

Applying Logistic Regression Models on Business Cycle Prediction

Ulf Hamberg[⊗] and David Verständig[✉]

The prediction of business cycles in the real economy is an important but challenging task. For this purpose we construct a logistic regression model that we use to predict turning points in the business cycles. With this approach we can determine probabilities of the economy being in a recession or an expansion in each time period. These probabilities are derived from a number of variables that should have an influence on the real economy. Our estimated model can predict business cycle turning points with a reasonable accuracy and could therefore be used as one potential tool for economic policy decision makers.

Key Words: logistic regression models, business cycle prediction, turning points

Tutor: David Domeij

Examinator: Richard Friberg

Discussant: Karin Hederörs Eriksson (20163), Anna Sandberg (20118)

Presentation: 25 March 2009, Stockholm School of Economics, Room 343

We would like to thank our tutor David Domeij for all his help and advice as well as Per-Olov Edlund for his valuable input on econometrical models.

☺ 20124@student.hhs.se

✉ 20194@student.hhs.se

Table of Contents

1. Introduction.....	1
2. Literature Overview	3
3. The Logistic Regression	5
4. The Variables	7
4.1 Periodicity	7
4.2 The Dependent Variable	7
4.3 The Independent Variables	9
5. Results.....	14
5.1 Main Model with Quarterly Data.....	14
5.2 Complementary Model with Monthly Data	17
5.3 Sensitivity Analysis	18
5.4 Out-of-Sample Evaluation	21
6. Analysis	22
7. Conclusion	25
7.1 Discussion.....	25
8. References.....	28
Appendices.....	30

1. Introduction

Since the dawn of mankind people have been fascinated by the future and tried to figure out what lies ahead. Everyone, from the farmer trying to predict the weather to the king who wants to know where his enemies will attack have, with different means, tried to find out what will happen in a near or distant future. In economics, forecasting has always been an area of great interest which has drawn a lot of attention. With the introduction of computers and modern technology this interest has exploded and forecasting in a vast array of areas has seen a rapid development.

Being one of the most central economic variables, GDP has been the subject of significant efforts when it comes to developing new methods of forecasting. Being able to predict the movements of the business cycle is of great importance for almost everybody in society. Households want to know how salaries and employment rates will develop, companies are interested in whether the demand of their goods will shrink or grow and governments need to know which way the cycle is turning when making policy decisions. For policymakers, both government officials and central bankers, it is particularly important to have a correct picture of the current state of the business cycle. An incorrect assessment of the economic conditions might lead to policy measures which in best case are ineffective and in worst case directly harmful. Of especially high interest for these agents is the ability to identify when the business cycle shifts from a period of expansion to one of contraction and vice versa, since these phases demand entirely different policy actions. The identification of these turning points is therefore of particular interest to economists.

An approach to this problem that has been gaining support lately is the estimation of logistic regression models and the closely related probit model, in which the probability that the economy is in either a contraction or an expansion is calculated. Several papers using this method have been written in the U.S. using American data. However the leading variables and the optimal model for the U.S. economy might differ quite substantially from the variables that are important for small open economies. Research done on data from such economies might therefore be of great interest and deepen the understanding of these models.

The aim of this thesis is to contribute to the research on business cycle forecasting in small open economies and develop a logistic regression model for this purpose that is adjusted to Swedish conditions, which is something that to the best of our knowledge has not been done before.

In more detail this means:

- Identification of suitable leading variables
- Setting up a logistic regression model using different lags and periodicity of these variables
- Evaluating this model on the basis of different tests and parameters as well as the comparison with previous work done on this topic in order to reach a conclusion on the usefulness of such a model for forecasting purposes in small open economies such as Sweden

We come up with a model that we consider to have a reasonable predictive accuracy when evaluated both in- and out-of-sample. We also discuss the sensitivity of the results when changing periodicity and the dependent variable and we conclude that a model based on quarterly data is fairly robust. However, as with all models, we believe that ours should be used sensibly and merely as one tool in the forecasting toolbox.

The paper is organized as follows: in section two we briefly give an account of previous work on business cycle prediction in general and the use of probit and logistic regression models within this framework in particular. In section three we present the rationale and theory of logistic regression and in section four we discuss our model input, the dependent and the independent variables. In section five we present the results of our empirical study and these results are analysed more in detail in section six. Section seven concludes and considers the practical use of our model.

2. Literature Overview

Due to its practical importance, the prediction of business cycles has been a subject of interest for economists since the 1950's and there has been extensive research done on the estimation and evaluation of models with this purpose.

One of the earliest models, the Leading Indicator model, is widely used and has an advocate in the famous business cycle researcher, Geoffrey Moore. The idea is to create an index, the so called Composite Leading Indicator (CLI), which is essentially an indexation of different time series that correlate with the business cycle. This method has received quite a lot of attention and has proved to be a rather reliable predictor especially when the CLI increases or decreases for several consecutive months (Zarnowitz and Moore 1982).

