimated over samples of different length, because the GI is only available since the first week of January 2004, while the IC are available since 1967. Indeed, an exercise comparing the forecasting performance of models estimated on the short sample only (starting in 2004) would be of little practical relevance if models estimated on the longer sample (starting in 1967) were better at predicting the unemployment rate.

We find that models augmented with the GI significantly outperform the more traditional ones in predicting the US unemployment rate: when forecasting at one step ahead the mean squared error (MSE) of our best model using GI as a leading indicator (0.023) is 28% lower than the best model not including it and estimated on the same sample. The best Google model estimated on the short sample outperforms alternative models estimated on the long sample; even in this comparison, the best Google model shows a MSE that is 18% lower than the best non-Google model. These results are rather striking since Google models estimated on the short sample use only 4 years of data, while the ones using the long sample are estimated on a time series that is 10 times bigger (more than 40 years). Relative forecast accuracy increases at longer forecast horizons: at three steps ahead, when using the GI the MSE decreases by 40% compared to the best alternative model estimated on the same sample, and by 22% when considering models estimated on the long sample.

Furthermore, we select the best models in terms of the lowest MSE and assess their out-of-sample forecast ability by testing for equal forecast accuracy and superior predictive ability using respectively Diebold and Mariano’s (1995) test and the Model Confidence Set (MCS) test by Hansen et al. (2011). Our results show that not only the best model in terms of lowest MSE always includes GI as a leading indicator, but also that models with GI estimated over the short sample (i.e. from 2004 onwards) beat models estimated over the long sample (i.e. from 1967 onwards) using the IC as a leading indicator. Moreover, around one third of the best models selected in the final MCS adopt the GI as the leading indicator.

Our results also hold after a number of *robustness checks*. In fact, the main results hold when conducting the horse race in different out-of-sample intervals that start before,

during and after the recent recession. When we forecast in the middle of the recession the performance of the GI as a leading indicator for unemployment is even more striking: around two thirds of the Google-based models enter the final MCS. We also repeat the forecast horse race for each of the 50 US states plus District of Columbia (DC) rather than at the federal level, finding that, when forecasting at one- and two-step-ahead, the best five models include the GI among the explanatory variables in 70.2% and in 62% of the cases, respectively. We also test the forecasting properties of two alternative, and less popular, job-search-related keywords, “*collect unemployment*” and “*job center*” finding that the latter improves the performance of standard time series models estimated on the same in-sample when forecasting at one, two and three steps ahead. We also use as a leading indicator the first principal component of the three GIs adopted in the paper, finding that the forecasting performance of our forecasting models improves even further.

Moreover, we provide a falsification test, checking the forecasting performance of an alternative Google-based indicator that shows the highest correlation with the unemployment rate in-sample, but captures the interest for a keyword that is completely unrelated to job-search activities. Models augmented with this fake GI indicator never rank among the best models in terms of forecasting ability, providing indirect evidence for the relevance of web-search data when the underlying keywords have a direct link with unemployment and job search.

Finally, we construct a group of quarterly forecasts of the unemployment rate using the best models from our horse-race over the monthly series and compare them with the quarterly predictions released by the Survey of Professional Forecasters (SPF) conducted by the Federal Reserve Bank of Philadelphia. Conditioning on the same information set, models using the GI outperform the professionals’ forecasts, showing a lower MSE by 67%.

The innovative data source employed in this article has already been used in epidemiology and in different fields of economics (Edelman, 2012). The first article using Google data (Ginsberg et al., 2009) estimates the weekly ‘*influenza*’ activity in the US using an index of health seeking behavior equal to the incidence of *influenza*-related internet queries. Da et al. (2011) show the relevance of Google data as a direct and timely mea-

sure of investors' attention for a sample of Russel 3000 stocks. Billari et al. (2012) use web-search data related to fertility as a leading indicator of the US birth rate. Baker and Fradkin (2011) develop a job-search activity index to analyze the reaction of job-search intensity to changes in unemployment benefit duration in the US.

To the best of our knowledge, this is the first paper using this kind of internet indicator to forecast the monthly unemployment rate in the US. Askitas and Zimmermann (2009) were the first ones using Google data to forecast unemployment with an application to Germany. However, there have also been some works for other countries, in particular for Italy (D'Amuri, 2009) and Israel (Suhoy, 2009), while Choi and Varian (2012) use web-search data to forecast consumer behavior and initial unemployment claims for the US. Central Banks are also starting to publish reports on the suitability of Google data to complement more standard economic indicators (see for example Artola and Galan, 2012, McLaren and Shanbhorge, 2011 and Troy et al., 2012 respectively for Spain, the UK and Australia). Based on our results for the unemployment rate, we believe that there will be further applications using Google query data in other fields of economics.

The paper is organized as follows: In Section 2 we describe the data used to predict the US unemployment rate, with a particular emphasis on the GI. In Section 3 we discuss the models employed to predict the US unemployment rate, while in Section 4 we compare the out-of-sample performance of such models. In Section 5 we show that the superior predictive performance of Google-based models is confirmed (i) when using different out-of-sample intervals that start before, during and after the recent recession; (ii) when forecasting at the state rather than at the federal level; (iii) in comparison with non-linear models; and (iv) by a falsification test. In Section 6 we compare our predictions with those of the Survey of Professional Forecasters, while Section 7 concludes.

2 Data

The data used in this paper come from different sources. The seasonally adjusted monthly unemployment rate is the one released by the Bureau of Labor Statistics (BLS) and comes from the Current Employment Statistics and the Local Area Unemployment Statistics for

the national and the state level, respectively. Unemployment rates for month t refer to individuals who do not have a job, but are available for work, in the week including the 12th day of month t and who have looked for a job in the prior 4 weeks ending with the reference week. For the federal level the available sample is 1948.1-2011.6, while for the state level the data on unemployment are available from 1976.1 to 2011.6. We complement these data with a well-known leading indicator for the unemployment rate (see for example Montgomery et al. 1998): the weekly seasonally-adjusted IC released by the U.S. Department of Labor,³ available since 1967.1 for the US and since 1986.12 for the single states.

The exogenous variable specific to this study is the weekly GI which summarizes the job searches performed through the Google website. The GI represents how many web searches have been done for a particular keyword in week t in a given geographical area r (i.e., $V_{t,r}$) relative to the total number of web searches performed through Google in the same week and area ($T_{t,r}$). The search index for week t is thus given by $GI_{t,r} = \frac{V_{t,r}}{T_{t,r}}$. Absolute values of the index are not publicly available, since Google normalizes the index $GI_{t,r}$ to 100 in the week in which it reaches the maximum level. Data are gathered using IP addresses only if the number of searches exceeds a certain threshold. Repeated queries from a single IP address within a short time are eliminated. The data are available almost in real time starting with the week ending on January 10, 2004. Our main indicator summarizes the incidence of queries including the keyword “*jobs*” on total queries performed through Google in the relevant week (this index is labeled GI henceforth).⁴

We choose to use the keyword “*jobs*” as the main indicator of job-search activities mainly for two reasons. First, we found that the keyword “*jobs*” was the most popular among different job-search-related keywords. Absolute search volumes are not available,

³Since seasonally adjusted data are issued only at the national level, we have performed our own seasonal adjustment for the state-level data using Tramo-Seats.

⁴We have adjusted both the weekly and the monthly indicators for seasonality using Tramo-Seats. The type of seasonality of the Google data is completely different from the usual one we find in economic variables. Typically, there are yearly troughs at the end of each year because the total number of queries is inflated by Christmas-related searches. The data, available free of charge, were downloaded on July 17, 2011.

but it is possible to identify the most popular keywords looking at relative incidences. In Figure 1, we plot the monthly averages for the values of the GI for the keywords “facebook”, “youtube”, and “jobs”; we also plot the values for two alternative job-search-related keywords “collect unemployment” and “job center” (henceforth labeled G2 and G3), whose forecast performance is tested in Section 5.4. We notice that “facebook” touches the highest incidence among the keywords, while the GI for “jobs” is constantly around the value of 10. This means that, when searches for “facebook” were at their peak, there was still one keyword search for “jobs” for every ten searches for “facebook”, which is, incidentally, the most popular keyword of all. The results are similar when conducting the comparison with the keyword “youtube”, another popular search, that reaches a maximum level above 40 during the considered interval. The other alternative job-search-related keywords we consider (“collect unemployment” and “job center”) fair less well in terms of popularity, with very low relative incidences.

Apart from its popularity, the second reason why we chose the keyword “jobs” is that we believe that it is widely used across the broadest range of job seekers, and as a consequence is less sensitive to the presence of demand or supply shocks specific to subgroups of workers that could bias the values of the GI and its ability to predict the overall unemployment rate. Finally, it has to be noted that the numerator of the index contains all the keyword searches including the word “jobs”, such as “public jobs” or “California jobs”, for example. As a consequence, the index is based on a broader set of queries including the word “jobs”, some of which might actually be unrelated to job search. Such a measurement error is unlikely to be correlated with the monthly unemployment rate over time and should, if anything, reduce the predictive power of our leading indicator. Nevertheless, in order to improve the precision of our GI, we subtract from the numerator the keyword searches for “Steve Jobs”, a popular search including the word “jobs”.

The variable has other limitations: Individuals looking for a job through the internet (jobs available through the internet) may well be not randomly selected among job seekers (jobs). Moreover, the indicator captures overall job-search activities, that is the sum of searches performed by unemployed and employed people. This limitation is made more severe by the fact that, while unemployed job search is believed to follow the anti-cyclical

variation of job separation rates, on-the-job search is normally assumed to be cyclical. We acknowledge that this could introduce some bias in our GI; nevertheless such a bias should, if anything, reduce the precision of our forecasts.

In the empirical analysis we align the GI and IC data with the relevant weeks for the unemployment survey. When constructing the GI or the IC for month t , we take into account the week including the 12th of the month and the three preceding weeks, exactly the same interval used to calculate the unemployment rate for month t reported in the official statistics. When there are more than four weeks between the reference week of month t and the following one in month $t + 1$, we do not use either the GI or the IC for the week that is not used by the official statistics in order to calculate the unemployment rate (see Figure 2 for a visual description of the alignment procedure).

Table 1 reports the descriptive statistics for the monthly US unemployment rate and both leading indicators (IC and the GI, both weekly and monthly) for the short sample (2004.1-2011.6). The monthly unemployment rate was equal on average to 6.5% during this interval, ranging between a minimum of 4.4% and a maximum of 10.1%. The series has a right-skewed distribution and a high kurtosis which make it non-normal as suggested by the Jarque-Bera test that almost always rejects the null hypothesis of normality. IC and GI share similar features, being non-normal and right-skewed, both at the weekly and the monthly level.

In Figure 3 and 4, we plot separately the monthly unemployment rate and our exogenous variables adopted as leading indicators over the relevant sample periods. In Figure 3, we plot the unemployment rate and the IC over the long sample (1967.1-2011.6), according to the availability of IC. Figure 4 depicts instead the unemployment rate along with the IC as well as the GI for “*jobs*” over the short sample. These latter indices are rescaled with respect to the maximum weekly value of each series over the sample. In both cases the two series show similar patterns, with both IC and the GI seeming to be leading indicators for the unemployment rate. This behavior is confirmed by the correlations: focusing on the short sample, we notice that both the GI and the IC are highly correlated with the level of the unemployment rate. For the IC at various lags up to the second, the correlation is between 0.83 and 0.88, while for the GI the correlation is always

greater than 0.90.⁵

In particular, the correlations of the GI for “*jobs*” with the unemployment rate are higher than those of the IC the leading indicator widely accepted by the literature. This is true both for the contemporaneous correlation and when considering one or two lags, suggesting that the Google-based indicator can be rather helpful when predicting unemployment.

In the literature many works impose the presence of a unit root or induce stationarity with a particular transformation - see for example Rothman (1998) who induces stationarity with a log-linear de-trended transformation ($u_t^{LLD} = \log(u_t) - \hat{a} - \hat{b}t$) and checks his out-of-sample results with the HP-filtered unemployment in $\log(u_t^{LHP})$). Montgomery et al. (1998) model the level of the monthly unemployment rate arguing that unit-root non-stationarity is hard to justify for the US unemployment rate because it is a rate that varies within a limited range. Similarly, Koop and Potter (1999) argue that since the unemployment rate is bounded between 0 and 1, it cannot exhibit global unit root behavior.⁶ As argued by Koop and Potter (1999) the bounded nature of the unemployment rate should guarantee a bounded behavior and therefore makes pre-testing for the unit root unnecessary. And of course, the same would apply to our GIs, given the fact that their weekly series are bounded between 0 and 100.

We have nevertheless checked for non-stationarity of the monthly US unemployment rate by computing a univariate unit root test for the integration of the series which is robust to structural breaks, outliers and non-linearities. In fact, as pointed out by Choi and Moh (2007), standard unit-root tests are known to be biased towards the non-rejection of the null of a unit-root when they are applied to time series with strong non-linear dynamics (such as the unemployment rate). We have thus performed the Range Unit Root test (RUR) suggested by Aparicio et al. (2006) which is a fully non-parametric unit-root test constructed using the running ranges of the series. This test is invariant to monotonic transformations of the series of interest and is robust to important parameter

⁵For the sake of brevity we have decided not to report the results on correlations and other results which are however available in the online Appendix.

⁶To make the series unbounded, Koop and Potter (1999) use the logistic transformation ($u_t^{logit} = \log(\frac{u_t}{1-u_t})$) suggested also by Wallis (1987).

shifts due to outliers or structural breaks.⁷

When we apply the RUR and the Forward-Backward RUR⁸ test on the level of the US monthly unemployment rate we find that for the long sample, i.e. 1967.1-2011.6, we fail to reject the null of unit root. In fact, the RUR test is 1.644 (with left-tail critical value of 1.30 and right-tail critical value of 3.34 at 5%) and the FB-RUR is 2.479 (with left-tail critical value of 1.87 and right-tail critical value of 3.34 at 5%). Nevertheless, with the short sample, i.e. 2004.1-2011.6, we reject the null of a unit root. The RUR test is equal to 3.795 (with left-tail critical value of 1.17 and right-tail critical value of 3.18 at 5%), while the FB-RUR test is equal to 4.696 (with left-tail critical value of 1.80 and right-tail critical value of 4.35 at 5%).

Given the fact that we are more interested in the short sample where the GI is available, we adopted the more agnostic approach of Koop and Potter (1999) or Montgomery et al. (1998). Therefore we have decided not to explicitly restrict our models to the stationary regime and we will present all our forecasting results using the levels of the monthly US unemployment rate as in Montgomery et al. (1998) and Proietti (2003).

⁷Given a series of interest y_t , Aparicio et al. (2006) considered the recursive ranges $R_i^y = y_{i,i} - y_{1,i}$, where $y_{1,i} = \min\{y_1, y_2, \dots, y_T\}$ and $y_{i,i} = \max\{y_1, y_2, \dots, y_T\}$. The Range Unit-Root test, $J_0^{(T)}$ is given as:

$$J_0^{(T)} = \frac{1}{\sqrt{T}} \sum_{i=2}^T \mathbf{1} \left(\Delta R_i^{(y)} > 0 \right) \quad (1)$$

where $\mathbf{1} \left(\Delta R_i^{(y)} > 0 \right)$ is the indicator function, taking value 1 when the change in the range is positive and zero otherwise. Thus the test determines the number of level crossings of the data, obtained by taking the difference of the extremes in an ever-growing sample of the series. Under the null of a unit root, $J_0^{(T)}$ converges to a non-degenerate unimodal random variable which peaks at the value 2. On the contrary, when the series is stationary, $J_0^{(T)}$ converges to 0 in probability. Therefore, we can use the left tail of the distribution of $J_0^{(T)}$ to discriminate between a stationary and a non-stationary series without a trend and the right tail if the variable is stationary with a linear trend. Critical values for the test are calculated from 20,000 replications of the null model of a random walk with normal increments.

⁸Aparicio et al. (2006) also suggest the Forward-Backward RUR (FB-RUR) test which is based on the reversed realizations of the sample of observations, $y'_t = y_{T-t+1}$, and is given as:

$$J_*^{(T)} = \frac{1}{\sqrt{2T}} \sum_{i=2}^T \left[\mathbf{1} \left(\Delta R_i^{(y)} > 0 \right) + \mathbf{1} \left(\Delta R_i^{(y')} > 0 \right) \right] \quad (2)$$

which improves upon the RUR test when additive outliers are present.

3 Forecasting models

In our forecasting exercise we compare a total of more than 500 linear ARMA models for the US unemployment rate u_t .

To start with, we estimate 384 models that can be grouped into three broad categories:

- a) models *not including* the GI as an exogenous variable and estimated on the long sample (in-sample 1967.1-2007.2; out-of-sample 2007.3-2011.6)
- b) models *not including* the GI as an exogenous variable but estimated on the short sample, for which Google data are available (in-sample 2004.1-2007.2; out-of-sample 2007.3-2011.6)
- c) models *including* the GI as an exogenous variable and estimated over the short sample (in-sample 2004.1-2007.2; out-of-sample 2007.3-2011.6).

Within these three broad groups we estimate exactly the same set of models, both in terms of lag specification and of exogenous variables included, with the GI indicator added as an additional independent variable in the last, otherwise identical, set of models.

The rationale for repeating our forecasting exercise along three dimensions is straightforward. The inclusion of the GI among the exogenous variables limits the length of the estimation interval, given that the indicator is available since January 2004 only. An exercise comparing the forecasting performance of models estimated on samples starting in 2004.1 could be able to assess the predictive power of the GI, but it would be of little practical relevance if models estimated on the longer sample were better at predicting unemployment rate dynamics.

Within the three groups we estimate pure time series $AR(p)$ and $ARMA(p, q)$ models, with at most 2 lags for p and q , for a total of four models ($AR(1)$, $AR(2)$, $ARMA(1,1)$ and $ARMA(2,2)$).

In addition, we augment these basic specifications with exogenous leading indicators, i.e. $ARMAX(p, q)$:

$$\phi(L)u_t = \mu + x_t'\beta + \theta(L)\varepsilon_t \quad (3)$$

where x'_t is a vector with a first column of ones and one or more columns of leading indicators. These indicators should help in improving the predictions of the US unemployment rate.

In particular, following Montgomery et al. (1998) we use as a leading indicator (both on the short and the long sample) the monthly IC, i.e. IC_t , their weekly levels ($IC_{w1,t}$, $IC_{w2,t}$, $IC_{w3,t}$, and $IC_{w4,t}$) and their first and second lags. All the models are estimated adding seasonal multiplicative factors to account for residual seasonality.⁹ In Table 2, we summarize all the groups of models within the short and the long sample.¹⁰

In our pseudo-out-of-sample exercise we consider the situation that real forecasters face when they produce their predictions and the future values of the exogenous variables (x_t) need to be forecast. At any given date, we have run our forecasting horse-race using only the information that was really available at that time. Therefore, we have adopted simple AR(1) models to predict x_t , so that we could use such predictions as inputs in our forecasting models. For robustness, we have considered several different models.¹¹ The results are quite similar and are therefore unreported for the sake of brevity. They are available from the authors upon request.

4 Out-of-Sample Forecasting Comparison

4.1 Main results

After having introduced the set of models included in our analysis, this Section assesses their forecasting performance in the out-of-sample interval 2007.3-2011.6.

In Table 3 we rank the best 15 models for the US monthly unemployment rate in terms of lowest Mean Squared Errors (MSE) at one, two and three steps ahead. At any forecast horizon, the best model always includes the GI for “jobs” (i.e., G1) among the exogenous

⁹In particular, we used a seasonal multiplicative autoregressive factor $SAR(12)$ for AR models and both an AR and MA seasonal $SMA(12)$ for ARMA models.

¹⁰In all our forecasting exercises we use a rolling window. However we have also performed our forecasting horse-race using a recursive scheme. The results are similar to those with a rolling scheme and are not reported for the sake of brevity, but they are available upon request.

¹¹We have adopted an AR(2), ARMA(1,1) and ARMA(2,2).

variables. At one-step-ahead, the best model is an ARMAX(2,2) augmented with the IC for unemployment benefits and with the value of G1, both with one lag and taken at their value for the fourth week (i.e., the one including the 12th of each month, in which the BLS survey is conducted). The best model with no Google data estimated on the same in-sample (2004.1-2007.2) is an ARX(1) with one lag of the IC for the fourth week and the seasonal factor; this model ranks 141st in the forecast comparison, with a Mean Squared Error that is equal to 0.032, a value 23% higher than the best model using Google (0.026). Models estimated on the longer in-sample (1967.1-2007.2), for which Google data are not available, show a better forecasting performance; in this case, the best model (an ARMAX(2,2) with two lags for the IC and a seasonal factor) ranks 7th in the forecast comparison, but its MSE is still 8% higher than the best Google-model estimated over the short sample. As expected, MSEs of the predictions rise for all models when forecasting at longer horizons. Nevertheless, the gap in favor of Google-based models widens. At two steps ahead, the best Google-based model (an ARX(1) with the first lag of the monthly IC and G1 plus the seasonal factor) has a MSE of 0.06; the best non-Google model estimated on the same in-sample has a 28% higher MSE, ranking 149th in the forecast comparison, while this gap reduces to 10% for the best non-Google model estimated on the long sample. These results are rather striking since Google models estimated on the short sample use only 4 years of data, while those using the long sample are estimated on a time series which is 10 times bigger (40 years).