Filardo (2004) gives a short overview of the development of four different business cycle prediction models: the Leading Indicators Model, Neftçi's Sequential Probability Model, the Probit Model and Stock and Watson's Experimental Recession Indexes. He evaluates how well these models were able to predict the turning point in the U.S. in 2001. Filardo comes to the conclusion that the Leading Indicators and Neftçi's models are the most useful ones, whereas Stock and Watson's is less powerful. The probit model is regarded to be reasonably accurate but its reliability on real time data is somewhat questioned.

However, the use of the probit model is strongly advocated by Estrella and Mishkin (1998) as the most logical and simple method to use when predicting turning points in business cycles. They create a model out of a wide range of variables such as interest rates and spreads, stock prices, money aggregates and macro indicators and use them to predict the US business cycle. They find that especially the yield spread is a reliable predictor of real activity but also several other variables such as stock prices and leading macroeconomic indicators can act as significant predictors. They also compare the probit model with other methods and show that the simple probit model in many aspects outperforms different complex CLI-methods.

Chin, Geweke and Miller (2000) develop a probit model to estimate the probabilities of turning points in the business cycle. When evaluating this model both in-sample and out-of-sample, they find that this type of model has some distinctive advantages to other standard economic models mostly because the probit model focuses directly on the turning points and is unaffected by the unpredictability of business cycle swings and the nonstationarity in these turning points.

A similar study with quarterly data is provided by Gaudreault and Lamy (2001). They use 35 different indicators to predict business cycles up to three quarters ahead. They evaluate four different probit models and show that there are significant difficulties associated with the prediction of turning points in the economy. They get their best result for a prediction zero quarters ahead, i.e. when predicting the current economic situation and this result is consistent for all four models. The single best model overall is the one where the current coincident index is combined with a two quarter lag of the change in Fed funds rate.

A more recent discussion paper using the probit model is Haltmaier (2008) who estimates business cycle predicting models for seven countries; the US, Canada, Japan, UK, Germany, Mexico, Korea and Taiwan. She uses monthly data for the variables oil price, stock price, real and nominal exchange rate, real activity, a leading indicator and the yield spread with up to six lags. Her results indicate that especially the yield spread and the share price are well performing predictors for almost every country, whereas the coefficients of the exchange rates and the oil price are insignificant in most cases. She concludes that probit models are useful as general guidance, especially for governments when deciding on economic policies.

3. The Logistic Regression

As seen in the literature review, there are several different methods that can be used when forecasting business cycles. In this thesis we use a model closely related to the binary probit model in which the probability of an event occurring is calculated with the help of different predictors. When choosing this model we take into account the findings of Chin, Geweke and Miller (2000) and Estrella and Mishkin (1998) who state that a probit model is the most appropriate model to use when predicting business cycle turning points. According to Chin, Geweke and Miller the probit model has a clear advantage over the usual standard methods in that it predicts turning points directly instead of indirectly through the estimation of future GDP, generally resulting in a higher degree of accuracy.

The model actually used in this thesis is, as mentioned above, not the probit model but a closely related version called the logistic-regression model. This model has the same basic implications as the probit but with the advantage of being somewhat more intuitive and easier to compute (Gujarati 2003, p. 614). More specifically we use what is called a binary logistic regression where binary means that the dependent variable can take only one of two values, “0” or “1”. In our case, a time period classified as an expansion is labelled “0” and a contraction is labelled “1”.

The results of the logistic regression are however not as easy to interpret as the usual linear OLS-regression. When computing a logistic regression, a parameter estimate comparable to the Beta-coefficient in an OLS-regression is obtained. This coefficient is called a logistic, or a logit, coefficient and it can be interpreted as the estimated change in the logit of a one unit increase in the explanatory variable. The logit is another expression for the log odds that an event is occurring or, in other words, the dependent variable taking the value “1”. The log odds is the natural log of the odds ratio which is the probability (P) that the event will occur in a given time period (i) to the probability that the event will not occur or in mathematical terms:

$$\ln\left(\frac{P_i}{1-P_i}\right) = Z_i \quad \text{where } Z_i \text{ is the logit.}$$

The probability of an event occurring (“1”) in a given period of time can be calculated by plugging in the observed values of the explanatory variables (X) together with the estimated logistic coefficients (β) in the formula, $\beta_1 + \beta_2 X_i = Z_i$, thus giving us the logit (Z_i). Expressing the logit to base e then gives us the odds ratio ($\frac{P_i}{1-P_i} = e^{Z_i}$) for the event happening.