The advantage for Google-based models further increases when forecasting at three steps ahead; in this case the advantage in terms of lower MSE is 19.8% and 55.0% compared to the best non-Google models estimated on the long and the short sample respectively. Figure 5 depicts the forecast errors of the best models overall, the best non-Google models over the long sample and the short sample in addition to the forecast errors from the three non-linear models¹² used. The three panels depict the last recession with a shaded area. As we can see from the top panel which relates to 1-step-ahead forecast errors from model number 493 (best model overall), model 128 (best non-Google model over the long sample), model 148 (best non-Google model over the short sample) and the

¹²See section 5.3 for details on these models.

three non-linear models, at the start of the recession all models seem to perform quite similarly. As soon as the recession starts to hit with Lehman Brothers' bankruptcy all the non-linear models and the non-Google model estimated over the long sample start to under-predict the unemployment rate, while the non-Google model estimated over the short sample tends to over-predict the unemployment rate. Instead the model using the GI manages to produce the best predictions with the lowest forecast errors. After the end of the recession, all models seem to fair similarly well, except for non-linear models which alternate periods of under-prediction with moments of over-predictions. Nevertheless, the best model using the GI still has a forecast error which is the closest to the zero line. A similar picture arises from the middle and the bottom panel where we depict the forecast errors for the same models at two and three steps ahead, respectively. For forecast horizons longer than one month, when the recession starts to intensify, non-linear models and the non-Google model estimated over the long sample tend to under-predict even further, while the non-Google model estimated over the short sample severely over-predicts.

These results point unambiguously to the predictive power of leading indicators based on Google data, with the advantage over standard time series models increasing with the length of the forecast horizon. In subsection 4.2 we discuss the results of formal tests of equal forecast accuracy and superior predictive ability to disentangle the best models in terms of forecasting performance.

4.2 Formal tests of forecast accuracy

The literature on US unemployment forecasting has thus far only considered the ratios of the mean squared errors between a competitor model and a benchmark model to evaluate each model's forecast ability. Nevertheless, after the seminal papers by Diebold and Mariano (1995) and West (1996), the community of forecasters has increasingly understood the importance of correctly testing for out-of-sample equal forecast accuracy. West (2006) provides a recent survey of the tests of equal forecast accuracy, while Busetti and Marcucci (2013) provide extensive Monte Carlo evidence on the best tests of equal forecast accuracy or forecast encompassing to be used by the practitioners in any specific forecasting framework. To provide a more formal assessment of the forecasting properties of

each model in our horse-race, we use the best model in terms of lowest MSE as the benchmark model and perform two tests. The first is a two-sided DM test for the null of equal forecast accuracy between the benchmark and the competitor.¹³ We use the two-sided version of the DM test because some models are nested and others are non-nested making the direction of the alternative hypothesis unknown. Using the two-sided version of the test we can thus compare both nested and non-nested models, as is our case where the exogenous variable often differs from one model to another and only a subset of models are really nested. Furthermore, we use the DM because it can be compared to standard critical values of the Gaussian distribution.

From Table 3 we can see that the best model in terms of the lowest MSE always beats the non-linear competitors estimated over the long sample in predicting the unemployment rate and almost always outperforms when compared to models not using the GI and estimated over the short sample. The DM test cannot reject the null of equal forecast accuracy only when the best model using the GI is compared to models estimated over the long sample (and thus using an in-sample that is 10 times bigger). However, we have to highlight the fact that being the simplest test of equal forecast accuracy, the DM is also the least powerful test that could have been employed. Therefore, even in this case we have been rather conservative. Had we adopted a more powerful test than the DM, we could have had even better results with much more frequent rejections of the null of equal forecast accuracy between our benchmark model which uses the GI and the competitors.

However, the DM test is only based on a pairwise comparison of forecasting models where one model is selected as the benchmark. Since we are comparing a large number of model-based forecasts we should account for all the possible pairwise comparisons using a test based on multiple comparisons. In order to be agnostic also on the choice of the benchmark we decided to compare the whole set of models jointly with the MCS test suggested by Hansen et al. (2011), a test based on multiple comparisons that does not imply the choice of a benchmark model. The MCS is in fact defined as the set that

¹³The DM test is based on the loss differential between the benchmark (model 0) and the k -th competitor, i.e. $d_t = e_{0,t}^2 - e_{k,t}^2$, where $e_{k,t}$ is model k 's forecast error and $e_{0,t}$ is the benchmark model's forecast error. To test the null of equal forecast accuracy $H_0 : E(d_t) = 0$, we employ the DM statistic $DM = P^{1/2} \bar{d} / \hat{\sigma}_{DM}$, where $\bar{d}$ is the average loss differential, P is the out-of-sample size, and $\hat{\sigma}_{DM}$ is the square-root of the long-run variance of d_t . Under the null, the DM test is distributed as a Gaussian.

contains the best models in terms of superior forecast accuracy without any assumptions about the true (benchmark) model. The MCS allows the researcher to identify, from a universe of model-based forecasts, a subset of models, equivalent in terms of superior ability, which outperform all the other competing models at a given confidence level α . The other thing we should note is that the MCS is a test of conditional predictive ability. As such, it allows a unified treatment of nested and non-nested models taking into account estimation technique, parameter uncertainty, ratio of estimation and evaluation sample, and data heterogeneity.¹⁴

The MCS results are reported in the last column of each panel of Table 3 for every forecast horizon. A 1 indicates that the model in the row is included in the final MCS, while a 0 means that the model is otherwise not included. We set the confidence level for the MCS to $\alpha = 0.05$, the block length to 10 and the number of bootstrap samples to 300. Such number might appear small but it is sufficient to identify the MCS. We did not choose a bigger number because using the range statistic we are comparing all possible pairwise combinations between model i and j and given the large number of models in our forecasting exercise a higher number of bootstrap samples would make the computation of the test more cumbersome. Looking at Table 3 at 1-, 2-, and 3-month-ahead forecast horizons, we can notice that the final MCS always includes all the best 15 models using G1 as the leading indicator at all forecast horizons. We can also notice that the group of best 15 models is largely dominated by Google-based models at all forecast horizons. Table 4 shows the number of models selected in the final MCS by category (Google, No Google, Short and Long Sample). From the left panel of Table we can notice that around

¹⁴Let us denote the initial set of k -step-ahead forecasts $\mathcal{M}^0 : \{f_{i,t+k} \in \mathcal{M}^0 \ \forall i = 1, \dots, M\}$, where $t = 0, 1, \dots, T-1$, T is the out-of-sample size and M is the number of models. The starting hypothesis is that all forecasts in the set $\mathcal{M}^0$ have equal forecasting performance, measured by a loss function $L_{i,t} = L(u_t, f_{i,t})$, where u_t is the unemployment rate and $f_{i,t}$ is the corresponding forecast at time t from model i . Let $d_{ij,t} = L_{i,t} - L_{j,t} \ \forall i, j = 1, \dots, M$ define the relative performance of forecast i and j . The null hypothesis for the MCS test takes the form $H_{0,\mathcal{M}^0} : E(d_{ij,t}) = 0 \ \forall i, j = 1, \dots, M$. We use the ‘range’ statistic defined as $T_R = \max_{i,j \in \mathcal{M}} |t_{ij}|$ where $t_{ij} = \bar{d}_{ij} / \sqrt{\hat{v} \hat{a}r(\bar{d}_{ij})}$ represents the standardized relative performance of forecast i with respect to forecast j , and $\bar{d}_{ij} = T^{-1} \sum_{t=1}^T d_{ij,t}$ is the sample average loss difference between forecast i and j . To obtain the distribution under H_0 a stationary bootstrap scheme is used. If H_0 is rejected, an elimination rule removes the forecast with the largest t_{ij} . This process is repeated until non-rejection of the null occurs, thus allowing the construction of $(1 - \alpha)$ -confidence set for the best forecasts in $\mathcal{M}^0$.

a quarter of the models using the GI is included in the final MCS for this in-sample at 1-step-ahead. Google-based models make up almost half of the final MCS at 2-step-ahead and one-third at 3-step-ahead. Again, we believe that these results are indeed astonishing given that Google-based models use only a limited amount of information compared to non-Google models estimated over the long sample.

5 Robustness checks

In this Section we provide the following *robustness checks* for the main results presented so far: (i) We vary the out-of-sample intervals for the forecast evaluation showing that main results hold when starting the forecast evaluation interval before, during and after the Great Recession; (ii) we repeat the forecast horse race for each of the 50 US states plus DC rather than at the federal level; (iii) we test the performance of alternative non-linear models not employing Google data; (iv) we test the forecasting properties of two alternative, and less popular, job-search-related keywords; and (v) we provide a falsification test. All these tests confirm, directly or indirectly, the very good performance of Google-augmented models when forecasting the monthly US unemployment rate.

5.1 Different in-sample/out-of-sample

As a first robustness check we compare the forecasting properties of our preferred models which adopt the GI as the leading indicator across different combinations of in-sample and out-of-sample periods. The rationale behind this is to check the robustness of our results to different business cycle conditions. This is of particular interest given that our out-of-sample includes the onset of the Great Recession; in which the unemployment rate sharply increased by about four percentage points; and the subsequent period of slow growth and high, but rather stable, unemployment. Choosing appropriate out-of-samples for our forecast comparison, we can test whether the superior performance of Google-augmented models is due to a good performance during a peculiar time period, or if its predictive ability is confirmed throughout the business cycle.

In particular, we conduct the forecast comparison of subsection 4.1 on two alterna-

tive out-of-samples: One starting with the NBER recession following the bankruptcy of Lehman (2008.10-2011.6) and another one starting with the end of that recession (2009.7-2011.6). Results of the forecast horse race, reported in Table 5, confirm the superior predictive performance of Google-based models: In both sub-samples, models including the indicator of internet job-search activity always show lower MSE at one, two and three steps ahead. Compared to the basic results presented in subsection 4.1, the gap in favor of Google-based models actually increases when considering these two different out-of-sample intervals: The best 10 models in terms of lowest MSE always include the GI, irrespective of the out-of-sample and of the forecast horizon.

Even with respect to the final MCS, Google-based models tend to outperform the others. Looking at Table 5, we can notice that the final MCS always includes the best 15 models adopting G1 as the leading indicator across all forecast horizons. Looking at the number of models selected in the final MCS, from the middle panel of Table 4 we can notice that around two thirds of the models using G1 are included in the final MCS for the in-sample terminating right after Lehman bankruptcy at all forecast horizons. This highlights the power of Google data to help forecast the unemployment rate when the business climate is particularly pessimistic and when having good forecasts matter the most. For the last in-sample terminating at the end of the last recession we can notice that around a quarter of the models using G1 are included in the final MCS across all forecast horizons. Again, even with such a short out-of-sample almost 25% of the best models entering in the final MCS use the GI for “*jobs*”.

5.2 State level forecasts

As an additional robustness check for the predictive properties of the GI, we estimated the same 520 linear models for each of the 50 states plus DC, assessing the percentage of states for which the best model in terms of lower MSE is the one using the GI. The descriptive statistics for the monthly unemployment rate, the IC and the GI for each state are in line with those discussed for the US and are not reported for the sake of brevity but are available on request.

In Table 6 we report for each state the best forecast obtained without Google (both on

the long and the short sample) and with the GI based on the keyword “*jobs*”. As in the previous cases, the forecast comparison takes place at 1, 2 and 3 steps ahead and over the out-of-sample 2007.2-2011.6, the baseline in our forecast comparison. The percentage of best 5 models adopting the GI as a leading indicator is equal to 70.2% when forecasting at one step ahead, and 62.0% at two steps ahead. Only when forecasting at three steps ahead does the percentage of states for which the best model includes the GI fall below 50% (to 39.2%, to be precise).

5.3 Nonlinear models

Most of the previous literature on unemployment forecasting in the US suggests using non-linear models to better approximate the long-term dynamic structure of its time series (see Montgomery et al., 1998 and Rothman, 1998). In particular, Montgomery et al. (1998) argue that Threshold Autoregressive (TAR) models can better approximate the unemployment rate dynamics especially during economic contractions, while linear ARMA models generally give a better representation of its short-term dynamics. To test the predictive ability of our best models which use the GI, we also included in the forecast comparison some non-linear models which are typically used in the literature. We have estimated three non-linear time series models. The first is a self-exciting threshold autoregression (SETAR) model which takes the following form:

$$u_t = [\phi_{01} + \phi_{11}u_{t-1} + \phi_{21}u_{t-2}] I(u_{t-1} \leq c) + [\phi_{02} + \phi_{12}u_{t-1} + \phi_{22}u_{t-2}] I(u_{t-1} > c) + \varepsilon_t \quad (4)$$

where $I(.)$ is the indicator function and c is the value of the threshold.

The SETAR models endogenously identify two different regimes given by the threshold variable u_{t-1} . In particular, following Rothman (1998) we adopted a SETAR model with two lags for each regime.

The second non-linear model used to forecast the unemployment rate is a logistic smooth transition autoregressive (LSTAR) model which is a generalization of the SETAR.

The LSTAR model takes the form

$$u_t = [\phi_{01} + \phi_{11}u_{t-1} + \phi_{21}u_{t-2}] [1 - G(\gamma, c, u_{t-1})] + [\phi_{02} + \phi_{12}u_{t-1} + \phi_{22}u_{t-2}] G(\gamma, c, u_{t-1}) + \varepsilon_t \quad (5)$$

where $G(\gamma, c, u_{t-1}) = [1 + \exp(-\gamma \prod_{k=1}^K (u_t - c_k))]^{-1}$ is the logistic transition function, $\gamma > 0$ is the slope parameter and c is the location parameter. In this model the change from one regime to the other is much smoother than in the SETAR model.

The third non-linear model employed to predict the US unemployment rate is an additive autoregressive model (AAR) of the following form

$$u_t = \mu + \sum_{i=1}^m s_i(u_{t-(i-1)d}) + \varepsilon_t \quad (6)$$

where s_i are smooth functions represented by penalized cubic regression splines. The AAR model is a generalized additive model that combines additive models and generalized linear models. These models maximize the quality of prediction of a target variable from various distributions, by estimating a non-parametric function of the predictor variables which are connected to the dependent variable via a link function (see Hastie and Tibshirani, 1990). We have included this additional model to enlarge our out-of-sample comparison to non-parametric models which were found to be superior in predicting the US unemployment rate by Golan and Perloff (2004).

Panel E of Table 3 reports the MSE, DM test and MCS test for 1- to 3-month-ahead forecasts from these three non-linear models estimated only up to the second lag for the long sample (in-sample 1967.1-2007.2, out-of-sample 2007.3-2011.6). At 1-month ahead the best non-linear model is the SETAR(2) which ranks 402nd, the second best is the LSTAR(2)(424th) and the third best is the AAR(2) (441st). Results do not improve at longer forecast horizons, and in particular these non-linear models are never included in the MCS except at one-step-ahead for the out-of-sample starting at the end of the most recent NBER recession (see right panel of Table 4). In addition, the DM test always rejects the null of equal forecast accuracy. We can thus conclude that our simple linear

model using our preferred leading indicator (GI) outperforms standard non-linear models estimated over the long sample across all forecast horizons.

5.4 Alternative keywords

As a further robustness check we analyze the properties of our forecasting models using not only our preferred GI for “*jobs*”, but also other keywords that are closely related to job search. In particular we look at the GIs for “*collect unemployment*” and “*job center*” (respectively labeled G2 and G3). As already discussed in Section 2 the volume of searches underlying these two keywords is much smaller compared to that for “*jobs*” (see Figure 1), but nevertheless it is interesting to check whether even in this case, models augmented with Google data are still good at predicting the unemployment rate. In Figure 6 we plot the dynamics of the monthly GIs along with the monthly US unemployment rate; visual inspection reveals a similar pattern for these two alternative leading indicators and the time series we are forecasting. The two keywords are very highly correlated with the contemporaneous unemployment rate (0.97 and 0.96, respectively). The descriptive statistics for each of the two indexes, both at the monthly and the weekly level, are reported in Table 1.

In Table 7 we show the results of pairwise forecast comparisons for each keyword, identical to the ones conducted for the main keyword “*jobs*”. When using these alternative and less-representative keywords the forecast performance deteriorates compared with our preferred keyword. Google-augmented models estimated on the short sample are now never able to improve the forecasting performance of non-Google models estimated on the long sample. Nevertheless, when conducting the comparison among models estimated on the same short-interval, many best models are augmented with Google data. In particular, the best model at one-step-ahead includes the GI for “*collect unemployment*”; models augmented with the GI for the keyword “*job center*” always outperform non-Google models, at all forecast horizons. However, using the GI for these two keywords does not add that much to the forecasting performance of these models. For example, in the final MCS only a few Google-models (around 10%) are selected (see Table 4).

As a final step, we extract the first principal component (labeled G5) of the three

Google indices analyzed so far, and we test the forecasting performance of this last leading indicator. This by construction summarizes all the information in the three leading indicators maximizing their variance. We get very interesting results: When combining all the information of the three indices in one leading indicator we get the best forecasting performances overall. At one-step-ahead the best model now becomes an ARX(1) with one lag of the IC and the G5 for the fourth week. Its MSE is lower than the best models based on the GI for “*jobs*”, and thus also of the best non-Google models. Compared to the best Google models exploiting the GI for “*jobs*”, gains in terms of lower MSE range between 7.2% (at three-step-ahead) and 18.3% (at two-step-ahead). Even with these models which adopt the Google factor all the best 15 models always enter the final MCS. Furthermore, as we can see from Table 4, the percentage of Google-based models which enter the final MCS is 80%, 45% and 35% at one-, two- and three-step-ahead, respectively. We find similar percentages with the other two in-samples that terminate in the middle and at the end of the recent recession.

5.5 Falsification test

In this section we provide a falsification test for the main results of this paper. In particular, we test the forecasting power of an alternative Google-based indicator, that is chosen to be the one with the highest correlation with the time series of the monthly US unemployment rate in the in-sample, but is not necessarily related to job search. We can identify this keyword thanks to the fact that Google has developed a new application, called ‘Google Correlate’¹⁵ able to identify, within a specified time interval, the web searches for keywords that (i) show the highest correlation with a given keyword search, and (ii) show the highest correlation with a given time series. In particular, we isolated the time series of the US monthly unemployment rate and we used this application to find the keyword search that, among all searches conducted through the search engine, was mostly correlated with it in our in-sample (2004.1-2007.2). We found that this series was the GI for the keyword ‘*dos*’, showing a correlation with the US unemployment rate

¹⁵Available at www.google.com/trends/correlate/. See Mohebbi et al., 2011 for details on this application.

of 0.98 in the relevant in-sample, but otherwise with no logical connection to job-search: ‘*dos*’ is an acronym for the US Department Of State or for Disk Operating System. We use this alternative web-search indicator (labeled as G4 in Table 8) to conduct a horse-race forecast comparison that is identical to the main one, whose results were presented in Section 4.1. Looking at Table 8, we can see that models augmented with this fake GI indicator never rank among the best 15 models of the forecast comparison across all forecast horizons (1-, 2-, and 3-step-ahead), providing indirect evidence for the relevance of the web-search data when the underlying keywords have a direct link with job search.

6 Comparison with the Survey of Professional Fore- casters

As an final robustness check, we compare the forecasts of our best model with the results of the Survey of Professional Forecasters (SPF), a quarterly survey of about 30 professionals, conducted by the Federal Reserve Bank of Philadelphia.¹⁶ The survey publishes estimates of the quarterly evolution of a set of macroeconomic variables approximately in the middle of the quarter.¹⁷ We construct three time series of predictions based on SPF data: One obtained with the best forecast¹⁸ in each quarter (SPF^{best}), one with the mean of the forecasts (SPF^{mean}), and one with the median (SPF^{median}). Conditioning on the same information set, we compare these forecasts with the ones obtained by six different models, chosen among those with the best forecasting performance. We define these best models as (i) our best model overall (the one using the GI); (ii) the best model among those not using the GI (NGI_L) over the long sample; and (iii) the best model among those not using the GI over the short sample (NGI_S). To these three groups of best models we add three additional groups of non-linear models based on (i) the SETAR(2), (ii) the LSTAR(2) and (iii) the AAR(2) model.

From each model x we compute three series of quarterly forecasts: 1) $x^{1st-month}$ are

¹⁶<http://www.phil.frb.org/research-and-data/real-time-center/survey-of-professional-forecasters/>.

¹⁷The SPF is issued around the 15th of February, May, August and November of each year.

¹⁸The best individual forecast is calculated ex-post once the real values for u_t are known.

the 1-month-ahead forecasts computed in the last month of each quarter before the one we want to forecast.¹⁹ The prediction for the whole quarter is equal to the forecast for the first month of the quarter. 2) $x^{2nd-month}$ are the 2-month-ahead forecasts computed in the last month of the quarter before, with the estimate for the whole quarter being equal to the estimate for the second, central, month. Both these forecasts are very conservative with respect to those of SPF, since the SPF is issued on the 15th of the second month of each reference quarter, thus around 45 days after our estimates are produced. Finally, 3) x^{Comb} are the quarterly forecasts computed as the average of the realized unemployment rate for the first month and the 1- and 2-month-ahead forecasts generated at the end of the first month of the reference quarter. These latter forecasts are less conservative because they use all the information available in the first month of the quarter; nevertheless, they still exploit the same information set available to the Professional Forecasters at the time of the Survey.

Does our model with Google outperform the professionals? It does, by a considerable margin, if we consider that it only uses a very short sample. In Table 9 we report the MSE for the eighteen best models and the three aggregations of SPF forecasts over the period 2007Q2-2011Q2 along with the DM test where the benchmark is either the best model, that is the model with the lowest MSE (in boldface), $G^{1st-month}$ or $G^{2nd-month}$. It is evident that the model including the GI, and exploiting the same information set (G_{comb}) outperforms all the three SPF forecasts, having a MSE that is about two thirds lower than the best SPF forecast (SPF_{median}). The DM test shows that the benchmark model (G_{comb}) is significantly better than all the other competitors except for the best non-Google time series models.

Figure 7 depicts the forecast errors from the best six models (those with the lowest MSE in Table 9) in addition to the mean and median SPF forecasts. It is rather clear that the model including the GI has the best performance in most periods, and in particular when the current recession worsened after the Lehman collapse in 2008Q4. The model including the GI tends to give forecast errors that are close to zero, while both the mean

¹⁹For example, if we want to forecast the quarterly unemployment rate for 2008Q2, at 2008.3 we compute the 1-month-ahead forecast from one of our three best models.

and the median of the SPF tend to under-predict the real unemployment rate. This means that our simple linear ARMA models with the GI as a leading indicator outperforms the predictions of the professional forecasters also during contractions, when the social impact of a high unemployment rate is even greater and the loss attached to high and positive forecast errors is maximal.

7 Conclusions

In this paper we suggest the use of the Google index (GI), based on internet job searches performed through Google, as the best leading indicator to predict the US monthly unemployment rate.

Popular time series specifications augmented with this indicator definitely improve their out-of-sample forecasting performance at one-, two- and three-month horizons. Our results from the out-of-sample horse-race with more than five hundred linear and non-linear specifications show that the best models in terms of lowest MSE are always those using the GI as the leading indicator. These models also fair better in comparison to other similar models estimated on the same or longer time spans and using the initial claims (IC) as a leading indicator. Our results hold when the forecast comparison takes place over an out-of-sample that starts before, during and after the Great Recession, and hold also at the state rather than at the federal level. Conditioning on the same information set, the best Google-augmented predictions also outperform the forecasts released in the Survey of Professional Forecasters conducted by the Philadelphia Fed.

Notwithstanding its limited time availability (Google data are available since January 2004) we believe that the GI should routinely be included in time series models to predict unemployment dynamics. We fully expect that the use of internet-based data will become widespread in economic research in the near future.

References

- Aparicio, F., A. Escribano, and A. Garcia** 2006. “Range Unit-Root (RUR) Tests: Robust against Nonlinearities, Error Distributions, Structural Breaks and Outliers”, *Journal of Time Series Analysis*, 27(4): 545–576.
- Artola, C. and Galan, E.** 2012. “Tracking the Future on the Web: Construction of Leading Indicators using Internet searches”, *Documentos Ocasionales*, (1203), Bank of Spain.
- Askitas, N. and Zimmermann, K. F.** 2009. “Google Econometrics and Unemployment Forecasting”, *Applied Economics Quarterly*, 55: 107–120.
- Baker, S. and Fradkin, A.** 2011. “What Drives Job Search? Evidence from Google Search Data”, *Discussion Paper No. 10-020*, Stanford Institute for Economic Research.
- Barnichon, R. and Nekarda, C. J.** 2012. “The Ins and Outs of Forecasting Unemployment: Using Labor Force Flows to Forecast the Labor Market”, *Brookings Papers on Economic Activity*, forthcoming.
- Billari, F. and D’Amuri, F. and Marcucci, J.** 2012. “Forecasting Births Using Google”, *Oxford University*, mimeo.
- Busetti, F. and Marcucci, J.** 2013. “Comparing Forecast Accuracy: A Monte Carlo Investigation”, *International Journal of Forecasting*, 29: 13–27.
- Choi, C.Y., and Y.K. Moh** 2007. “How useful are tests for unit-root in distinguishing from stationary but non-linear processes?”, *Econometrics Journal*, 10: 82–112.
- Choi, H. and Varian, H.** 2012. “Predicting the Present with Google Trends”, *Economic Record*, 88: 2–9.
- Da, Z. and Engelberg, J. and Pengjie, G.** 2011. “In Search of Attention”, *Journal of Finance*, 5: 1461–1499.
- D’Amuri, F.** 2009. “Predicting unemployment in short samples with internet job search query data”, MPRA working paper n. 18403.
- DeLong, J. B. and Summers, L. H.** 1986. “Are Business Cycles Symmetrical?”, in *The American Business Cycle, Continuity and Changes*, ed. R. J. Gorton, Chicago: University of Chicago Press for NBER.

- Diebold, F. X. and Mariano, R. S.** 1995. “Comparing Predictive Accuracy”, *Journal of Business & Economic Statistics*, 13: 253–263.
- Edelman, B.** 2012. “Using Internet Data for Economic Research”, *Journal of Economic Perspectives*, 26(2): 189–206.
- Ginsberg, J., Mohebbi, M. H., Patel, R. S., Brammer L., Smolinski, M. S. and Brilliant L.** 2009. “Detecting Influenza epidemics using Search Engine Query Data”, *Nature*, 457: 1012–1014.
- Golan, A. and Perloff, J. M.** 2004. “Superior Forecasts of the U.S. Unemployment Rate Using a Nonparametric Method”, *The Review of Economics and Statistics*, February, 86(1): 433–438.
- Hansen, P. R.** 2005. “A Test for Superior Predictive Ability”, *Journal of Business and Economic Statistics*, 23: 365–380.
- Hansen, P. R., Lunde A. and Nason J. M.** 2011. “The Model Confidence Set”, *Econometrica*, 79(2): 453–497.
- Hastie, T. J., Tibshirani, R. J.** 1990. “Generalized Additive Models”, Chapman and Hall Ltd., London.
- Koop, G., and Potter, S. M.** 1999. “Dynamic Asymmetries in U.S. Unemployment”, *Journal of Business and Economic Statistics*, 17(3): 298–312.
- McLaren, N. and Shanbhorge, R.** 2011. “Using Internet Data as Economic Indicators”, *Bank of England Quarterly Bulletin*, Second Quarter.
- McQueen, G., and Thorley, S.** 1993. “Asymmetric Business Cycle Turning Points”, *Journal of Monetary Economics*, 31: 341–362.
- Montgomery, A. L., Zarnowitz, V., Tsay, R. S. and Tiao, G. C.** 1998. “Forecasting the U.S. Unemployment Rate”, *Journal of the American Statistical Association*, June, 93(442): 478–493.
- Neftci, S. N.** 1984. “Are Economic Time Series Asymmetric Over the Business Cycles?”, *Journal of Political Economy*, 85: 281–291.
- Proietti, T.** 2003. “Forecasting the US Unemployment Rate”, *Computational Statistics & Data Analysis*, 42: 451–476.
- Rothman, P.** 1998. “Forecasting Asymmetric Unemployment Rates”, *The Review of*

Economics and Statistics, February, 80(1): 164–168.

Sichel, D. E. 1993. “Business Cycle Asymmetry: A Deeper Look”, *Economic Enquiry*, 31: 224–236.

Suhoy, T. 2009. “Query Indices and a 2008 Downturn ”, *Bank of Israel Discussion Paper (2009.06)*.

Troy, G. and D. P. and D. S. 2012. “Electronic Indicators of Economic Activity”, *Australian Reserve Bank - Economic Bulletin*, June 1–12.

Wallis, K. 1987. “Time Series Analysis of Bounded Economic Variables”, *Journal of Time Series Analysis*, 8: 115–123.

West, K. D. 1996. “Asymptotic inference about predictive ability”, *Econometrica*, 64: 1067–1084.

West, K. D. 2006. “Forecast Evaluation”, 100-134, in *Handbook of Economic Forecasting*, Vol. 1, G. Elliott, C.W.J. Granger and A. Timmerman (eds), Amsterdam: Elsevier.

Table 1: Descriptive statistics: sample 2004.1-2011.6

	Mean	Median	Max	Min	Std. Dev.	Skew.	Kurt.	Jarque-Bera	Obs.
u_t	6.542	5.448	10.147	4.391	2.091	0.613	1.597	13.021***	90
IC_t	1581.0	1396.5	2580.0	1152.0	372.8	1.074	3.138	17.374***	90
$IC_{W1,t}$	394.6	352.5	659.0	282.0	94.1	1.101	3.262	18.435***	90
$IC_{W2,t}$	394.9	349.5	650.0	289.0	94.3	1.070	3.192	17.304***	90
$IC_{W3,t}$	395.5	354.5	655.0	298.0	90.7	1.072	3.233	17.430***	90
$IC_{W4,t}$	395.9	352.0	642.0	283.0	96.9	0.999	2.918	15.007***	90
$G1_t$	68.8	65.2	83.8	56.2	9.2	0.237	1.462	9.711***	90
$G1_{W1,t}$	67.7	64.3	84.1	55.0	9.4	0.360	1.542	9.924***	90
$G1_{W2,t}$	68.3	66.0	89.4	55.7	8.9	0.398	1.852	7.314**	90
$G1_{W3,t}$	69.6	66.9	91.3	55.6	10.0	0.337	1.745	7.609**	90
$G1_{W4,t}$	69.2	65.8	88.5	54.5	9.6	0.308	1.604	8.735**	90
$G2_t$	32.3	17.1	81.5	1.9	25.6	0.619	1.633	12.746***	90
$G2_{W1,t}$	31.5	16.4	79.1	1.4	25.4	0.718	1.803	13.109***	90
$G2_{W2,t}$	31.4	18.5	76.8	-7.4	25.5	0.541	1.700	10.610***	89
$G2_{W3,t}$	33.5	17.4	91.2	-4.5	27.5	0.597	1.721	11.490***	90
$G2_{W4,t}$	33.0	19.0	81.9	-0.9	26.1	0.581	1.632	12.083***	90
$G3_t$	57.9	51.2	84.3	40.0	14.5	0.660	1.818	11.780***	90
$G3_{W1,t}$	57.8	52.6	88.0	33.5	14.7	0.561	1.955	8.819**	90
$G3_{W2,t}$	56.9	50.6	91.5	35.9	15.4	0.594	1.987	9.135**	90
$G3_{W3,t}$	58.4	53.7	89.7	25.0	15.7	0.493	1.884	8.321**	90
$G3_{W4,t}$	58.4	52.1	87.3	41.3	14.7	0.672	1.855	11.695***	90
$G4_t$	51.7	48.5	85.0	26.5	18.1	0.272	1.654	7.905**	90
$G4_{W1,t}$	51.7	50.3	86.0	25.1	18.0	0.233	1.688	7.269**	90
$G4_{W2,t}$	51.7	49.1	82.9	24.9	18.2	0.273	1.695	7.502**	90
$G4_{W3,t}$	51.6	48.2	86.7	25.3	18.6	0.274	1.676	7.699**	90
$G4_{W4,t}$	51.9	49.5	82.4	27.6	18.3	0.262	1.594	8.443**	90

Notes: u_t is the US monthly unemployment rate in levels. IC indicates the monthly initial claims, while $G1$, $G2$, $G3$, $G4$, and $G5$ are the monthly averages of the weekly Google indexes for keywords ‘jobs’, ‘collect unemployment’, ‘job center’, ‘dos’ (the false index), and the first principal component of the first three Google indexes used as leading indicators. The subscripts Wj indicate the j^{th} week. ***, ** and * indicate rejection of the null of normality at 1, 5 and 10%, respectively.

Table 2: Forecasting Models: $\phi(L)y_t = \mu + x'_t\beta + \theta(L)\varepsilon_t$ for the unemployment rate

Long sample: 1967.1-2011.6														Short Sample: 2004.1-2011.6															
AR(1) # AR(2) # ARMA(1,1) # ARMA(2,2) # ARMA(1,1) # ARMA(2,2) #														AR(1) # AR(2) # ARMA(1,1) # ARMA(2,2) #															
w/o LI														w/o LI															
u_{t-1} 1 u_{t-k} 1 $u_{t-1}, \varepsilon_{t-1}$ 1 $u_{t-k}, \varepsilon_{t-k}$ 1 u_{t-1} 1 u_{t-k} 1 $u_{t-1}, \varepsilon_{t-1}$ 1 $u_{t-k}, \varepsilon_{t-k}$ 1														u_{t-1} 1 u_{t-k} 1 $u_{t-1}, \varepsilon_{t-1}$ 1 $u_{t-k}, \varepsilon_{t-k}$ 1 u_{t-1} 1 u_{t-k} 1 $u_{t-1}, \varepsilon_{t-1}$ 1 $u_{t-k}, \varepsilon_{t-k}$ 1															
w/ LI x_t														w/ LI x_t															
(t)																													
IC		✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1
IC _{wj}		✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4
G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC, G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC _{wj} , G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
(t-1)																													
IC		✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1
IC _{wj}		✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4
G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC, G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC _{wj} , G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
(t-2)																													
IC		✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1
IC _{wj}		✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4
G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC, G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC _{wj} , G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
(t-3)																													
IC		✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1	✓	1
IC _{wj}		✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4	✓	4
G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC, G		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
IC _{wj} , G _{wj}		-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
$j = 1, \dots, 4; k = 1, 2 - w / \text{ or } w / \text{ o SAR / SMA}$																													

$j = 1, \dots, 4; k = 1, 2 - w/$ or w/o SAR/SMA

Notes: # indicates the number of models in each group. The subscript $wj, j = 1, \dots, 4$ denotes the weekly leading indicators. A ✓ denotes that the model in that group adopts the row variable as a leading indicator. A - indicates that the model does not adopt the row variable as leading indicator. For models with both IC and the GI at the weekly level we have 5 models at each lag because we consider also a model with all the weekly ICs and GIs.

Table 3: Results for US unemployment rate in levels (u_t) - forecasting with AR(1) auxiliary model. Rolling scheme. Best 15 models in terms of lowest MSE: GI for “jobs” (G^1)

1-step-ahead					2-step-ahead					3-step-ahead				
Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS
Best 15 models with GI - G1 - Out-of-Sample: 2007.3-2011.6														
Panel A1: Best models overall					Panel A2: Best models overall					Panel A3: Best models overall				
493 ARMAX(2,2) - $IC_{w4,t-1} - G^1_{w4,t-1}$	0.026	1	0.00	1	302 ARX(1) - $IC_{t-1} - G^1_{t-1} - SA$	0.060	1	0.00	1	302 ARX(1) - $IC_{t-1} - G^1_{t-1} - SA$	0.111	1	0.00	1
296 ARX(1) - $IC_{t-1} - G^1_{t-1}$	0.027	2	0.03	1	368 ARX(2) - $IC_{t-1} - G^1_{t-1} - SA$	0.060	2	0.08	1	368 ARX(2) - $IC_{t-1} - G^1_{t-1} - SA$	0.114	2	0.17	1
362 ARX(2) - $IC_{t-1} - G^1_{t-1}$	0.027	3	0.09	1	362 ARX(2) - $IC_{t-1} - G^1_{t-1}$	0.060	3	0.10	1	362 ARX(2) - $IC_{t-1} - G^1_{t-1}$	0.122	3	0.64	1
274 ARX(1) - $IC_t - G^1_t$	0.027	4	0.19	1	296 ARX(1) - $IC_{t-1} - G^1_{t-1}$	0.061	4	0.60	1	412 ARMAX(1,1) - $IC_t - G^1_t - SA$	0.123	4	0.66	1
483 ARMAX(2,2) - $G^1_{w4,t-1}$	0.027	5	0.37	1	295 ARX(1) - $IC_{w4,t-1} - G^1_{w4,t-1}$	0.063	5	0.54	1	296 ARX(1) - $IC_{t-1} - G^1_{t-1}$	0.123	5	1.71*	1
295 ARX(1) - $IC_{w4,t-1} - G^1_{w4,t-1}$	0.027	6	0.31	1	276 ARX(1) - $IC_{w4,t-1} - G^1_{w4,t-1} - SA$	0.064	6	0.50	1	390 ARX(2) - $IC_{t-2} - G^1_{t-2} - SA$	0.127	6	1.09	1
128 ARMAX(2,2) - $IC_{t-2} - SA$	0.028	7	0.29	1	301 ARX(1) - $IC_{w4,t-1} - G^1_{w4,t-1} - SA$	0.065	7	0.69	1	472 ARMAX(2,2) - $IC_t - G^1_t$	0.129	7	0.90	1
340 ARX(2) - $IC_t - G^1_t$	0.028	8	0.30	1	383 ARX(2) - $IC_{w4,t-2} - G^1_{w4,t-2}$	0.065	8	0.72	1	291 ARX(1) - $G^1_{t-1} - SA$	0.129	8	0.95	1
280 ARX(1) - $IC_t - G^1_t - SA$	0.028	9	0.44	1	274 ARX(1) - $IC_t - G^1_t$	0.065	9	1.15	1	313 ARX(1) - $G^1_{t-2} - SA$	0.130	9	0.93	1
302 ARX(1) - $IC_{t-1} - G^1_{t-1} - SA$	0.028	10	0.40	1	307 ARX(1) - $G^1_{w4,t-2}$	0.065	10	0.48	1	307 ARX(1) - $G^1_{w4,t-2}$	0.131	10	0.86	1
270 ARX(1) - $IC_{w1,t} - G^1_{w1,t}$	0.028	11	0.42	1	390 ARX(2) - $IC_{t-2} - G^1_{t-2} - SA$	0.066	11	0.85	1	406 ARMAX(1,1) - $IC_t - G^1_t$	0.131	11	0.91	1
122 ARMAX(2,2) - $IC_{w4,t-2}$	0.028	12	0.42	1	280 ARX(1) - $IC_t - G^1_t - SA$	0.066	12	1.41	1	312 ARX(1) - $G^1_{w4,t-2} - SA$	0.131	12	0.86	1
286 ARX(1) - G^1_{t-1}	0.028	13	0.44	1	122 ARMAX(2,2) - $IC_{w4,t-2}$	0.066	13	0.63	1	280 ARX(1) - $IC_t - G^1_t - SA$	0.132	13	2.14**	1
273 ARX(1) - $IC_{w4,t} - G^1_{w4,t}$	0.028	14	0.53	1	361 ARX(2) - $IC_{w4,t-1} - G^1_{w4,t-1}$	0.066	14	0.70	1	122 ARMAX(2,2) - $IC_{w4,t-2}$	0.133	14	1.13	1
279 ARX(1) - $IC_{w4,t} - G^1_{w4,t} - SA$	0.028	15	0.54	1	472 ARMAX(2,2) - $IC_t - G^1_t$	0.066	15	0.71	1	324 ARX(1) - $IC_{t-2} - G^1_{t-2} - SA$	0.133	15	2.16**	1
Panel B1: Best models w Google (Short sample)					Panel B2: Best models w Google (Short sample)					Panel B3: Best models w Google (Short sample)				
493 ARMAX(2,2) - $IC_{w4,t-1} - G^1_{w4,t-1}$	0.026	1	0.00	1	302 ARX(1) - $IC_{t-1} - G^1_{t-1} - SA$	0.060	1	0.00	1	302 ARX(1) - $IC_{t-1} - G^1_{t-1} - SA$	0.111	1	0.00	1
296 ARX(1) - $IC_{t-1} - G^1_{t-1}$	0.027	2	0.03	1	368 ARX(2) - $IC_{t-1} - G^1_{t-1} - SA$	0.060	2	0.08	1	368 ARX(2) - $IC_{t-1} - G^1_{t-1} - SA$	0.114	2	0.17	1
362 ARX(2) - $IC_{t-1} - G^1_{t-1}$	0.027	3	0.09	1	362 ARX(2) - $IC_{t-1} - G^1_{t-1}$	0.060	3	0.10	1	362 ARX(2) - $IC_{t-1} - G^1_{t-1}$	0.122	3	0.64	1
Panel C1: Best models w/o Google (Short sample)					Panel C2: Best models w/o Google (Short sample)					Panel C3: Best models w/o Google (Short sample)				
148 ARX(1) - $IC_{w4,t} - SA$	0.032	141	1.47	1	258 ARMAX(2,2) - $IC_{w4,t-2} - SA$	0.083	149	1.78*	0	258 ARMAX(2,2) - $IC_{w4,t-2} - SA$	0.172	149	2.23**	0
143 ARX(1) - $IC_{w4,t}$	0.033	150	1.53	1	178 ARX(2) - $IC_{w4,t} - SA$	0.088	204	2.44**	0	218 ARMAX(1,1) - $IC_{w4,t-1} - SA$	0.190	252	2.72***	0
144 ARX(1) - IC_t	0.033	201	1.74*	1	153 ARX(1) - $IC_{w4,t-1}$	0.088	205	2.68***	0	248 ARMAX(2,2) - $IC_{w4,t-1} - SA$	0.200	270	2.87***	0
Panel D1: Best models w/o Google (Long sample)					Panel D2: Best models w/o Google (Long sample)					Panel D3: Best models w/o Google (Long sample)				
128 ARMAX(2,2) - $IC_{t-2} - SA$	0.028	7	0.29	1	122 ARMAX(2,2) - $IC_{w4,t-2}$	0.066	13	0.63	1	122 ARMAX(2,2) - $IC_{w4,t-2}$	0.133	14	1.13	1
122 ARMAX(2,2) - $IC_{w4,t-2}$	0.028	12	0.42	1	17 ARX(1) - $IC_{w4,t} - SA$	0.068	26	0.72	1	47 ARX(2) - $IC_{w4,t} - SA$	0.135	18	0.99	1
123 ARMAX(2,2) - IC_{t-2}	0.029	18	0.48	1	77 ARMAX(1,1) - $IC_{w4,t} - SA$	0.069	31	0.71	1	17 ARX(1) - $IC_{w4,t} - SA$	0.136	19	1.02	1
Panel E1: Best nonlinear models (Long sample)					Panel E2: Best nonlinear models (Long sample)					Panel E3: Best nonlinear models (Long sample)				
129 SETAR(2)	0.040	402	2.20**	0	129 SETAR(2)	0.127	464	3.04***	0	129 SETAR(2)	0.282	455	3.50***	0
130 LSTAR(2)	0.042	424	2.16**	0	130 LSTAR(2)	0.130	472	2.88***	0	130 LSTAR(2)	0.288	460	3.24***	0
131 AAR(2)	0.043	441	2.28**	0	131 AAR(2)	0.132	475	2.84***	0	131 AAR(2)	0.293	464	3.21***	0

Notes: Long sample: 1967.1-2011.6; short sample: 2004.1-2011.6; out of sample: 2007.3-2011.6. The first column reports model number, **Model** is the model type, **MSE** is the mean squared error, **Rk** is the ranking with respect to the lowest MSE, **DM** is the Diebold and Mariano test for the null hypothesis of equal predictive accuracy (Diebold and Mariano, 1995), and **MCS** is the Model Confidence Set approach by Hansen, Lunde and Nason (2011). G^1_t is the Google Index for keyword ‘jobs’. The column MCS has a 1 when the row model is included in the final model confidence set at 5% confidence level and a 0 otherwise. Models from 1 to 128 are non-Google models estimated over the long sample; models from 129 to 131 are non-linear models estimated over the long sample; models from 132 to 259 are non-Google models estimated over the short sample; models from 260 to 523 are Google models estimated over the sort sample. In all panels ***, ** and * indicate rejection at 1, 5 and 10%, respectively.