Having the odds ratio and rearranging, one obtains the probability of the event occurring:

$$\frac{e^{z_i}}{1 + e^{z_i}} = P_i$$

A high positive logistic coefficient (or a high value of X_i) thus makes the logit larger and thereby increases the probability of an event occurring in a certain time period.

The probability obtained above can then be used in forecasting purposes to predict for example turning points in business cycles. If the probability rises above a given threshold, the model forecasts that the economy will fall into recession and thus has reached a turning point (peak). The next turning point (trough) then occurs when the probability again slides below the threshold. The predictive accuracy of the model can thus be easily evaluated by comparing the predicted recessions/expansions with the actual outcome in the economy. In this thesis we have used a threshold of 50 percent but it is worth mentioning that the threshold can be adjusted depending on how you weigh the risk of conducting a Type 1 error (not predicting a recession that actually occurs) against conducting a Type 2 error (predicting a recession that does not occur). A high threshold means that the risk for Type 2 errors is low but correspondingly means a high risk for Type 1 errors.

To come up with the best possible logistic regression model we use a method called stepwise backward elimination which in a predictive model as ours is a straightforward way of reaching the highest possible predictive accuracy (Menard 1995). In stepwise backward elimination one begins with a logistic regression model that includes all variables that are regarded to be of interest. Then, through repeated iterations, variables which are not significant estimates, are eliminated with one variable being eliminated in each iteration. By doing this one eventually ends up with a model only including significant variables, which means that for each variable the null hypothesis that its coefficient is zero can be rejected. In our thesis we choose to use a 10 percent significance level as the cut off level for removal of variables as suggested by Menard.

4. The Variables

4.1 Periodicity

When setting up a predictive model the periodicity of the data can have a significant effect on the results. Common periodicities are for example “daily”, “monthly”, “quarterly” or “yearly” data. When estimating GDP movements, quarterly data is the most common and for all purposes most useful periodicity. This is partly due to the availability of this kind of data, but quarterly figures are also preferable since the time series become less variable, thus reducing noise and usually leading to a better goodness of fit in the model (Estrella and Mishkin 1998). Following this and the lack of monthly Swedish data for some variables we choose to build our main model using a quarterly periodicity. However, we also make a complementary model on a monthly basis as suggested in Haltmaier (2008), which allows for a more timely estimation and a much larger sample of observations. Having a complementary model also allows us to make comparisons with our main findings regarding predictive accuracy and goodness of fit. However, as mentioned above, data for some variables is not available on a monthly basis. This means that for some variables each quarter has to be transformed into three corresponding months, thereby assuming that each monthly value is a third of the corresponding quarterly value.

4.2 The Dependent Variable

The following discussion concerning the definition of a business cycle is in accordance with the views expressed by the Swedish National Institute of Economic Research in Konjunkturläget (2005). They conclude that business cycles can be defined in different ways, but that there is a consensus among economists today, that the most useful measure is the difference between economic activity in the economy as measured by GDP and the potential level of activity. The potential level is the level of GDP that the economy can reach with full resource utilization.

The point at which actual GDP is furthest above potential GDP is called a peak. When the peak has been reached, a period of slower or negative growth begins which ends in the point when actual GDP is furthest below potential, called a trough. During the period between a peak and a trough the economy is said to be in a recession or a contraction. A recession is followed by a period of rising growth called an expansion which culminates in the next peak. A business cycle is defined as the whole period between two peaks or two troughs, including one contraction and one expansion. The length of the cycle varies wildly but is considered to last around three to eight years.

In order to develop a predictive model, a numerical measure of the business cycle is needed as the dependent variable. Using the definitions above, the situation in the economy can be summarized in the output gap. The output gap is calculated as actual GDP minus potential GDP, a positive output gap meaning that the economy is in a boom and a negative gap indicating an underutilization of resources. The difficulty when calculating the output gap lies in the calculation of potential GDP, which is affected by many different factors, among them labour force participation and labour productivity. To obtain an accurate number on potential GDP, economists use a variety of different indicators such as capacity utilization in the industry, job vacancies and the development of salaries. However, even with the use of the most advanced methods the exact size of the output gap is fairly difficult to measure and the results tend to vary between calculations done by different agencies and organizations. However, for the purpose of this thesis, in which we are more interested in predicting turning points than the absolute level of resource utilization, we consider the output gap to be a useful and viable variable.

In order to use a logistic regression model when analyzing business cycles one needs to classify each time period as either being a part of an expansion or a contraction. To be able to define each period as either an expansion or a contraction we therefore need data on the occurrence of peaks and troughs. For Sweden, this data is not available for the full time period in any major database and we therefore develop our own definition of what can be considered as troughs and peaks based on a method used in similar work done on unemployment by Chin, Geweke and Miller (2000).