Table 4: Number of models in the MCS

	IS: 2004.1-2007.2			IS: 2004.1-2008.9			IS: 2004.1-2009.6		
	OOS: 2007.2-2011.6			OOS: 2008.10-2011.6			OOS: 2009.7-2011.6		
	1-step	2-step	3-step	1-step	2-step	3-step	1-step	2-step	3-step
Models in MCS among those with GI - G1									
No-Google (Long sample)	116	15	43	4	3	9	121	38	6
Non-linear (Long sample)	0	0	0	0	0	0	3	0	0
No-Google (Short sample)	9	0	0	0	0	1	3	0	0
Google (Short sample)	37	12	18	14	7	11	35	13	1
All	162	27	61	18	10	21	162	51	7
Models in MCS among those with GI - G2									
No-Google (Long sample)	118	16	45	41	9	8	121	51	7
Non-linear (Long sample)	0	0	0	0	0	0	3	1	0
No-Google (Short sample)	15	0	2	12	4	2	3	0	0
Google (Short sample)	11	0	0	11	4	2	15	11	3
All	144	16	47	64	17	12	142	63	10
Models in MCS among those with GI - G3									
No-Google (Long sample)	117	16	41	20	0	5	118	39	8
Non-linear (Long sample)	0	0	0	0	0	0	3	0	0
No-Google (Short sample)	14	0	0	5	2	1	3	2	2
Google (Short sample)	15	1	4	10	17	7	11	20	2
All	146	17	45	35	19	13	135	61	12
Models in MCS among those with GI - G5 (Principal Component)									
No-Google (Long sample)	2	9	34	35	8	6	120	29	2
Non-linear (Long sample)	0	0	0	0	0	0	3	0	0
No-Google (Short sample)	0	0	0	11	4	0	3	1	0
Google (Short sample)	7	7	18	18	18	13	21	37	7
All	9	16	52	64	30	19	147	67	9

Notes: The table shows for each set of models (estimated over the long or the short sample and with or without the GI) the number of models which are selected in the final MCS at 1-, 2- and 3-step-ahead. IS is the in-sample and OOS is the out-of-sample. *G1*, *G2*, *G3*, and *G5* are the monthly averages of the weekly Google indexes for keywords ‘jobs’, ‘collect unemployment’, ‘job center’, and the first principal component of the first three Google indexes used as leading indicators.

Table 5: Results for US unemployment rate in levels (ur_t) - forecasting with AR(1) auxiliary model. Rolling scheme. Best 15 models in terms of lowest MSE with G1 across different out-of-samples.

1-step-ahead				2-step-ahead				3-step-ahead						
Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS
Best 15 models with GI = GI1 - Out-of-Sample: 2008.10-2011.6														
Panel A1: Best models overall					Panel A2: Best models overall					Panel A3: Best models overall				
471 ARMAX(2, 2) - IC _{w4,t} - G ¹ _{w4,t}	0.019	1	0.00	1	471 ARMAX(2, 2) - IC _{w4,t} - G ¹ _{w4,t}	0.040	1	0.00	1	515 ARMAX(2, 2) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.079	1	0.00	1
516 ARMAX(2, 2) - IC _{t-2} - G ¹ _{t-2}	0.022	2	0.77	1	515 ARMAX(2, 2) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.044	2	0.59	1	427 ARMAX(1, 1) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.082	2	0.45	1
515 ARMAX(2, 2) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.023	3	1.17	1	427 ARMAX(1, 1) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.044	3	0.67	1	449 ARMAX(1, 1) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.083	3	0.86	1
427 ARMAX(1, 1) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.023	4	1.18	1	449 ARMAX(1, 1) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.046	4	0.87	1	521 ARMAX(2, 2) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.083	4	0.31	1
295 ARX(1) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.023	5	0.89	1	383 ARX(2) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.047	5	0.96	1	471 ARMAX(2, 2) - IC _{w4,t} - G ¹ _{w4,t}	0.084	5	0.50	1
361 ARX(2) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.024	6	1.16	1	389 ARX(2) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.050	6	1.11	1	389 ARX(2) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.086	6	1.04	1
296 ARX(1) - IC _{t-1} - G ¹ _{t-1}	0.024	7	1.30	1	450 ARMAX(1, 1) - IC _{t-2} - G ¹ _{t-2}	0.051	7	1.58	1	383 ARX(2) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.087	7	1.52	1
493 ARMAX(2, 2) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.024	8	1.78*	1	317 ARX(1) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.051	8	1.15	1	323 ARX(1) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.092	8	1.01	1
383 ARX(2) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.024	9	1.32	1	521 ARMAX(2, 2) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.051	9	1.29	1	317 ARX(1) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.093	9	1.40	1
450 ARMAX(1, 1) - IC _{t-2} - G ¹ _{t-2}	0.024	10	1.54	1	295 ARX(1) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.052	10	1.33	1	455 ARMAX(1, 1) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.095	10	0.88	1
449 ARMAX(1, 1) - IC _{w4,t-2} - G ¹ _{w4,t-2}	0.025	11	1.74*	1	493 ARMAX(2, 2) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.052	11	1.65*	1	301 ARX(1) - IC _{w4,t-1} - G ¹ _{w4,t-1} - SA	0.095	11	1.46	1
428 ARMAX(1, 1) - IC _{t-1} - G ¹ _{t-1}	0.025	12	1.60	1	516 ARMAX(2, 2) - IC _{t-2} - G ¹ _{t-2}	0.053	12	2.03**	1	433 ARMAX(1, 1) - IC _{w4,t-1} - G ¹ _{w4,t-1} - SA	0.098	12	1.22	1
362 ARX(2) - IC _{t-1} - G ¹ _{t-1}	0.025	13	1.42	1	301 ARX(1) - IC _{w4,t-1} - G ¹ _{w4,t-1} - SA	0.053	13	1.29	1	367 ARX(2) - IC _{w4,t-1} - G ¹ _{w4,t-1} - SA	0.099	13	1.42	1
318 ARX(1) - IC _{t-2} - G ¹ _{t-2}	0.025	14	1.58	1	323 ARX(1) - IC _{w4,t-2} - G ¹ _{w4,t-2} - SA	0.053	14	1.23	1	295 ARX(1) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.102	14	1.88*	1
273 ARX(1) - IC _{w4,t} - G ¹ _{w4,t}	0.025	15	1.44	1	361 ARX(2) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.054	15	1.56	1	361 ARX(2) - IC _{w4,t-1} - G ¹ _{w4,t-1}	0.104	15	1.79*	1
Panel B1: Best models w/o Google (Long sample)					Panel B2: Best models w/o Google (Long sample)					Panel B3: Best models w/o Google (Long sample)				
122 ARMAX(2, 2) - IC _{w4,t-2}	0.029	52	2.05**	1	122 ARMAX(2, 2) - IC _{w4,t-2}	0.077	68	2.41**	1	122 ARMAX(2, 2) - IC _{w4,t-2}	0.148	47	2.26**	0
123 ARMAX(2, 2) - IC _{t-2}	0.031	74	2.15**	1	77 ARMAX(1, 1) - IC _{w4,t} - SA	0.081	83	2.09**	1	77 ARMAX(1, 1) - IC _{w4,t} - SA	0.153	57	1.93**	1
17 ARX(1) - IC _{w4,t} - SA	0.031	77	2.28**	1	107 ARMAX(2, 2) - IC _{w4,t} - SA	0.083	90	2.16**	1	107 ARMAX(2, 2) - IC _{w4,t} - SA	0.159	71	1.98**	0
Panel C1: Best models w/o Google (Short sample)					Panel C2: Best models w/o Google (Short sample)					Panel C3: Best models w/o Google (Short sample)				
143 ARX(1) - IC _{w4,t}	0.031	82	2.33**	1	233 ARMAX(2, 2) - IC _{w4,t}	0.080	79	2.89***	0	233 ARMAX(2, 2) - IC _{w4,t}	0.165	81	3.26***	0
203 ARMAX(1, 1) - IC _{w4,t}	0.032	99	2.31**	1	203 ARMAX(1, 1) - IC _{w4,t}	0.082	87	3.03***	1	203 ARMAX(1, 1) - IC _{w4,t}	0.173	100	3.07***	0
153 ARX(1) - IC _{w4,t-1}	0.032	104	2.51**	1	143 ARX(1) - IC _{w4,t}	0.087	108	3.26***	0	208 ARMAX(1, 1) - IC _{w4,t} - SA	0.183	127	3.11***	0
Panel D1: Best nonlinear models (Long sample)					Panel D2: Best nonlinear models (Long sample)					Panel D3: Best nonlinear models (Long sample)				
129 SETAR(2)	0.045	475	2.66***	0	129 SETAR(2)	0.156	502	3.36***	0	129 SETAR(2)	0.347	495	3.47***	0
130 LSTAR(2)	0.048	494	2.71***	0	130 LSTAR(2)	0.166	505	3.31***	0	130 LSTAR(2)	0.367	499	3.35***	0
131 AAR(2)	0.049	499	2.78***	0	131 AAR(2)	0.172	506	3.31***	0	131 AAR(2)	0.379	501	3.37***	0

(Continued on next page)

Table 5 – continued

1-step-ahead			2-step-ahead			3-step-ahead		
Model	MSE	Rk DM	MCS	Model	MSE	Rk DM	MCS	Model
Best 15 models with $GI = G1$ - Out-of-Sample: 2009.7-2011.6			Best 15 models with $GI = G1$ - Out-of-Sample: 2009.7-2011.6			Best 15 models with $GI = G1$ - Out-of-Sample: 2009.7-2011.6		
Panel A1: Best models overall			Panel A2: Best models overall			Panel A3: Best models overall		
516 ARMAX(2, 2) - $IC_{t-2} - G_{t-2}^1$	0.023	1 0.00	1	471 ARMAX(2, 2) - $IC_{w4,t} - G_{w4,t}^1$	0.040	1 0.00	1	450 ARMAX(1, 1) - $IC_{t-2} - G_{t-2}^1$
270 ARX(1) - $IC_{w1,t} - G_{w1,t}^1$	0.024	2 0.19	1	450 ARMAX(1, 1) - $IC_{t-2} - G_{t-2}^1$	0.045	2 0.58	1	471 ARMAX(2, 2) - $IC_{w4,t} - G_{w4,t}^1$
473 ARMAX(2, 2) - $IC_{w1,t} \dots IC_{w4,t} - G_{w1,t}^1 \dots G_{w4,t}^1$	0.025	3 0.36	1	516 ARMAX(2, 2) - $IC_{t-2} - G_{t-2}^1$	0.046	3 0.73	1	516 ARMAX(2, 2) - $IC_{t-2} - G_{t-2}^1$
512 ARMAX(2, 2) - $IC_{w1,t-2} - G_{w1,t-2}^1$	0.025	4 0.24	1	429 ARMAX(1, 1) - $IC_{w1,t-1} \dots IC_{w4,t-1} - G_{w1,t-1}^1 \dots G_{w4,t-1}^1$	0.048	4 1.30	1	407 ARMAX(1, 1) - $IC_{w1,t} \dots IC_{w4,t} - G_{w1,t}^1 \dots G_{w4,t}^1$
458 ARMAX(2, 2) - $G_{w1,t}^1$	0.025	5 0.20	1	363 ARX(2) - $IC_{w1,t-1} \dots G_{w4,t-1}^1$	0.048	5 1.28	1	473 ARMAX(2, 2) - $IC_{w1,t} \dots IC_{w4,t} - G_{w1,t}^1 \dots G_{w4,t}^1$
275 ARX(1) - $IC_{w1,t} \dots IC_{w4,t} - G_{w1,t}^1 \dots G_{w4,t}^1$	0.025	6 0.30	1	473 ARMAX(2, 2) - $IC_{w1,t-1} \dots IC_{w4,t-1} - G_{w1,t-1}^1 \dots G_{w4,t-1}^1$	0.049	6 1.17	1	427 ARMAX(1, 1) - $IC_{w4,t-1} - G_{w4,t-1}^1$
276 ARX(1) - $IC_{w1,t} - G_{w1,t}^1 - SA$	0.025	7 0.27	1	427 ARMAX(1, 1) - $IC_{w4,t-1} - G_{w4,t-1}^1$	0.049	7 1.20	1	428 ARMAX(1, 1) - $IC_{t-1} - G_{t-1}^1$
336 ARX(2) - $IC_{w1,t} - G_{w1,t}^1$	0.025	8 0.41	1	512 ARMAX(2, 2) - $IC_{w1,t-2} - G_{w1,t-2}^1$	0.050	8 0.84	1	472 ARMAX(2, 2) - $IC_{t-1} - G_{t-1}^1$
461 ARMAX(2, 2) - $G_{w4,t}^1$	0.026	9 0.28	1	449 ARMAX(1, 1) - $IC_{w4,t-2} - G_{w4,t-2}^1$	0.052	9 1.32	1	449 ARMAX(1, 1) - $IC_{w4,t-2} - G_{w4,t-2}^1$
471 ARMAX(2, 2) - $IC_{w4,t} - G_{w4,t}^1$	0.026	10 0.62	1	369 ARX(2) - $IC_{w1,t-1} \dots IC_{w4,t-1} - G_{w1,t-1}^1 \dots G_{w4,t-1}^1$	0.052	10 1.36	1	429 ARMAX(1, 1) - $IC_{w4,t-1} - G_{w4,t-1}^1$
480 ARMAX(2, 2) - $G_{w1,t-1}^1$	0.026	11 0.31	1	303 ARX(1) - $IC_{w1,t-1} \dots IC_{w4,t-1} - G_{w1,t-1}^1 \dots G_{w4,t-1}^1$	0.052	11 1.38	1	406 ARMAX(1, 1) - $IC_{t-1} - G_{t-1}^1$
274 ARX(1) - $IC_{t-1} - G_{t-1}^1$	0.026	12 0.57	1	361 ARX(2) - $IC_{w4,t-1} - G_{w4,t-1}^1$	0.052	12 1.13	1	405 ARMAX(1, 1) - $IC_{w4,t} - G_{w4,t}^1$
483 ARMAX(2, 2) - $G_{w4,t-1}^1$	0.026	13 0.40	1	407 ARMAX(1, 1) - $IC_{w1,t} \dots IC_{w4,t} - G_{w1,t}^1 \dots G_{w4,t}^1$	0.052	13 1.69*	1	411 ARMAX(1, 1) - $IC_{w4,t} - G_{w4,t}^1$
33 ARX(1) - $IC_{t-2} - SA$	0.026	14 0.54	1	295 ARX(1) - $IC_{w4,t-1} - G_{w4,t-1}^1$	0.052	14 1.07	1	383 ARX(2) - $IC_{w4,t-2} - G_{w4,t-2}^1$
38 ARX(1) - $IC_{t-2} - SA$	0.026	15 0.54	1	383 ARX(2) - $IC_{w4,t-2} - G_{w4,t-2}^1$	0.053	15 1.35	1	515 ARMAX(2, 2) - $IC_{w4,t-2} - G_{w4,t-2}^1$
Panel B1: Best models w/o Google (Long sample)			Panel B2: Best models w/o Google (Long sample)			Panel B3: Best models w/o Google (Long sample)		
33 ARX(1) - IC_{t-2}	0.026	14 0.54	1	87 ARMAX(1, 1) - $IC_{w4,t-1} - SA$	0.061	49 1.56	1	47 ARX(2) - $IC_{w4,t} - SA$
38 ARX(1) - $IC_{t-2} - SA$	0.026	15 0.54	1	97 ARMAX(1, 1) - $IC_{w4,t-2} - SA$	0.061	50 1.56	1	37 ARX(1) - $IC_{w4,t-2} - SA$
93 ARMAX(1, 1) - IC_{t-2}	0.026	16 0.54	1	32 ARX(1) - $IC_{w4,t-2}$	0.061	52 1.59	1	102 ARMAX(2, 2) - $IC_{w4,t} - G_{w4,t}^1$
Panel C1: Best models w/o Google (Short sample)			Panel C2: Best models w/o Google (Short sample)			Panel C3: Best models w/o Google (Short sample)		
232 ARMAX(2, 2) - $IC_{w3,t-2}$	0.032	234 1.36	0	249 ARMAX(2, 2) - $IC_{t-1} - SA$	0.069	165 1.82*	0	203 ARMAX(1, 1) - $IC_{w4,t}$
142 ARX(1) - $IC_{w3,t}$	0.032	240 1.55	0	259 ARMAX(2, 2) - $IC_{t-2} - SA$	0.070	189 1.74*	0	178 ARX(2) - $IC_{w4,t} - SA$
235 ARMAX(2, 2) - $IC_{w1,t} - SA$	0.032	247 1.14	0	143 ARX(1) - $IC_{w4,t}$	0.072	218 2.09**	0	233 ARMAX(2, 2) - $IC_{w4,t}$
Panel D1: Best nonlinear models (Long sample)			Panel D2: Best nonlinear models (Long sample)			Panel D3: Best nonlinear models (Long sample)		
129 SETAR(2)	0.030	204 0.75	1	129 SETAR(2)	0.074	245 1.41	0	129 SETAR(2)
130 LSTAR(2)	0.030	197 0.74	1	130 LSTAR(2)	0.073	244 1.41	0	130 LSTAR(2)
131 AAR(2)	0.030	203 0.72	1	131 AAR(2)	0.072	208 1.25	0	131 AAR(2)

Notes: Long sample: 1967.1-2011.6; short sample: 2004.1-2011.6; out of sample: 2007.3-2011.6. The first column reports model number, **Model** is the model type, **MSE** is the mean squared error, **Rk** is the ranking with respect to the lowest MSE, **DM** is the Diebold and Mariano test for the null hypothesis of equal predictive accuracy (Diebold and Mariano, 1995), and **MCS** is the Model Confidence Set approach by Hansen, Lunde and Nason (2011). G_{t-1}^1 is the Google Index for keyword 'jobs'. The column MCS has a 1 when the row model is included in the final model confidence set at 5% confidence level and a 0 otherwise. In all panels ***, ** and * indicate rejection at 1, 5 and 10%, respectively.

Table 6: Forecasting the unemployment rate by state: h -step-ahead state level forecasts with AR(1) auxiliary model. Out of sample 2007.2-2011.6.

State	Panel A - 1-step ahead				Panel B - 2-step ahead				Panel C - 3-step ahead			
	No Google		Google: G1		No Google		Google: G1		No Google		Google: G1	
	mod #	MSE	mod #	MSE	mod #	MSE	mod #	MSE	mod #	MSE	mod #	MSE
1	122	2.88E-03	459	2.92E-03	120	2.75E-02	469	2.63E-02	123	1.01E-01	469	1.06E-01
2	254	1.44E-03	508	1.28E-03	254	1.17E-02	503	8.11E-03	106	2.37E-02	503	1.99E-02
3	123	5.84E-03	489	5.67E-03	123	1.85E-02	489	1.89E-02	112	4.91E-02	481	5.07E-02
4	112	1.01E-03	464	9.21E-04	112	8.94E-03	481	8.48E-03	120	2.08E-02	459	2.98E-02
5	128	2.00E-03	503	1.37E-03	236	1.16E-02	486	1.12E-02	110	2.95E-02	470	3.37E-02
6	242	4.47E-03	483	3.83E-03	107	3.34E-02	503	2.92E-02	108	9.43E-02	507	9.35E-02
7	112	3.92E-04	455	4.73E-04	112	4.36E-03	503	4.70E-03	112	1.80E-02	513	1.84E-02
8	248	7.43E-03	332	7.45E-03	248	3.30E-02	332	3.22E-02	248	1.04E-01	332	8.77E-02
9	194	2.52E-03	373	2.45E-03	174	1.24E-02	343	1.11E-02	234	3.28E-02	473	3.06E-02
10	235	1.72E-03	512	1.58E-03	233	1.90E-02	512	1.62E-02	247	6.19E-02	507	4.37E-02
11	176	6.33E-03	343	5.72E-03	176	3.32E-02	458	2.70E-02	242	1.01E-01	502	7.13E-02
12	242	1.94E-03	338	1.71E-03	242	1.06E-02	498	8.59E-03	242	2.50E-02	498	2.44E-02
13	133	2.19E-03	499	2.21E-03	114	4.19E-03	467	4.57E-03	8	8.42E-03	480	1.09E-02
14	240	1.84E-03	503	1.12E-03	228	2.02E-02	503	1.01E-02	101	6.02E-02	503	3.58E-02
15	171	1.46E-02	349	1.35E-02	170	8.69E-02	380	7.75E-02	170	2.89E-01	380	2.38E-01
16	117	3.51E-03	266	3.93E-03	127	9.57E-03	277	1.02E-02	127	2.08E-02	277	1.85E-02
17	116	9.19E-04	503	8.83E-04	122	9.65E-03	512	9.04E-03	122	3.62E-02	512	3.72E-02
18	100	3.62E-03	503	3.57E-03	7	2.74E-02	498	2.68E-02	112	9.83E-02	503	9.55E-02
19	39	1.52E-02	516	1.27E-02	110	6.52E-02	360	6.12E-02	110	1.68E-01	448	1.61E-01
20	243	2.58E-03	508	1.71E-03	99	1.41E-02	503	1.26E-02	99	3.85E-02	503	6.28E-02
21	244	1.21E-03	503	9.27E-04	241	1.05E-02	503	7.73E-03	241	3.62E-02	503	2.96E-02
22	251	5.69E-04	491	5.32E-04	7	7.04E-03	478	6.50E-03	7	1.87E-02	500	1.91E-02
23	229	4.28E-03	503	2.89E-03	229	3.32E-02	503	2.11E-02	229	1.29E-01	508	8.57E-02
24	248	2.70E-03	497	1.50E-03	112	2.22E-02	513	1.17E-02	120	6.92E-02	513	4.43E-02
25	39	1.42E-02	518	1.50E-02	39	7.57E-02	497	4.96E-02	39	2.04E-01	387	1.06E-01
26	250	1.20E-03	501	1.11E-03	250	1.29E-02	502	1.30E-02	102	4.94E-02	457	5.11E-02
27	132	1.37E-03	352	1.29E-03	175	4.43E-03	352	4.50E-03	175	9.98E-03	330	1.03E-02
28	127	4.20E-04	497	4.84E-04	127	3.52E-03	508	4.11E-03	128	1.32E-02	508	1.54E-02
29	220	1.75E-02	496	1.67E-02	255	6.49E-02	496	5.82E-02	231	1.11E-01	516	1.08E-01
30	251	1.07E-03	513	9.29E-04	102	1.31E-02	519	1.16E-02	102	3.56E-02	508	3.72E-02
31	119	8.53E-04	459	8.75E-04	119	7.56E-03	503	8.05E-03	119	2.90E-02	503	3.50E-02
32	4	1.02E-02	377	9.03E-03	248	3.49E-02	488	3.30E-02	67	9.11E-02	458	9.13E-02
33	229	2.57E-03	459	2.01E-03	123	1.58E-02	327	1.22E-02	120	4.15E-02	327	3.68E-02
34	4	9.18E-03	464	7.38E-03	65	5.43E-02	464	4.48E-02	125	1.76E-01	464	1.67E-01
35	248	1.57E-03	456	1.60E-03	248	1.10E-02	456	1.19E-02	103	2.41E-02	470	2.92E-02
36	251	2.31E-03	513	1.75E-03	127	2.40E-02	513	1.70E-02	99	7.04E-02	503	6.46E-02
37	48	9.76E-03	325	1.02E-02	128	3.82E-02	457	3.93E-02	128	9.38E-02	458	1.06E-01
38	101	3.60E-03	503	3.09E-03	102	3.24E-02	503	2.93E-02	122	1.15E-01	470	1.32E-01
39	127	4.80E-03	349	4.12E-03	127	1.91E-02	503	1.76E-02	7	4.61E-02	502	4.00E-02
40	192	1.52E-03	354	1.56E-03	55	7.33E-03	486	6.50E-03	55	1.96E-02	508	2.26E-02
41	233	7.49E-03	486	4.76E-03	231	5.55E-02	503	3.22E-02	231	1.58E-01	503	1.04E-01
42	109	6.54E-04	480	6.53E-04	103	8.13E-03	491	8.47E-03	103	3.32E-02	491	3.71E-02
43	230	4.21E-03	503	3.56E-03	230	2.94E-02	503	2.05E-02	119	1.03E-01	503	6.76E-02
44	57	2.35E-03	338	2.31E-03	57	1.05E-02	338	1.09E-02	52	2.67E-02	469	3.17E-02
45	235	2.44E-03	473	2.51E-03	127	1.93E-02	344	1.83E-02	127	4.88E-02	470	5.65E-02
46	119	1.01E-03	503	9.24E-04	119	1.33E-02	503	1.35E-02	119	5.39E-02	503	7.12E-02
47	234	1.15E-03	503	8.74E-04	120	1.06E-02	503	7.79E-03	120	3.38E-02	503	3.45E-02
48	100	2.58E-03	503	1.80E-03	100	2.17E-02	503	1.70E-02	100	5.84E-02	470	7.45E-02
49	231	3.58E-03	460	2.75E-03	101	2.13E-02	503	2.38E-02	101	5.80E-02	502	7.46E-02
50	250	2.41E-03	503	1.90E-03	127	2.73E-02	503	2.34E-02	126	1.00E-01	503	1.21E-01
51	236	5.43E-04	459	5.74E-04	123	5.75E-03	459	6.80E-03	123	2.38E-02	513	3.27E-02
Percentage of best models with GI												
among first 5				70.2%				62.0%				39.2%
among first 10				65.7%				56.3%				35.7%
among first 15				60.9%				52.7%				33.1%

Notes: G1 is the GI for ‘jobs’, the only one available at the state level. In-sample ending with 2007.1; out of sample: 2007.2-2011.6. **State** reports the State code (we consider also District Columbia) **mod #** is model number, **MSE** reports the lowest mean squared error. In each row, the MSE in bold indicates the best model.

Table 7: Results for US unemployment rate in levels (ur_t) - forecasting with AR(1) auxiliary model. Rolling scheme. Best 15 models in terms of lowest MSE across different Google keywords.

1-step-ahead				2-step-ahead				3-step-ahead						
Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS
Best 15 models with G1 - G2 - Out-of-Sample: 2007.2-2011.6														
Panel A1: Best models overall					Panel A2: Best models overall					Panel A3: Best models overall				
128 ARMAX(2,2) - IC _{t-2} - SA	0.028	1	0.00	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.066	1	0.00	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.133	1	0.00	1
122 ARMAX(2,2) - IC _{w4,t-2}	0.028	2	0.26	1	17 ARX(1) - IC _{w4,t-2} - SA	0.068	2	0.53	1	47 ARX(2) - IC _{w4,t-2} - SA	0.135	2	0.24	1
123 ARMAX(2,2) - IC _{t-2}	0.029	3	0.59	1	47 ARX(2) - IC _{w4,t-2} - SA	0.069	3	0.62	1	17 ARX(1) - IC _{w4,t-2} - SA	0.136	3	0.30	1
17 ARX(1) - IC _{w4,t-2} - SA	0.029	4	0.61	1	12 ARX(1) - IC _{w4,t-2} - SA	0.069	4	0.64	1	72 ARMAX(1,1) - IC _{w4,t-2} - SA	0.138	4	0.47	1
77 ARMAX(1,1) - IC _{w4,t-2} - SA	0.029	5	0.60	1	77 ARMAX(1,1) - IC _{w4,t-2} - SA	0.069	5	0.75	1	12 ARX(1) - IC _{w4,t-2} - SA	0.138	5	0.52	1
12 ARX(1) - IC _{w4,t-2} - SA	0.029	6	0.83	1	72 ARMAX(1,1) - IC _{w4,t-2} - SA	0.069	6	0.67	1	42 ARX(2) - IC _{w4,t-2} - SA	0.138	6	0.52	1
127 ARMAX(2,2) - IC _{w4,t-2} - SA	0.029	7	0.87	1	42 ARX(2) - IC _{w4,t-2} - SA	0.069	7	0.73	1	77 ARMAX(1,1) - IC _{w4,t-2} - SA	0.140	7	0.94	1
72 ARMAX(2,2) - IC _{w4,t-2} - SA	0.029	8	1.04	1	107 ARMAX(2,2) - IC _{w4,t-2} - SA	0.071	8	1.17	1	102 ARMAX(2,2) - IC _{w4,t-2} - SA	0.143	8	0.87	1
77 ARMAX(1,1) - IC _{w4,t-2} - SA	0.029	9	0.89	1	123 ARMAX(2,2) - IC _{t-2}	0.071	9	1.63	1	123 ARMAX(2,2) - IC _{t-2}	0.143	9	1.44	1
47 ARX(2) - IC _{w4,t-2} - SA	0.030	10	1.04	1	102 ARMAX(2,2) - IC _{t-2} - SA	0.071	10	1.01	1	107 ARMAX(2,2) - IC _{w4,t-2} - SA	0.145	10	1.41	1
87 ARMAX(1,1) - IC _{w4,t-2} - SA	0.030	11	1.18	1	128 ARMAX(2,2) - IC _{t-2} - SA	0.072	11	1.20	1	52 ARX(2) - IC _{w4,t-2} - SA	0.146	11	1.10	1
42 ARX(2) - IC _{w4,t-2} - SA	0.030	12	1.18	1	52 ARX(2) - IC _{w4,t-2} - SA	0.072	12	1.16	1	82 ARMAX(1,1) - IC _{w4,t-2} - SA	0.146	12	1.12	1
78 ARMAX(1,1) - IC _{t-2} - SA	0.030	13	1.23	1	82 ARMAX(1,1) - IC _{w4,t-2} - SA	0.072	13	1.16	1	57 ARX(2) - IC _{w4,t-2} - SA	0.146	13	1.14	1
120 ARMAX(2,2) - IC _{w2,t-2}	0.030	14	1.25	1	87 ARMAX(1,1) - IC _{w4,t-2} - SA	0.073	14	1.42	0	22 ARX(1) - IC _{w4,t-2} - SA	0.147	14	1.15	1
18 ARX(1) - IC _{t-2} - SA	0.030	14	1.25	1	127 ARMAX(2,2) - IC _{w4,t-2} - SA	0.073	15	1.31	1	27 ARX(1) - IC _{w4,t-2} - SA	0.147	15	1.16	1
119 ARMAX(2,2) - IC _{w1,t-2}	0.030	15	1.23	1	Panel B2: Best models w Google (Short sample)					Panel B3: Best models w Google (Short sample)				
Panel B1: Best models w Google (Short sample)					276 ARX(1) - IC _{w1,t-2} - SA	0.0330	86	0.927	1	433 ARMAX(1,1) - IC _{w4,t-1} - G ² _{w4,t-1} - SA	0.258	167	2.37**	0
276 ARX(1) - IC _{w1,t-2} - SA	0.0330	86	0.927	1	434 ARMAX(1,1) - IC _{t-1} - G ² _{t-1} - SA	0.1147	135	0	0	G ² _{w4,t-1} - SA	0.253	161	1.84*	0
274 ARX(1) - IC _{t-2} - G ² _{t-2}	0.0336	113	1.020	1	SA	0.101	151	1.78*	0	SA	0.257	166	2.14**	0
270 ARX(1) - IC _{w1,t-2} - G ² _{w1,t-2}	0.0340	122	1.082	1	SA	0.101	151	1.78*	0	SA	0.257	166	2.14**	0
Panel C1: Best models w/o Google (Short sample)					Panel C2: Best models w/o Google (Short sample)					Panel C3: Best models w/o Google (Short sample)				
148 ARX(1) - IC _{w4,t-2} - SA	0.032	61	1.11	1	258 ARMAX(2,2) - IC _{w4,t-2} - SA	0.083	51	1.22	0	258 ARMAX(2,2) - IC _{w4,t-2} - SA	0.172	64	1.16	1
143 ARX(1) - IC _{w4,t-2}	0.033	68	1.17	1	178 ARX(2) - IC _{w4,t-2} - SA	0.088	78	1.65*	0	218 ARMAX(1,1) - IC _{w4,t-2} - SA	0.190	111	1.60	1
144 ARX(1) - IC _{t-2}	0.033	104	1.33	1	153 ARX(1) - IC _{w4,t-2} - SA	0.088	79	1.73*	0	248 ARMAX(2,2) - IC _{w4,t-2} - SA	0.200	119	1.84*	0
Panel D1: Best models w/o Google (Long sample)					Panel D2: Best models w/o Google (Long sample)					Panel D3: Best models w/o Google (Long sample)				
128 ARMAX(2,2) - IC _{t-2} - SA	0.028	1	0.00	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.066	1	0.00	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.133	1	0.00	1
122 ARMAX(2,2) - IC _{w4,t-2}	0.028	2	0.26	1	17 ARX(1) - IC _{w4,t-2} - SA	0.068	2	0.53	1	47 ARX(2) - IC _{w4,t-2} - SA	0.135	2	0.24	1
123 ARMAX(2,2) - IC _{t-2}	0.029	3	0.59	1	47 ARX(2) - IC _{w4,t-2} - SA	0.069	3	0.62	1	17 ARX(1) - IC _{w4,t-2} - SA	0.136	3	0.30	1
Panel E1: Best nonlinear models (Long sample)					Panel E2: Best nonlinear models (Long sample)					Panel E3: Best nonlinear models (Long sample)				
129 SETAR(2)	0.040	204	2.68***	0	129 SETAR(2)	0.127	255	3.11***	0	129 SETAR(2)	0.282	217	3.54***	0
130 LSTAR(2)	0.042	229	2.68***	0	130 LSTAR(2)	0.130	267	2.97***	0	130 LSTAR(2)	0.288	225	3.33***	0
131 AAR(2)	0.043	245	2.74***	0	131 AAR(2)	0.132	275	2.88***	0	131 AAR(2)	0.293	234	3.22***	0
(Continued on next page)														

Table 7 – continued

1-step-ahead				2-step-ahead				3-step-ahead						
Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS
Panel A1: Best models overall					Panel A2: Best models overall					Panel A3: Best models overall				
128 ARMAX(2, 2) - IC _{t-2} - SA	0.028	1	0.00	1	122 ARMAX(2, 2) - IC _{w4,t-2}	0.066	1	0.00	1	122 ARMAX(2, 2) - IC _{w4,t-2}	0.133	1	0.00	1
122 ARMAX(2, 2) - IC _{w4,t-2}	0.028	2	0.25	1	17 ARX(1) - IC _{w4,t-2} - SA	0.068	2	0.53	1	47 ARX(2) - IC _{w4,t-2} - SA	0.135	2	0.24	1
295 ARX(1) - IC _{w4,t-1} - G ³ _{w4,t-1}	0.028	3	0.20	1	77 ARMAX(1, 1) - IC _{w4,t-2} - SA	0.069	3	0.69	1	17 ARX(1) - IC _{w4,t-2} - SA	0.136	3	0.30	1
123 ARMAX(2, 2) - IC _{t-2}	0.029	4	0.59	1	47 ARX(2) - IC _{w4,t-2} - SA	0.069	4	0.62	1	72 ARMAX(1, 1) - IC _{w4,t-2} - SA	0.138	4	0.47	1
17 ARX(1) - IC _{w4,t-2} - SA	0.029	5	0.60	1	12 ARX(1) - IC _{w4,t-2}	0.069	5	0.64	1	12 ARX(1) - IC _{w4,t-2}	0.138	5	0.52	1
77 ARMAX(1, 1) - IC _{w4,t-2} - SA	0.029	6	0.55	1	72 ARMAX(1, 1) - IC _{w4,t-2}	0.069	6	0.67	1	42 ARX(2) - IC _{w4,t-2}	0.138	6	0.52	1
12 ARX(1) - IC _{w4,t-2}	0.029	7	0.82	1	42 ARX(2) - IC _{w4,t-2}	0.069	7	0.73	1	77 ARMAX(1, 1) - IC _{w4,t-2} - SA	0.140	7	0.94	1
127 ARMAX(2, 2) - IC _{w4,t-2} - SA	0.029	8	0.86	1	107 ARMAX(2, 2) - IC _{w4,t-2} - SA	0.071	8	1.17	1	102 ARMAX(2, 2) - IC _{w4,t-2}	0.143	8	0.87	1
301 ARX(1) - IC _{w4,t-1} - G ³ _{w4,t-1} - SA	0.029	9	0.35	1	123 ARMAX(2, 2) - IC _{t-2}	0.071	9	1.63	1	123 ARMAX(2, 2) - IC _{t-2}	0.143	9	1.44	1
72 ARMAX(1, 1) - IC _{w4,t-2}	0.029	10	1.04	1	102 ARMAX(2, 2) - IC _{w4,t-2}	0.071	10	1.01	1	107 ARMAX(2, 2) - IC _{w4,t-2} - SA	0.145	10	1.41	1
47 ARX(2) - IC _{w4,t-2} - SA	0.029	11	0.89	1	127 ARMAX(2, 2) - IC _{w4,t-2} - SA	0.072	11	1.24	1	52 ARX(2) - IC _{w4,t-2} - SA	0.146	11	1.10	1
87 ARMAX(1, 1) - IC _{w4,t-1} - SA	0.030	12	1.04	1	52 ARX(2) - IC _{w4,t-1}	0.072	12	1.16	1	82 ARMAX(1, 1) - IC _{w4,t-1}	0.146	12	1.12	1
42 ARX(2) - IC _{w4,t-2}	0.030	13	1.17	1	82 ARMAX(1, 1) - IC _{w4,t-1}	0.072	13	1.16	1	57 ARX(2) - IC _{w4,t-1} - SA	0.146	13	1.14	1
78 ARMAX(1, 1) - IC _{t-2} - SA	0.030	14	1.17	1	87 ARMAX(1, 1) - IC _{w4,t-1} - SA	0.072	14	1.42	0	22 ARX(1) - IC _{w4,t-1}	0.147	14	1.15	1
120 ARMAX(2, 2) - IC _{w2,t-2}	0.030	15	1.22	1	128 ARMAX(2, 2) - IC _{t-2} - SA	0.072	15	1.33	0	27 ARX(1) - IC _{w4,t-1} - SA	0.147	15	1.16	1
Panel B1: Best models w Google (Short sample)					Panel B2: Best models w Google (Short sample)					Panel B3: Best models w Google (Short sample)				
295 ARX(1) - IC _{w4,t-1} - G ² _{w4,t-1}	0.0285	3	0.195	1	515 ARMAX(2, 2) - IC _{w4,t-2} - G ² _{w4,t-2}	0.073	19	0.49	1	471 ARMAX(2, 2) - IC _{w4,t-2} - G ² _{w4,t-2}	0.172	63	1.05	0
301 ARX(1) - IC _{w4,t-1} - G ² _{w4,t-1} - SA	0.0292	9	0.353	0	429 ARMAX(1, 1) - IC _{w4,t-1} - G ² _{w4,t-1}	0.077	34	0.69	0	317 ARX(1) - IC _{w4,t-2} - G ² _{w4,t-2}	0.204	145	1.11	0
471 ARMAX(2, 2) - IC _{w4,t-2} - G ² _{w4,t-2}	0.0299	17	0.479	0	471 ARMAX(2, 2) - IC _{w4,t-2} - G ² _{w4,t-2}	0.077	35	0.94	0	433 ARMAX(1, 1) - IC _{w4,t-1} - G ² _{w4,t-1} - SA	0.168	50	1.28	0
Panel C1: Best models w/o Google (Short sample)					Panel C2: Best models w/o Google (Short sample)					Panel C3: Best models w/o Google (Short sample)				
148 ARX(1) - IC _{w4,t-2} - SA	0.032	75	1.11	1	258 ARMAX(2, 2) - IC _{w4,t-2} - SA	0.083	59	1.19	0	258 ARMAX(2, 2) - IC _{w4,t-2} - SA	0.172	65	1.15	1
143 ARX(1) - IC _{w4,t-2}	0.033	81	1.16	1	178 ARX(2) - IC _{w4,t-2} - SA	0.088	97	1.65*	0	218 ARMAX(1, 1) - IC _{w4,t-1} - SA	0.190	120	1.59	1
144 ARX(1) - IC _{t-2}	0.033	119	1.32	1	153 ARX(1) - IC _{w4,t-1}	0.088	98	1.73*	0	248 ARMAX(2, 2) - IC _{w4,t-1} - SA	0.200	134	1.79*	0
Panel D1: Best models w/o Google (Long sample)					Panel D2: Best models w/o Google (Long sample)					Panel D3: Best models w/o Google (Long sample)				
128 ARMAX(2, 2) - IC _{t-2} - SA	0.028	1	0.00	1	122 ARMAX(2, 2) - IC _{w4,t-2}	0.066	1	0.00	1	122 ARMAX(2, 2) - IC _{w4,t-2}	0.133	1	0.00	1
122 ARMAX(2, 2) - IC _{w4,t-2}	0.028	2	0.25	1	17 ARX(1) - IC _{w4,t-2} - SA	0.068	2	0.53	1	47 ARX(2) - IC _{w4,t-2} - SA	0.135	2	0.24	1
123 ARMAX(2, 2) - IC _{t-2}	0.029	4	0.59	1	77 ARMAX(1, 1) - IC _{w4,t-2} - SA	0.069	3	0.69	1	17 ARX(1) - IC _{w4,t-2} - SA	0.136	3	0.30	1
Panel E1: Best nonlinear models (Long sample)					Panel E2: Best nonlinear models (Long sample)					Panel E3: Best nonlinear models (Long sample)				
129 SETAR(2)	0.040	246	2.67***	0	129 SETAR(2)	0.127	302	3.11***	0	129 SETAR(2)	0.282	279	3.54***	0
130 LSTAR(2)	0.042	270	2.68***	0	130 LSTAR(2)	0.130	315	2.97***	0	130 LSTAR(2)	0.288	288	3.33***	0
131 AAR(2)	0.043	285	2.73***	0	131 AAR(2)	0.132	320	2.88***	0	131 AAR(2)	0.293	296	3.22***	0
(Continued on next page)														

Table 7 – continued

1-step-ahead					2-step-ahead					3-step-ahead				
Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS
Best 15 models with G1 - G5 - Out-of-Sample: 2007-2-2011.6														
Panel A1: Best models overall					Panel A2: Best models overall					Panel A3: Best models overall				
295 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.023	1	0.00	1	317 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.049	1	0.00	1	323 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.103	1	0.00	1
317 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.024	2	0.07	1	323 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.050	2	0.15	1	317 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.109	2	1.01	1
301 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.024	3	0.70	1	301 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.052	3	0.54	1	301 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.113	3	1.08	1
383 ARX(2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.025	4	0.46	1	389 ARX(2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.052	4	0.65	1	389 ARX(2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.119	4	1.49	1
515 ARMAX(2,2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.025	5	0.45	1	383 ARX(2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.052	5	0.75	1	295 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.119	5	1.62	1
323 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.025	6	0.73	1	295 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.053	6	0.68	1	515 ARMAX(2,2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.120	6	1.25	1
361 ARX(2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.025	7	0.70	1	515 ARMAX(2,2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.054	7	1.01	1	383 ARX(2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.124	7	1.53	1
389 ARX(2) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.026	8	0.65	1	367 ARX(2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.057	8	1.08	1	367 ARX(2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.131	8	1.74*	1
367 ARX(2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.026	9	0.78	1	361 ARX(2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.058	9	1.15	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.133	9	1.40	1
302 ARX(1) - IC _{t-1} - G ^{PC} _{t-1} - SA	0.026	10	0.75	1	427 ARMAX(1,1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.063	10	1.96**	1	361 ARX(2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.135	10	1.85*	1
275 ARX(1) - IC _{w1,t} ... IC _{w4,t} - G ^{PC} _{w1,t} ... G ^{PC} _{w4,t}	0.026	11	0.81	1	471 ARMAX(2,2) - IC _{w4,t} - G ^{PC} _{w4,t}	0.064	11	1.81*	1	47 ARX(2) - IC _{w4,t} - SA	0.135	11	1.33	1
499 ARMAX(2,2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1} - SA	0.026	12	0.70	1	302 ARX(1) - IC _{t-1} - G ^{PC} _{t-1} - SA	0.065	12	1.51	1	17 ARX(1) - IC _{w4,t} - SA	0.136	12	1.35	1
273 ARX(1) - IC _{w4,t} - G ^{PC} _{w4,t}	0.027	13	1.37	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.066	13	1.69*	1	72 ARMAX(1,1) - IC _{w4,t}	0.138	13	1.41	1
279 ARX(1) - IC _{w4,t} - G ^{PC} _{w4,t} - SA	0.027	14	1.49	1	493 ARMAX(2,2) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.067	14	1.74*	1	12 ARX(1) - IC _{w4,t}	0.138	14	1.43	1
296 ARX(1) - IC _{t-1} - G ^{PC} _{t-1}	0.027	15	0.99	1	433 ARMAX(1,1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1} - SA	0.067	15	2.08**	1	42 ARX(2) - IC _{w4,t}	0.138	15	1.43	1
Panel B1: Best models w Google (Short sample)					Panel B2: Best models w Google (Short sample)					Panel B3: Best models w Google (Short sample)				
295 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.023	1	0.00	1	317 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.049	1	0.00	1	323 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.103	1	0.00	1
317 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.024	2	0.07	1	323 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.050	2	0.15	1	317 ARX(1) - IC _{w4,t-2} - G ^{PC} _{w4,t-2}	0.109	2	1.01	1
301 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.024	3	0.70	1	301 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.052	3	0.54	1	301 ARX(1) - IC _{w4,t-1} - G ^{PC} _{w4,t-1}	0.113	3	1.08	1
Panel C1: Best models w/o Google (Short sample)					Panel C2: Best models w/o Google (Short sample)					Panel C3: Best models w/o Google (Short sample)				
148 ARX(1) - IC _{w4,t} - SA	0.032	104	2.35**	0	258 ARMAX(2,2) - IC _{w4,t-2} - SA	0.083	102	2.75***	0	258 ARMAX(2,2) - IC _{w4,t-2} - SA	0.172	95	2.60***	0
143 ARX(1) - IC _{w4,t}	0.033	106	2.41**	0	178 ARX(2) - IC _{w4,t} - SA	0.088	138	3.42***	0	218 ARMAX(1,1) - IC _{w4,t-1} - SA	0.190	166	3.02***	0
144 ARX(1) - IC _t	0.033	152	2.60***	0	153 ARX(1) - IC _{w4,t-1}	0.088	139	3.64***	0	248 ARMAX(2,2) - IC _{w4,t-1} - SA	0.200	182	3.47***	0
Panel D1: Best models w/o Google (Long sample)					Panel D2: Best models w/o Google (Long sample)					Panel D3: Best models w/o Google (Long sample)				
122 ARMAX(2,2) - IC _{w4,t-2}	0.028	22	1.61	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.066	13	1.69*	1	122 ARMAX(2,2) - IC _{w4,t-2}	0.133	9	1.40	1
17 ARX(1) - IC _{w4,t} - SA	0.028	27	1.72*	0	17 ARX(1) - IC _{w4,t} - SA	0.068	21	1.67*	1	47 ARX(2) - IC _{w4,t} - SA	0.135	11	1.33	1
77 ARMAX(1,1) - IC _{w4,t} - SA	0.029	31	1.53	0	77 ARMAX(1,1) - IC _{w4,t} - SA	0.069	22	1.57	1	17 ARX(1) - IC _{w4,t} - SA	0.136	12	1.35	1
Panel E1: Best nonlinear models (Long sample)					Panel E2: Best nonlinear models (Long sample)					Panel E3: Best nonlinear models (Long sample)				
129 SETAR(2)	0.040	303	3.00***	0	129 SETAR(2)	0.127	369	3.56***	0	129 SETAR(2)	0.282	358	3.76***	0
130 LSTAR(2)	0.042	340	2.91***	0	130 LSTAR(2)	0.130	383	3.29***	0	130 LSTAR(2)	0.288	368	3.43***	0
131 AAR(2)	0.043	358	2.99***	0	131 AAR(2)	0.132	398	3.26***	0	131 AAR(2)	0.293	379	3.42***	0

Notes: Long sample: 1967.1-2011.6; short sample: 2004.1-2011.6; out of sample: 2007.3-2011.6. The first column reports model number. **Model** is the model type, **MSE** is the mean squared error, **Rk** is the ranking with respect to the lowest MSE, **DM** is the Diebold and Mariano test for the null hypothesis of equal predictive accuracy (Diebold and Mariano, 1995), and **MCS** is the Model Confidence Set approach by Hansen, Lunde and Nason (2011). G_t^2 , G_t^3 , and G_t^4 are the leading indicators (Google Index for keyword 'collect unemployment', 'job center', and the first principal component of G1, G2 and G3, respectively). The column MCS has a 1 when the row model is included in the final model confidence set at 5% confidence level and a 0 otherwise. In all panels ***, ** and * indicate rejection at 1, 5 and 10%, respectively.

Table 8: Results for US unemployment rate in levels (ur_t) - forecasting with AR(1) auxiliary model. Rolling scheme. Falsification test.

Model	1-step-ahead			2-step-ahead			3-step-ahead		
	MSE	Rk	DM	MCS	Model	MSE	Rk	DM	MCS
Panel A1: Best models overall									
128 ARMAX(2,2) - $IC_{t-2} - SA$	0.028	1	0.00	1	122 ARMAX(2,2) - $IC_{w4,t-2}$	0.066	1	0.00	1
122 ARMAX(2,2) - $IC_{w4,t-2}$	0.028	2	0.19	1	17 ARX(1) - $IC_{w4,t-2}$	0.068	2	0.53	1
123 ARMAX(2,2) - IC_{t-2}	0.029	3	0.51	1	77 ARMAX(1,1) - $IC_{w4,t-2}$	0.069	3	0.89	1
17 ARX(1) - $IC_{w4,t-2}$	0.029	4	0.53	1	47 ARX(2) - $IC_{w4,t-2}$	0.069	4	0.62	1
77 ARMAX(1,1) - $IC_{w4,t-2}$	0.029	5	0.49	1	12 ARX(1) - $IC_{w4,t-2}$	0.069	5	0.64	1
12 ARX(1) - $IC_{w4,t-2}$	0.029	6	0.75	1	72 ARMAX(1,1) - $IC_{w4,t-2}$	0.069	6	0.67	1
72 ARMAX(1,1) - $IC_{w4,t-2}$	0.029	7	0.79	1	42 ARX(2) - $IC_{w4,t-2}$	0.069	7	0.73	1
127 ARMAX(2,2) - $IC_{w4,t-2} - SA$	0.029	8	0.96	1	107 ARMAX(2,2) - $IC_{w4,t-2} - SA$	0.071	8	0.87	1
72 ARMAX(1,1) - $IC_{w4,t-2}$	0.029	9	0.82	1	123 ARMAX(2,2) - IC_{t-2}	0.071	9	1.17	1
47 ARX(2) - $IC_{w4,t-2}$	0.029	10	0.98	1	102 ARMAX(2,2) - $IC_{w4,t-2}$	0.071	10	1.01	1
87 ARMAX(1,1) - $IC_{w4,t-2}$	0.030	11	1.09	1	127 ARMAX(2,2) - $IC_{w4,t-2} - SA$	0.072	11	1.24	1
42 ARX(2) - $IC_{w4,t-2}$	0.030	12	1.11	1	52 ARX(2) - $IC_{w4,t-2}$	0.072	12	1.16	1
78 ARMAX(1,1) - $IC_t - SA$	0.030	13	1.16	1	82 ARMAX(1,1) - $IC_{w4,t-1}$	0.072	13	1.16	1
120 ARMAX(2,2) - $IC_{w2,t-2}$	0.030	14	1.18	1	87 ARMAX(1,1) - $IC_{w4,t-1} - SA$	0.072	14	1.42	0
18 ARX(1) - $IC_t - SA$	0.030	15	1.18	1	22 ARX(1) - $IC_{w4,t-1}$	0.073	15	1.19	1
119 ARMAX(2,2) - $IC_{w1,t-2}$	0.030	15	1.18	1					
Panel B1: Best models w Google (Short sample)									
471 ARMAX(2,2) - $IC_{w4,t} - G_{w4,t}^4$	0.036	134	1.90*	1	449 ARMAX(1,1) - $IC_{w4,t-2} - G_{w4,t-2}^4$	0.086	62	1.48	0
292 ARX(1) - $IC_{w1,t-1} - G_{w1,t-1}^4$	0.039	157	2.20**	0	471 ARMAX(2,2) - $IC_{w4,t} - G_{w4,t}^4$	0.093	115	1.90*	0
345 ARX(2) - $IC_{w4,t} - G_{w4,t}^4$	0.039	158	2.85***	0	345 ARX(2) - $IC_{w4,t} - G_{w4,t}^4 - SA$	0.101	151	3.97***	0
Panel C1: Best models w/o Google (Short sample)									
148 ARX(1) - $IC_{w4,t} - SA$	0.032	61	1.08	1	258 ARMAX(2,2) - $IC_{w4,t-2} - SA$	0.083	48	1.19	0
143 ARX(1) - $IC_{w4,t}$	0.033	67	1.14	1	178 ARX(2) - $IC_{w4,t-2} - SA$	0.088	79	1.65*	0
144 ARX(1) - IC_t	0.033	102	1.29	1	153 ARX(1) - $IC_{w4,t-1}$	0.088	80	1.73*	0
Panel D1: Best models w/o Google (Long sample)									
128 ARMAX(2,2) - $IC_{t-2} - SA$	0.028	1	0.00	1	122 ARMAX(2,2) - $IC_{w4,t-2}$	0.066	1	0.00	1
122 ARMAX(2,2) - $IC_{w4,t-2}$	0.028	2	0.19	1	17 ARX(1) - $IC_{w4,t-2}$	0.068	2	0.53	1
123 ARMAX(2,2) - IC_{t-2}	0.029	3	0.51	1	77 ARMAX(1,1) - $IC_{w4,t-2}$	0.069	3	0.69	1
Panel E1: Best nonlinear models (Long sample)									
129 SETAR(2)	0.040	179	2.66***	1	129 SETAR(2)	0.127	226	3.11***	0
130 LSTAR(2)	0.042	201	2.67***	1	130 LSTAR(2)	0.130	238	2.97***	0
131 AAR(2)	0.043	211	2.72***	1	131 AAR(2)	0.132	244	2.88***	0

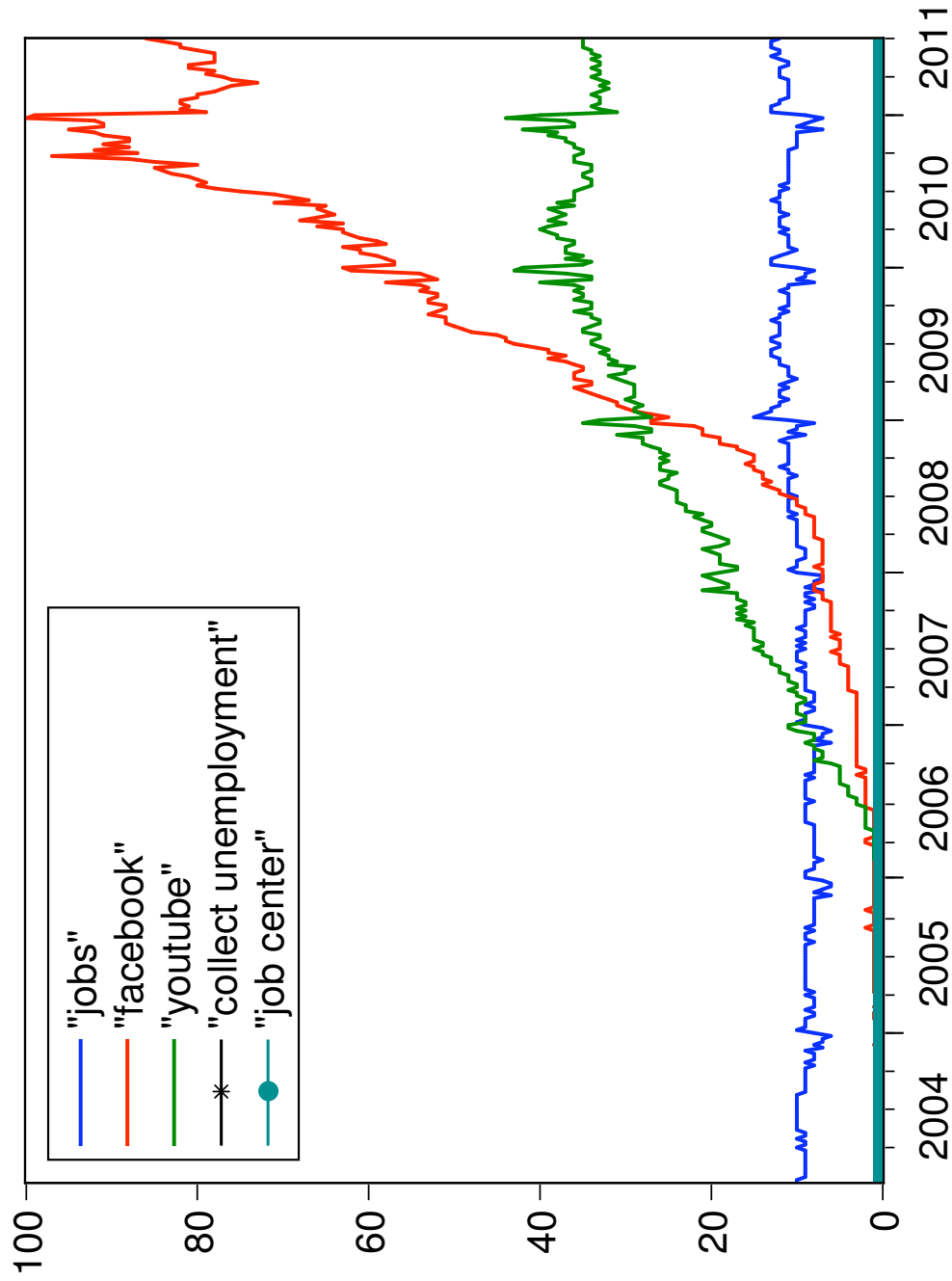
Notes: Long sample: 1967.1-2011.6; short sample 2004.1-2011.6; out of sample: 2007.3-2011.6. The first column reports model number. **Model** is the model type, **MSE** is the mean squared error, **Rk** is the ranking with respect to the lowest MSE, **DM** is the Diebold and Mariano test for the null hypothesis of equal predictive accuracy (Diebold and Mariano, 1995), and **MCS** is the Model Confidence Set approach by Hansen, Lunde and Nason (2011). $G_{w4,t}^4$ is the Google Index for keyword 'dos', the keyword whose GI is most correlated with the US unemployment rate but otherwise unrelated to job search, see section 5.5 for details). The column 'MCS' has a 1 when the row model is included in the final model confidence set at 5% confidence level and a 0 otherwise. In all panels ***, ** and * indicate rejection at 1, 5 and 10%, respectively.

Table 9: Forecasts of the quarterly US unemployment: comparison of the best models with the Survey of Professional Forecasters.

Model	MSE	Rank	DM (vs Best)	DM (vs $G^{1st-month}$)	DM (vs $G^{2nd-month}$)
SPF^{best}	0.2362	21	3.2764***	2.7584***	2.8581***
SPF^{mean}	0.0453	9	2.0589**	-0.3351	0.1073
SPF^{med}	0.0413	7	2.0943**	-0.6244	-0.1931
$G^{1st-month}$	0.0499	10	2.0538**		0.4305
$G^{2nd-month}$	0.0436	8	2.4386**	-0.4305	
G^{Comb}	0.0136	1		-2.0538**	-2.4386**
$NG_L^{1st-month}$	0.0822	13	2.5102**	1.6321	1.683*
$NG_L^{2nd-month}$	0.0586	11	2.2419**	0.4511	0.9749
NG_L^{Comb}	0.0180	2	0.9701	-1.9187*	-2.3255**
$NG_S^{1st-month}$	0.0769	12	2.3967**	2.0415**	1.5351
$NG_S^{2nd-month}$	0.1497	20	2.868***	2.239**	2.2873**
NG_S^{Comb}	0.0269	3	1.5281	-1.0637	-0.9759
$SETAR^{1st-month}$	0.1079	18	2.4294**	1.8975*	1.807*
$SETAR^{2nd-month}$	0.1011	15	2.4565**	1.8878*	1.7857*
$SETAR^{Comb}$	0.0347	4	2.5766***	-0.862	-0.5875
$LSTAR^{1st-month}$	0.1070	17	2.3766**	1.8363*	1.7627*
$LSTAR^{2nd-month}$	0.1005	14	2.3765**	1.7993*	1.7234*
$LSTAR^{Comb}$	0.0362	5	2.5419**	-0.8069	-0.4819
$AAR^{1st-month}$	0.1115	19	2.323**	1.7745*	1.7434*
$AAR^{2nd-month}$	0.1065	16	2.2844**	1.7005*	1.6834*
AAR^{Comb}	0.0370	6	2.4204**	-0.716	-0.4084

Notes: In this table we compare the SPF one-quarter-ahead unemployment forecasts with similar forecasts generated from our best models for u_t , i.e. models n. ..., and ... for 1-, 2- and 3-month-ahead forecasts, respectively. The out-of-sample period is 2007:Q2-2011:Q2. SPF^{best} is the best individual forecaster in the survey, SPF^{mean} is the mean of the forecasts, while SPF^{median} is the median. Models $x^{1st-month}$ are 1-month-ahead forecasts computed in the last month of the quarter before. Models $x^{2nd-month}$ are 2-month-ahead forecasts computed in the last month of the quarter before. Both these forecasts are very conservative since the SPF is issued on the 15th of the second month of each reference quarter. Models x^{Comb} compute the quarterly forecast as the average of the realized unemployment rate for the first month and the 1- and 2-month-ahead forecasts generated at the end of the first month of the reference quarter. The model with Google is the best model overall, the model with NG_L is the best model without Google on the long sample, while the model with subscript NG_S is the best model without Google in the short sample. SETAR, LSTAR and AAR are the corresponding non-linear models estimated over the full sample up to the second lag. In boldface we indicate the model with the minimum MSE, while in italics the next to the minimum MSE. The benchmark model for the DM and HLN tests is G^{Comb} . ***, **, * indicate rejection at 1, 5 and 10%, respectively.

Figure 1: Relative incidence of keyword searches through Google



Notes: The figure depicts the relative incidence of the web searches for the keyword 'jobs' adopted to construct our preferred Google index along with the other job-related searches 'collect unemployment' and 'job center' (which are almost nil), and the recently more popular 'facebook' and 'youtube' keywords over the relevant sample 2004.1-2011.6.

Figure 2: Exact timing of monthly US Unemployment rate calculation

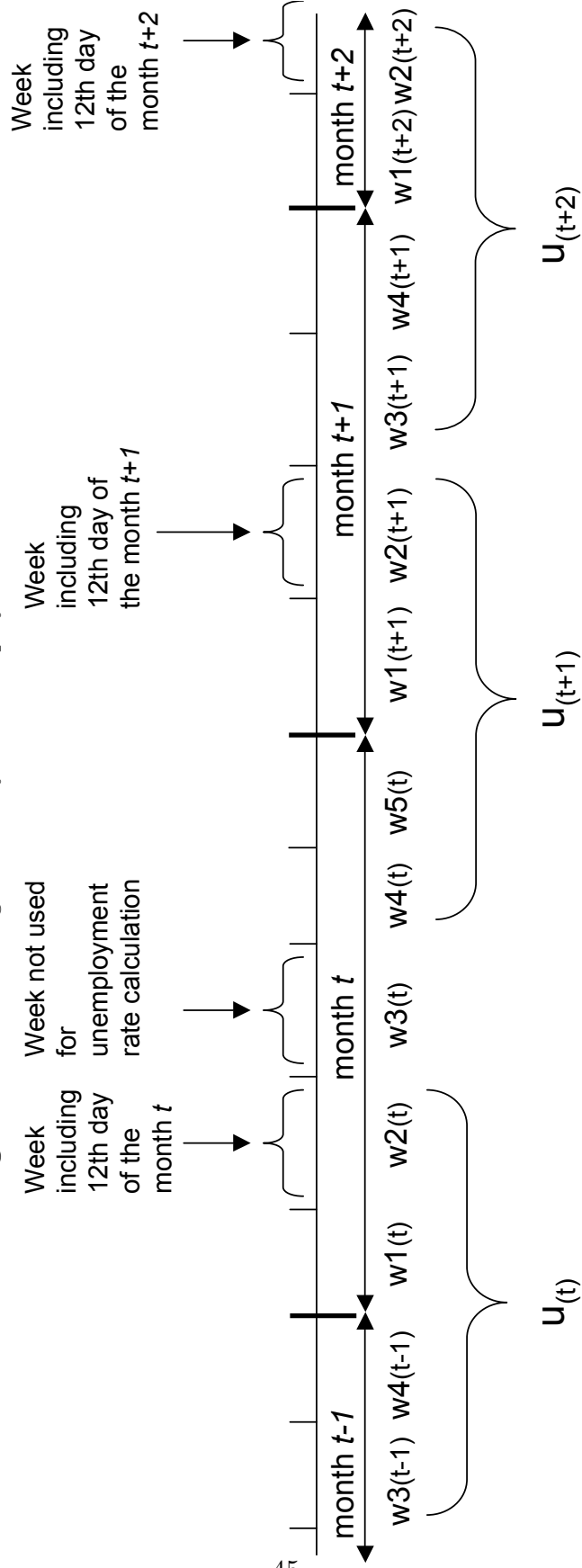
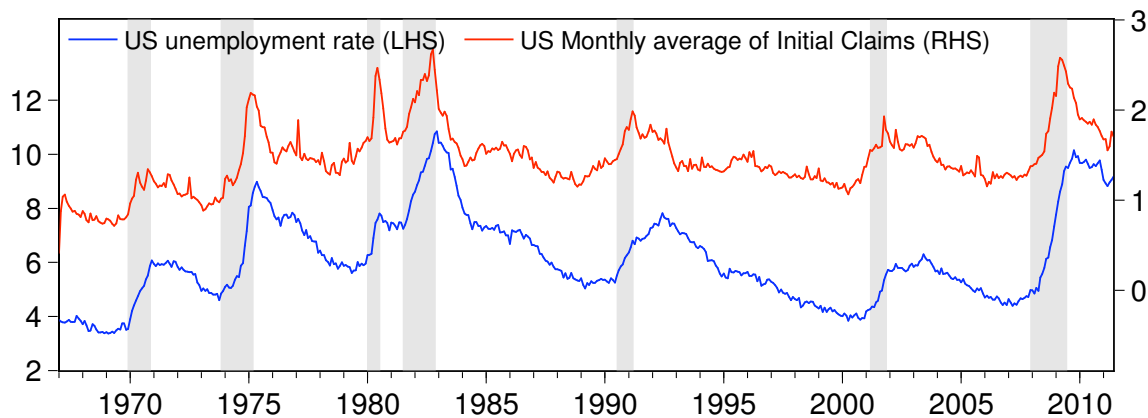
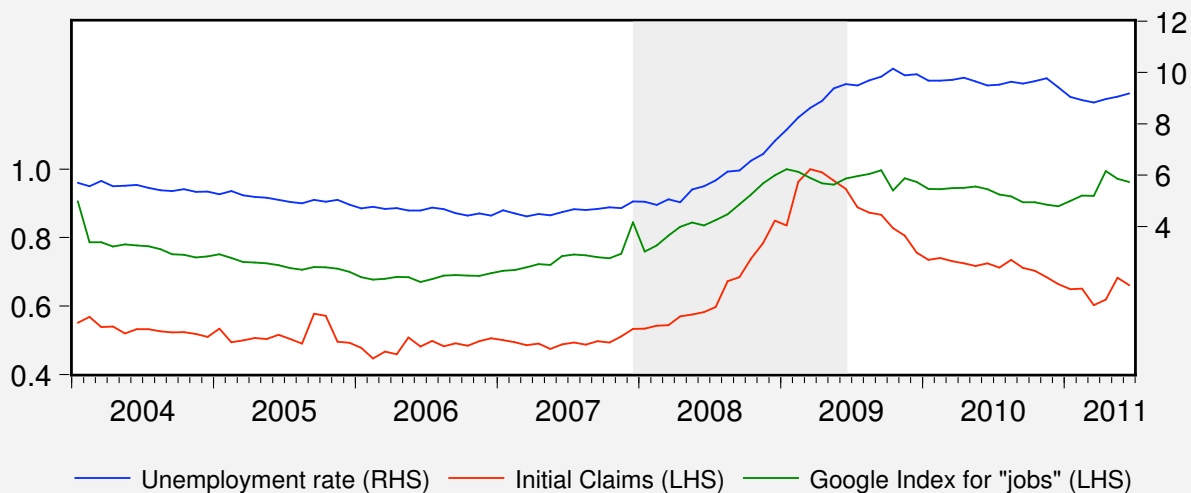


Figure 3: US Unemployment rate and Initial Claims: Long sample 1967.1-2011.6



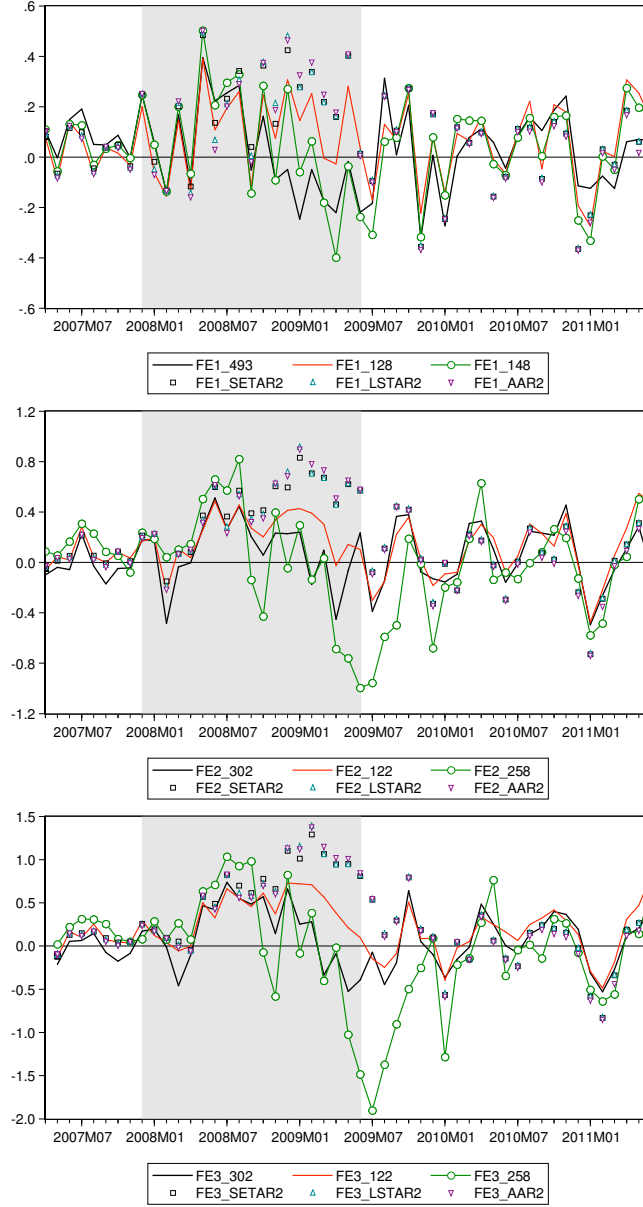
Notes: Shaded areas identify official NBER recessions.

Figure 4: US Unemployment rate, Initial claims and Google Index: Short sample 2004.1-2011.6



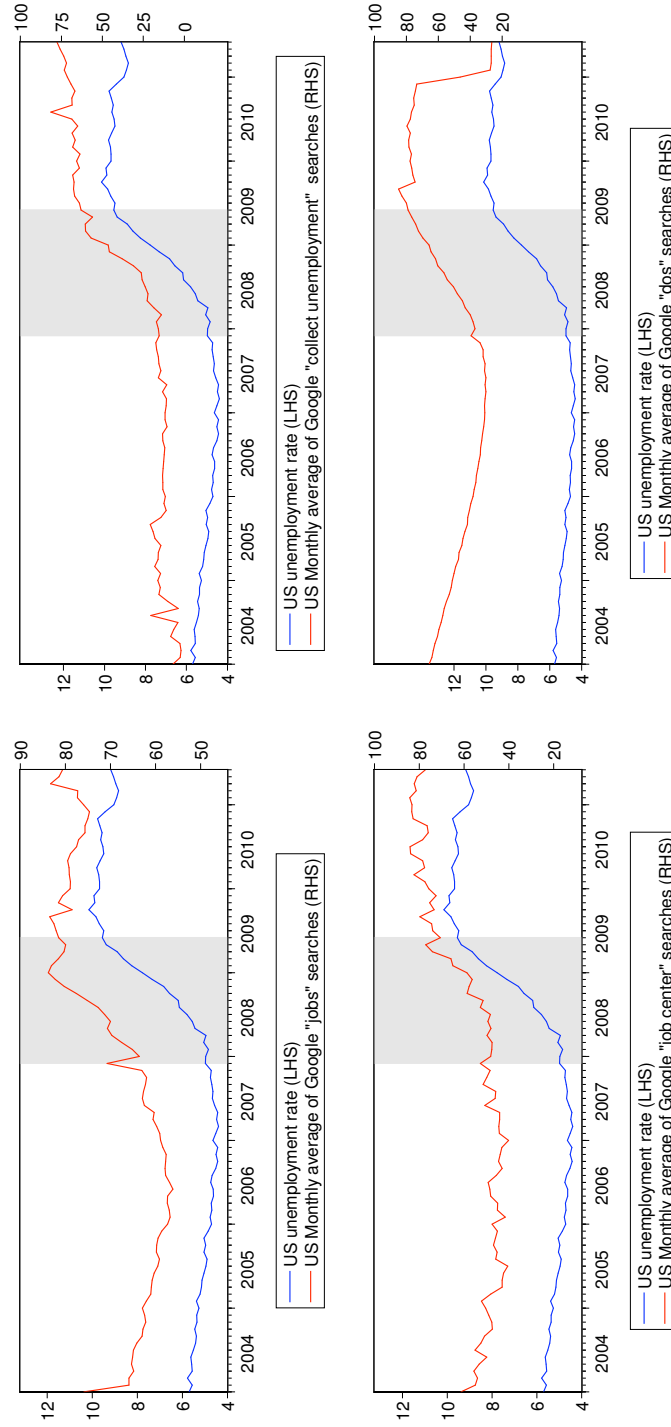
Notes: Shaded areas identify NBER recessions. The Initial claims are monthly averages rebased on their maximum over the sample 2004.1-2011.6. The Google Index is the monthly average of Google 'jobs' searches rebased on their weekly maximum value over the sample 2004.1-2011.6.

Figure 5: Forecast error comparison of the best models with and without the GI over the short and long sample and the non-linear models



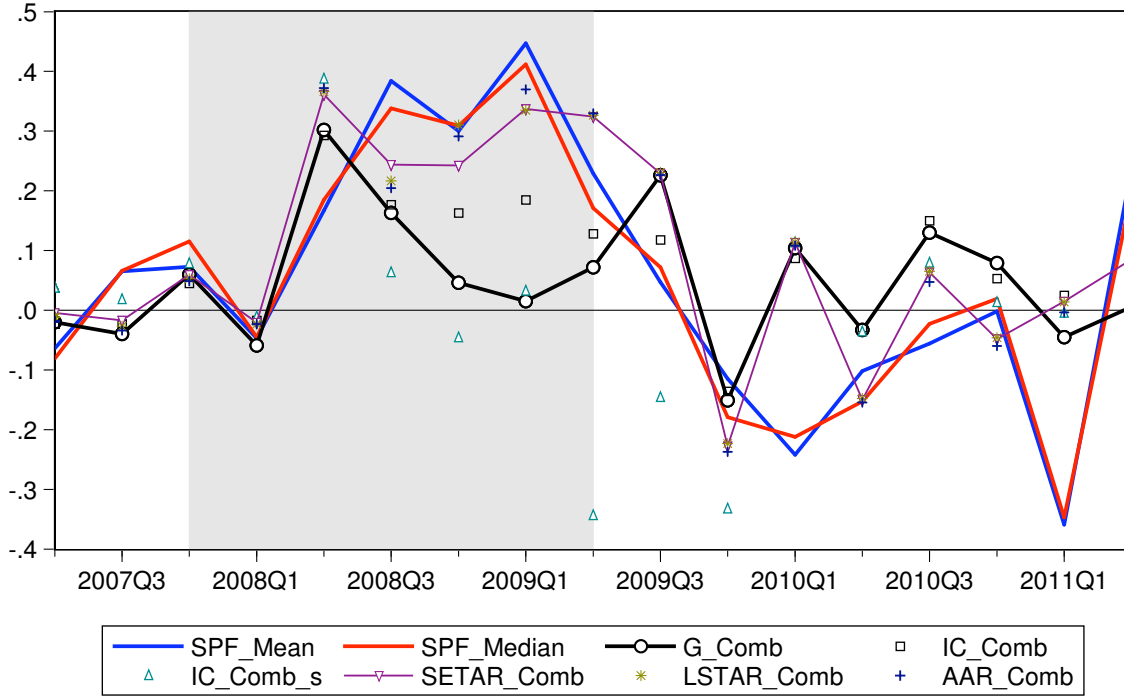
Notes: In this figure we compare the 1-, 2- and 3-step-ahead forecast errors generated by our best models in the top, middle and bottom panel, respectively. For each panel we present the forecast errors for the best models overall using the GI (i.e. models n. 493, 302 and 302 for 1-, 2- and 3-month-ahead forecasts, respectively), our best non-Google models over the long sample (i.e. models n. 128, 122 and 122 for 1-, 2- and 3-month-ahead forecasts, respectively) and our best non-Google models over the short sample (i.e. models n. 148, 258 and 258 for 1-, 2- and 3-month-ahead forecasts, respectively). The out-of-sample period is 2007.3-2011.6. SETAR, LSTAR and AAR are the corresponding non-linear models estimated over the long sample up to the second lag.

Figure 6: US Unemployment rate and Google indexes: Short sample 2004.1-2011.6



Notes: Shaded areas identify NBER recessions. The Google index is the monthly average of weekly Google searches for ‘jobs’, ‘collect unemployment’, ‘job center’, and ‘dos’ (the false index). Sample: 2004.1-2011.6.

Figure 7: Forecast errors from quarterly forecasts of the US unemployment rate: comparison of the best models with the Survey of Professional Forecasters



Notes: In this table we compare the SPF one-quarter-ahead unemployment forecasts with similar forecasts generated from our best models for u_t , i.e. models n. 261, 261 and 398 for 1-, 2- and 3-month-ahead forecasts, respectively. The out-of-sample period is 2007.Q2-2011.Q2. SPF^{best} is the best individual forecaster in the survey, SPF^{mean} is the mean of the forecasts, while SPF^{median} is the median. Models $x^{1st-month}$ are 1-month-ahead forecasts computed in the last month of the quarter before. Models $x^{2nd-month}$ are 2-month-ahead forecasts computed in the last month of the quarter before. Both these forecasts are very conservative because the SPF is issued on the 15th of the second month of each reference quarter. Models x^{Comb} compute the quarterly forecast as the average of the realized unemployment rate for the first month and the 1- and 2-month-ahead forecasts generated at the end of the first month of the reference quarter. The model with Google (G) is the best model overall, the model with the Initial Claims (IC) is the best model without Google, while the models with subscript IC_s is the best model without Google in the short sample. SETAR, LSTAR and AAR are the corresponding non-linear models estimated over the full sample up to the second lag.

RECENTLY PUBLISHED “TEMI” (*)

- N. 868 – *The economic costs of organized crime: evidence from southern Italy*, by Paolo Pinotti (April 2012).
- N. 869 – *Network effects of public transport infrastructure: evidence on Italian regions*, by Valter Di Giacinto, Giacinto Micucci and Pasqualino Montanaro (July 2012).
- N. 870 – *To misreport or not to report? The measurement of household financial wealth*, by Andrea Neri and Maria Giovanna Ranalli (July 2012).
- N. 871 – *Capital destruction, jobless recoveries, and the discipline device role of unemployment*, by Marianna Riggi (July 2012).
- N. 872 – *Selecting predictors by using Bayesian model averaging in bridge models*, by Lorenzo Bencivelli, Massimiliano Marcellino and Gianluca Moretti (July 2012).
- N. 873 – *Euro area and global oil shocks: an empirical model-based analysis*, by Lorenzo Forni, Andrea Gerali, Alessandro Notarpietro and Massimiliano Pisani (July 2012).
- N. 874 – *Evidence on the impact of R&D and ICT investment on innovation and productivity in Italian firms*, by Bronwyn H. Hall, Francesca Lotti and Jacques Mairesse (July 2012).
- N. 875 – *Family background, self-confidence and economic outcomes*, by Antonio Filippin and Marco Paccagnella (July 2012).
- N. 876 – *Banks’ reactions to Basel-III*, by Paolo Angelini and Andrea Gerali (July 2012).
- N. 877 – *Exporters and importers of services: firm-level evidence on Italy*, by Stefano Federico and Enrico Tosti (September 2012).
- N. 878 – *Do food commodity prices have asymmetric effects on euro-area inflation?*, by Mario Porqueddu and Fabrizio Venditti (September 2012).
- N. 879 – *Industry dynamics and competition from low-wage countries: evidence on Italy*, by Stefano Federico (September 2012).
- N. 880 – *The micro dynamics of exporting: evidence from French firms*, by Ines Buono and Harald Fadinger (September 2012).
- N. 881 – *On detecting end-of-sample instabilities*, by Fabio Buseti (September 2012).
- N. 882 – *An empirical comparison of alternative credit default swap pricing models*, by Michele Leonardo Bianchi (September 2012).
- N. 883 – *Learning, incomplete contracts and export dynamics: theory and evidence from French firms*, by Romain Aeberhardt, Ines Buono and Harald Fadinger (October 2012).
- N. 884 – *Collaboration between firms and universities in Italy: the role of a firm’s proximity to top-rated departments*, by Davide Fantino, Alessandra Mori and Diego Scalise (October 2012).
- N. 885 – *Parties, institutions and political budget cycles at the municipal level*, by Marika Cioffi, Giovanna Messina and Pietro Tommasino (October 2012).
- N. 886 – *Immigration, jobs and employment protection: evidence from Europe before and during the Great Recession*, by Francesco D’Amuri and Giovanni Peri (October 2012).
- N. 887 – *A structural model for the housing and credit markets in Italy*, by Andrea Nobili and Francesco Zollino (October 2012).
- N. 888 – *Monetary policy in a model with misspecified, heterogeneous and ever-changing expectations*, by Alberto Locarno (October 2012).

(*) Requests for copies should be sent to:

Banca d’Italia – Servizio Studi di struttura economica e finanziaria – Divisione Biblioteca e Archivio storico – Via Nazionale, 91 – 00184 Rome – (fax 0039 06 47922059). They are available on the Internet www.bancaditalia.it.

2009

- F. PANETTA, F. SCHIVARDI and M. SHUM, *Do mergers improve information? Evidence from the loan market*, Journal of Money, Credit, and Banking, v. 41, 4, pp. 673-709, **TD No. 521 (October 2004)**.
- M. BUGAMELLI and F. PATERNÒ, *Do workers' remittances reduce the probability of current account reversals?*, World Development, v. 37, 12, pp. 1821-1838, **TD No. 573 (January 2006)**.
- P. PAGANO and M. PISANI, *Risk-adjusted forecasts of oil prices*, The B.E. Journal of Macroeconomics, v. 9, 1, Article 24, **TD No. 585 (March 2006)**.
- M. PERICOLI and M. SBRACIA, *The CAPM and the risk appetite index: theoretical differences, empirical similarities, and implementation problems*, International Finance, v. 12, 2, pp. 123-150, **TD No. 586 (March 2006)**.
- R. BRONZINI and P. PISELLI, *Determinants of long-run regional productivity with geographical spillovers: the role of R&D, human capital and public infrastructure*, Regional Science and Urban Economics, v. 39, 2, pp. 187-199, **TD No. 597 (September 2006)**.
- U. ALBERTAZZI and L. GAMBACORTA, *Bank profitability and the business cycle*, Journal of Financial Stability, v. 5, 4, pp. 393-409, **TD No. 601 (September 2006)**.
- F. BALASSONE, D. FRANCO and S. ZOTTERI, *The reliability of EMU fiscal indicators: risks and safeguards*, in M. Larch and J. Nogueira Martins (eds.), Fiscal Policy Making in the European Union: an Assessment of Current Practice and Challenges, London, Routledge, **TD No. 633 (June 2007)**.
- A. CIARLONE, P. PISELLI and G. TREBESCHI, *Emerging Markets' Spreads and Global Financial Conditions*, Journal of International Financial Markets, Institutions & Money, v. 19, 2, pp. 222-239, **TD No. 637 (June 2007)**.
- S. MAGRI, *The financing of small innovative firms: the Italian case*, Economics of Innovation and New Technology, v. 18, 2, pp. 181-204, **TD No. 640 (September 2007)**.
- V. DI GIACINTO and G. MICUCCI, *The producer service sector in Italy: long-term growth and its local determinants*, Spatial Economic Analysis, Vol. 4, No. 4, pp. 391-425, **TD No. 643 (September 2007)**.
- F. LORENZO, L. MONTEFORTE and L. SESSA, *The general equilibrium effects of fiscal policy: estimates for the euro area*, Journal of Public Economics, v. 93, 3-4, pp. 559-585, **TD No. 652 (November 2007)**.
- Y. ALTUNBAS, L. GAMBACORTA and D. MARQUÉS, *Securitisation and the bank lending channel*, European Economic Review, v. 53, 8, pp. 996-1009, **TD No. 653 (November 2007)**.
- R. GOLINELLI and S. MOMIGLIANO, *The Cyclical Reaction of Fiscal Policies in the Euro Area. A Critical Survey of Empirical Research*, Fiscal Studies, v. 30, 1, pp. 39-72, **TD No. 654 (January 2008)**.
- P. DEL GIOVANE, S. FABIANI and R. SABBATINI, *What's behind "Inflation Perceptions"? A survey-based analysis of Italian consumers*, Giornale degli Economisti e Annali di Economia, v. 68, 1, pp. 25-52, **TD No. 655 (January 2008)**.
- F. MACCHERONI, M. MARINACCI, A. RUSTICHINI and M. TABOGA, *Portfolio selection with monotone mean-variance preferences*, Mathematical Finance, v. 19, 3, pp. 487-521, **TD No. 664 (April 2008)**.
- M. AFFINITO and M. PIAZZA, *What are borders made of? An analysis of barriers to European banking integration*, in P. Alessandrini, M. Fratianni and A. Zazzaro (eds.): The Changing Geography of Banking and Finance, Dordrecht Heidelberg London New York, Springer, **TD No. 666 (April 2008)**.
- A. BRANDOLINI, *On applying synthetic indices of multidimensional well-being: health and income inequalities in France, Germany, Italy, and the United Kingdom*, in R. Gotoh and P. Dumouchel (eds.), Against Injustice. The New Economics of Amartya Sen, Cambridge, Cambridge University Press, **TD No. 668 (April 2008)**.
- G. FERRERO and A. NOBILI, *Futures contract rates as monetary policy forecasts*, International Journal of Central Banking, v. 5, 2, pp. 109-145, **TD No. 681 (June 2008)**.
- P. CASADIO, M. LO CONTE and A. NERI, *Balancing work and family in Italy: the new mothers' employment decisions around childbearing*, in T. Addabbo and G. Solinas (eds.), Non-Standard Employment and Qualità of Work, Physica-Verlag. A Springer Company, **TD No. 684 (August 2008)**.
- L. ARCIERO, C. BIANCOTTI, L. D'AURIZIO and C. IMPENNA, *Exploring agent-based methods for the analysis of payment systems: A crisis model for StarLogo TNG*, Journal of Artificial Societies and Social Simulation, v. 12, 1, **TD No. 686 (August 2008)**.
- A. CALZA and A. ZAGHINI, *Nonlinearities in the dynamics of the euro area demand for M1*, Macroeconomic Dynamics, v. 13, 1, pp. 1-19, **TD No. 690 (September 2008)**.
- L. FRANCESCO and A. SECCHI, *Technological change and the households' demand for currency*, Journal of Monetary Economics, v. 56, 2, pp. 222-230, **TD No. 697 (December 2008)**.

- G. ASCARI and T. ROPELE, *Trend inflation, taylor principle, and indeterminacy*, Journal of Money, Credit and Banking, v. 41, 8, pp. 1557-1584, **TD No. 708 (May 2007)**.
- S. COLAROSSO and A. ZAGHINI, *Gradualism, transparency and the improved operational framework: a look at overnight volatility transmission*, International Finance, v. 12, 2, pp. 151-170, **TD No. 710 (May 2009)**.
- M. BUGAMELLI, F. SCHIVARDI and R. ZIZZA, *The euro and firm restructuring*, in A. Alesina e F. Giavazzi (eds): Europe and the Euro, Chicago, University of Chicago Press, **TD No. 716 (June 2009)**.
- B. HALL, F. LOTTI and J. MAIRESSE, *Innovation and productivity in SMEs: empirical evidence for Italy*, Small Business Economics, v. 33, 1, pp. 13-33, **TD No. 718 (June 2009)**.

2010

- A. PRATI and M. SBRACIA, *Uncertainty and currency crises: evidence from survey data*, Journal of Monetary Economics, v. 57, 6, pp. 668-681, **TD No. 446 (July 2002)**.
- L. MONTEFORTE and S. SIVIERO, *The Economic Consequences of Euro Area Modelling Shortcuts*, Applied Economics, v. 42, 19-21, pp. 2399-2415, **TD No. 458 (December 2002)**.
- S. MAGRI, *Debt maturity choice of nonpublic Italian firms*, Journal of Money, Credit, and Banking, v.42, 2-3, pp. 443-463, **TD No. 574 (January 2006)**.
- G. DE BLASIO and G. NUZZO, *Historical traditions of civicness and local economic development*, Journal of Regional Science, v. 50, 4, pp. 833-857, **TD No. 591 (May 2006)**.
- E. IOSSA and G. PALUMBO, *Over-optimism and lender liability in the consumer credit market*, Oxford Economic Papers, v. 62, 2, pp. 374-394, **TD No. 598 (September 2006)**.
- S. NERI and A. NOBILI, *The transmission of US monetary policy to the euro area*, International Finance, v. 13, 1, pp. 55-78, **TD No. 606 (December 2006)**.
- F. ALTISSIMO, R. CRISTADORO, M. FORNI, M. LIPPI and G. VERONESE, *New Eurocoin: Tracking Economic Growth in Real Time*, Review of Economics and Statistics, v. 92, 4, pp. 1024-1034, **TD No. 631 (June 2007)**.
- U. ALBERTAZZI and L. GAMBACORTA, *Bank profitability and taxation*, Journal of Banking and Finance, v. 34, 11, pp. 2801-2810, **TD No. 649 (November 2007)**.
- M. IACOVIELLO and S. NERI, *Housing market spillovers: evidence from an estimated DSGE model*, American Economic Journal: Macroeconomics, v. 2, 2, pp. 125-164, **TD No. 659 (January 2008)**.
- F. BALASSONE, F. MAURA and S. ZOTTERI, *Cyclical asymmetry in fiscal variables in the EU*, Empirica, **TD No. 671**, v. 37, 4, pp. 381-402 **(June 2008)**.
- F. D'AMURI, O. GIANMARCO I.P. and P. GIOVANNI, *The labor market impact of immigration on the western german labor market in the 1990s*, European Economic Review, v. 54, 4, pp. 550-570, **TD No. 687 (August 2008)**.
- A. ACCETTURO, *Agglomeration and growth: the effects of commuting costs*, Papers in Regional Science, v. 89, 1, pp. 173-190, **TD No. 688 (September 2008)**.
- S. NOBILI and G. PALAZZO, *Explaining and forecasting bond risk premiums*, Financial Analysts Journal, v. 66, 4, pp. 67-82, **TD No. 689 (September 2008)**.
- A. B. ATKINSON and A. BRANDOLINI, *On analysing the world distribution of income*, World Bank Economic Review, v. 24, 1, pp. 1-37, **TD No. 701 (January 2009)**.
- R. CAPPARIELLO and R. ZIZZA, *Dropping the Books and Working Off the Books*, Labour, v. 24, 2, pp. 139-162, **TD No. 702 (January 2009)**.
- C. NICOLETTI and C. RONDINELLI, *The (mis)specification of discrete duration models with unobserved heterogeneity: a Monte Carlo study*, Journal of Econometrics, v. 159, 1, pp. 1-13, **TD No. 705 (March 2009)**.
- L. FORNI, A. GERALI and M. PISANI, *Macroeconomic effects of greater competition in the service sector: the case of Italy*, Macroeconomic Dynamics, v. 14, 5, pp. 677-708, **TD No. 706 (March 2009)**.
- V. DI GIACINTO, G. MICUCCI and P. MONTANARO, *Dynamic macroeconomic effects of public capital: evidence from regional Italian data*, Giornale degli economisti e annali di economia, v. 69, 1, pp. 29-66, **TD No. 733 (November 2009)**.
- F. COLUMBA, L. GAMBACORTA and P. E. MISTRULLI, *Mutual Guarantee institutions and small business finance*, Journal of Financial Stability, v. 6, 1, pp. 45-54, **TD No. 735 (November 2009)**.
- A. GERALI, S. NERI, L. SESSA and F. M. SIGNORETTI, *Credit and banking in a DSGE model of the Euro Area*, Journal of Money, Credit and Banking, v. 42, 6, pp. 107-141, **TD No. 740 (January 2010)**.
- M. AFFINITO and E. TAGLIAFERRI, *Why do (or did?) banks securitize their loans? Evidence from Italy*, Journal

- of Financial Stability, v. 6, 4, pp. 189-202, **TD No. 741 (January 2010)**.
- S. FEDERICO, *Outsourcing versus integration at home or abroad and firm heterogeneity*, Empirica, v. 37, 1, pp. 47-63, **TD No. 742 (February 2010)**.
- V. DI GIACINTO, *On vector autoregressive modeling in space and time*, Journal of Geographical Systems, v. 12, 2, pp. 125-154, **TD No. 746 (February 2010)**.
- L. FORNI, A. GERALI and M. PISANI, *The macroeconomics of fiscal consolidations in euro area countries*, Journal of Economic Dynamics and Control, v. 34, 9, pp. 1791-1812, **TD No. 747 (March 2010)**.
- S. MOCETTI and C. PORELLO, *How does immigration affect native internal mobility? new evidence from Italy*, Regional Science and Urban Economics, v. 40, 6, pp. 427-439, **TD No. 748 (March 2010)**.
- A. DI CESARE and G. GUAZZAROTTI, *An analysis of the determinants of credit default swap spread changes before and during the subprime financial turmoil*, Journal of Current Issues in Finance, Business and Economics, v. 3, 4, pp., **TD No. 749 (March 2010)**.
- P. CIPOLLONE, P. MONTANARO and P. SESTITO, *Value-added measures in Italian high schools: problems and findings*, Giornale degli economisti e annali di economia, v. 69, 2, pp. 81-114, **TD No. 754 (March 2010)**.
- A. BRANDOLINI, S. MAGRI and T. M. SMEEDING, *Asset-based measurement of poverty*, Journal of Policy Analysis and Management, v. 29, 2, pp. 267-284, **TD No. 755 (March 2010)**.
- G. CAPPELLETTI, *A Note on rationalizability and restrictions on beliefs*, The B.E. Journal of Theoretical Economics, v. 10, 1, pp. 1-11, **TD No. 757 (April 2010)**.
- S. DI ADDARIO and D. VURI, *Entrepreneurship and market size. the case of young college graduates in Italy*, Labour Economics, v. 17, 5, pp. 848-858, **TD No. 775 (September 2010)**.
- A. CALZA and A. ZAGHINI, *Sectoral money demand and the great disinflation in the US*, Journal of Money, Credit, and Banking, v. 42, 8, pp. 1663-1678, **TD No. 785 (January 2011)**.

2011

- S. DI ADDARIO, *Job search in thick markets*, Journal of Urban Economics, v. 69, 3, pp. 303-318, **TD No. 605 (December 2006)**.
- F. SCHIVARDI and E. VIVIANO, *Entry barriers in retail trade*, Economic Journal, v. 121, 551, pp. 145-170, **TD No. 616 (February 2007)**.
- G. FERRERO, A. NOBILI and P. PASSIGLIA, *Assessing excess liquidity in the Euro Area: the role of sectoral distribution of money*, Applied Economics, v. 43, 23, pp. 3213-3230, **TD No. 627 (April 2007)**.
- P. E. MISTRULLI, *Assessing financial contagion in the interbank market: maximum entropy versus observed interbank lending patterns*, Journal of Banking & Finance, v. 35, 5, pp. 1114-1127, **TD No. 641 (September 2007)**.
- E. CIAPANNA, *Directed matching with endogenous markov probability: clients or competitors?*, The RAND Journal of Economics, v. 42, 1, pp. 92-120, **TD No. 665 (April 2008)**.
- M. BUGAMELLI and F. PATERNÒ, *Output growth volatility and remittances*, Economica, v. 78, 311, pp. 480-500, **TD No. 673 (June 2008)**.
- V. DI GIACINTO e M. PAGNINI, *Local and global agglomeration patterns: two econometrics-based indicators*, Regional Science and Urban Economics, v. 41, 3, pp. 266-280, **TD No. 674 (June 2008)**.
- G. BARONE and F. CINGANO, *Service regulation and growth: evidence from OECD countries*, Economic Journal, v. 121, 555, pp. 931-957, **TD No. 675 (June 2008)**.
- R. GIORDANO and P. TOMMASINO, *What determines debt intolerance? The role of political and monetary institutions*, European Journal of Political Economy, v. 27, 3, pp. 471-484, **TD No. 700 (January 2009)**.
- P. ANGELINI, A. NOBILI e C. PICILLO, *The interbank market after August 2007: What has changed, and why?*, Journal of Money, Credit and Banking, v. 43, 5, pp. 923-958, **TD No. 731 (October 2009)**.
- L. FORNI, A. GERALI and M. PISANI, *The Macroeconomics of Fiscal Consolidation in a Monetary Union: the Case of Italy*, in Luigi Paganetto (ed.), Recovery after the crisis. Perspectives and policies, VDM Verlag Dr. Muller, **TD No. 747 (March 2010)**.
- A. DI CESARE and G. GUAZZAROTTI, *An analysis of the determinants of credit default swap changes before and during the subprime financial turmoil*, in Barbara L. Campos and Janet P. Wilkins (eds.), The Financial Crisis: Issues in Business, Finance and Global Economics, New York, Nova Science Publishers, Inc., **TD No. 749 (March 2010)**.
- A. LEVY and A. ZAGHINI, *The pricing of government guaranteed bank bonds*, Banks and Bank Systems, v. 6, 3, pp. 16-24, **TD No. 753 (March 2010)**.

- G. GRANDE and I. VISCO, *A public guarantee of a minimum return to defined contribution pension scheme members*, *The Journal of Risk*, v. 13, 3, pp. 3-43, **TD No. 762 (June 2010)**.
- P. DEL GIOVANE, G. ERAMO and A. NOBILI, *Disentangling demand and supply in credit developments: a survey-based analysis for Italy*, *Journal of Banking and Finance*, v. 35, 10, pp. 2719-2732, **TD No. 764 (June 2010)**.
- G. BARONE and S. MOCETTI, *With a little help from abroad: the effect of low-skilled immigration on the female labour supply*, *Labour Economics*, v. 18, 5, pp. 664-675, **TD No. 766 (July 2010)**.
- A. FELETTIGH and S. FEDERICO, *Measuring the price elasticity of import demand in the destination markets of italian exports*, *Economia e Politica Industriale*, v. 38, 1, pp. 127-162, **TD No. 776 (October 2010)**.
- S. MAGRI and R. PICO, *The rise of risk-based pricing of mortgage interest rates in Italy*, *Journal of Banking and Finance*, v. 35, 5, pp. 1277-1290, **TD No. 778 (October 2010)**.
- M. TABOGA, *Under/over-valuation of the stock market and cyclically adjusted earnings*, *International Finance*, v. 14, 1, pp. 135-164, **TD No. 780 (December 2010)**.
- S. NERI, *Housing, consumption and monetary policy: how different are the U.S. and the Euro area?*, *Journal of Banking and Finance*, v.35, 11, pp. 3019-3041, **TD No. 807 (April 2011)**.
- V. CUCINIELLO, *The welfare effect of foreign monetary conservatism with non-atomistic wage setters*, *Journal of Money, Credit and Banking*, v. 43, 8, pp. 1719-1734, **TD No. 810 (June 2011)**.
- A. CALZA and A. ZAGHINI, *welfare costs of inflation and the circulation of US currency abroad*, *The B.E. Journal of Macroeconomics*, v. 11, 1, Art. 12, **TD No. 812 (June 2011)**.
- I. FAIELLA, *La spesa energetica delle famiglie italiane*, *Energia*, v. 32, 4, pp. 40-46, **TD No. 822 (September 2011)**.
- R. DE BONIS and A. SILVESTRINI, *The effects of financial and real wealth on consumption: new evidence from OECD countries*, *Applied Financial Economics*, v. 21, 5, pp. 409-425, **TD No. 837 (November 2011)**.

2012

- F. CINGANO and A. ROSOLIA, *People I know: job search and social networks*, *Journal of Labor Economics*, v. 30, 2, pp. 291-332, **TD No. 600 (September 2006)**.
- G. GOBBI and R. ZIZZA, *Does the underground economy hold back financial deepening? Evidence from the italian credit market*, *Economia Marche, Review of Regional Studies*, v. 31, 1, pp. 1-29, **TD No. 646 (November 2006)**.
- S. MOCETTI, *Educational choices and the selection process before and after compulsory school*, *Education Economics*, v. 20, 2, pp. 189-209, **TD No. 691 (September 2008)**.
- A. ACCETTURO and G. DE BLASIO, *Policies for local development: an evaluation of Italy's "Patti Territoriali"*, *Regional Science and Urban Economics*, v. 42, 1-2, pp. 15-26, **TD No. 789 (January 2006)**.
- F. BUSETTI and S. DI SANZO, *Bootstrap LR tests of stationarity, common trends and cointegration*, *Journal of Statistical Computation and Simulation*, v. 82, 9, pp. 1343-1355, **TD No. 799 (March 2006)**.
- S. NERI and T. ROPELE, *Imperfect information, real-time data and monetary policy in the Euro area*, *The Economic Journal*, v. 122, 561, pp. 651-674, **TD No. 802 (March 2011)**.
- A. ANZUINI and F. FORNARI, *Macroeconomic determinants of carry trade activity*, *Review of International Economics*, v. 20, 3, pp. 468-488, **TD No. 817 (September 2011)**.
- R. CRISTADORO and D. MARCONI, *Household savings in China*, *Journal of Chinese Economic and Business Studies*, v. 10, 3, pp. 275-299, **TD No. 838 (November 2011)**.
- A. FILIPPIN and M. PACCAGNELLA, *Family background, self-confidence and economic outcomes*, *Economics of Education Review*, v. 31, 5, pp. 824-834, **TD No. 875 (July 2012)**.

FORTHCOMING

- M. BUGAMELLI and A. ROSOLIA, *Produttività e concorrenza estera*, *Rivista di politica economica*, **TD No. 578 (February 2006)**.
- P. SESTITO and E. VIVIANO, *Reservation wages: explaining some puzzling regional patterns*, *Labour*, **TD No. 696 (December 2008)**.

- P. PINOTTI, M. BIANCHI and P. BUONANNO, *Do immigrants cause crime?*, Journal of the European Economic Association, **TD No. 698 (December 2008)**.
- F. LIPPI and A. NOBILI, *Oil and the macroeconomy: a quantitative structural analysis*, Journal of European Economic Association, **TD No. 704 (March 2009)**.
- F. CINGANO and P. PINOTTI, *Politicians at work. The private returns and social costs of political connections*, Journal of the European Economic Association, **TD No. 709 (May 2009)**.
- Y. ALTUNBAS, L. GAMBACORTA, and D. MARQUÉS-IBÁÑEZ, *Bank risk and monetary policy*, Journal of Financial Stability, **TD No. 712 (May 2009)**.
- G. BARONE and S. MOCETTI, *Tax morale and public spending inefficiency*, International Tax and Public Finance, **TD No. 732 (November 2009)**.
- S. FEDERICO, *Headquarter intensity and the choice between outsourcing versus integration at home or abroad*, Industrial and Corporate Change, **TD No. 742 (February 2010)**.
- I. BUONO and G. LALANNE, *The effect of the Uruguay Round on the intensive and extensive margins of trade*, Journal of International Economics, **TD No. 835 (February 2011)**.
- G. BARONE, R. FELICI and M. PAGNINI, *Switching costs in local credit markets*, International Journal of Industrial Organization, **TD No. 760 (June 2010)**.
- E. COCOZZA and P. PISELLI, *Testing for east-west contagion in the European banking sector during the financial crisis*, in R. Matoušek; D. Stavárek (eds.), *Financial Integration in the European Union*, Taylor & Francis, **TD No. 790 (February 2011)**.
- A. DE SOCIO, *Squeezing liquidity in a “lemons market” or asking liquidity “on tap”*, Journal of Banking and Finance, **TD No. 819 (September 2011)**.
- M. AFFINITO, *Do interbank customer relationships exist? And how did they function in the crisis? Learning from Italy*, Journal of Banking and Finance, **TD No. 826 (October 2011)**.
- O. BLANCHARD and M. RIGGI, *Why are the 2000s so different from the 1970s? A structural interpretation of changes in the macroeconomic effects of oil prices*, Journal of the European Economic Association, **TD No. 835 (November 2011)**.
- S. FEDERICO, *Industry dynamics and competition from low-wage countries: evidence on Italy*, Oxford Bulletin of Economics and Statistics, **TD No. 877 (settembre 2012)**.