The data used is the output-gap as a share of potential GDP from the Swedish National Institute of Economic Research on a quarterly basis beginning with the first quarter of 1985. This can be seen graphically in *Graph A1*. In order to be considered a possible peak, the output-gap of a specific quarter has to be “higher” than both the previous and the following quarter. Furthermore, since a business cycle lasts for three to eight years, no peak can be closer to another peak than three years. When this is the case the “lowest” peak is excluded and classified as part of a continuing expansion. By using this method we arrive at a quarter which has a “higher” value of output-gap than the adjacent quarters and is at least three years away from the nearest peak.

Also, following our definition, the next turning point after a peak has to be a trough which means that if two peaks follow after one another the “lowest” is excluded. Using these steps and the same reasoning for troughs we identify seven turning points; four peaks and three troughs. With these findings we classify all time periods between 1985-Q1 and 2008-Q2 as being either in an expansion or a contraction. All dates for peaks and troughs are found in *Table 1*. Important to note is that a peak, if it fulfils all the conditions

above, can occur even though GDP is below potential. In other words a peak can coincide with a negative output gap.

Table 1: Business Cycle Peaks and Troughs Since 1985

Peaks	1989-Q4	1995-Q2	2000-Q3	2006-Q3
Troughs	1993-Q4	1996-Q4	2003-Q2	

4.3 The Independent Variables

In order to develop a model that is not solely depending on one type of variables we try to incorporate a wide range of different kinds of independent variables. Such a broad model should be less vulnerable to sector specific volatility than one that has a very narrow focus and might therefore be more reliable over time.

As a starting point there should of course be a theoretically justified relationship between the leading independent variable and the dependent variable. However, even when it is possible to include a variable from a theoretical point of view, reliable data might not be readily available. Also, due to the risk of overfitting the model, it may not be desirable to include too many independent variables. The problem with overfitting is that the model gives good in-sample results but may be a lot worse at predicting beyond the data of the sample (Estrella and Mishkin 1998).

Taking these considerations into account six different explanatory variables are chosen; changes in a stock price index, oil price, a confidence indicator, European GDP and residential building starts as well as the spread between a short-term and a long-term interest rate. The chosen variables and the sources can be seen in *Table A1*. Similar to the approach by Haltmaier (2008) we use the yearly change in the variables to get a measure that reflects the relative development in the variables.

In general, financial variables such as prices of different financial instruments are associated with expectations of future economic events. Therefore, it seems necessary to include at least some financial variables in our model.

Intuitively, share prices should be an important predictor of the business cycle. In the long term the price of a share in a listed company is directly related to the future earnings, or rather the value of the discounted future dividends, of that company. If a slowdown in the economy is expected, the profits, and therefore the dividends, are normally expected to fall as well, which will lead to a decrease in share prices. In a small country like Sweden with relatively few listed companies it makes sense to use a share

price index that is as broad as possible to reduce the impact of firm-specific risk in accordance with normal portfolio theory (Bodie, Kane and Marcus 2005, p. 225). We therefore use the OMX Stockholm PI (abbreviated: *OMX*) which is the main Swedish stock index.

Another often used predictor of business cycles is the spread between the long-term and the short-term interest rate, the yield curve. More specifically, a positive yield curve is generally associated with a future increase in real economic activity and a downward sloping term structure is seen as negative sign for real economic growth (Estrella and Hardouvelis 1991). This relation is rather intuitive when considering the rationale behind the interest rate spread. The shape of the yield curve is related to the current supply and demand of short- and long-term bonds. Normally the demand for long-term bonds is lower than for short-term bonds for several reasons. The higher demand for short-term bonds depends at least partially of the higher risks of long-term bonds. Investing in long-term bonds means tying up money for a longer period of time as well as greater price uncertainty and higher default risk. Also, short-term bonds are usually more liquid which results in a liquidity premium for long-term bonds. This means that prices for short-term bonds are higher than for long-term bonds and, since the price of a bond is inversely related to the yield, the yields of short-term bonds are lower than for long-term bonds (Fabozzi 2007, pp. 94-122). This relation is distorted when a recessionary period is expected. In that case, holding long-term bonds is more attractive due to an increase in short-term uncertainty. This results in a decrease in the long-term yield and a flattening of the yield curve. Eventually, if the risk of recession is large enough, the yield curve is inverted, i.e. downward sloping, indicating a very negative short-term outlook on the economy. The spread variable (*Spread*) used in this thesis is defined as the difference between the ten-year Swedish government bond yield and the interest rate of three-month Swedish T-bills.

Prices of commodities that are important for households and industry can have a significant effect on the business cycle. Some of the most important commodities are those connected to energy. Higher energy prices will result in higher costs for individuals as well as companies, affecting disposable income and profit margins and is therefore likely to have a long-term effect on the economy. The main reason for this popular belief is the oil price shock in the 1970's that resulted in low growth as well as high inflation and unemployment (Blanchard and Gali 2007). However, Blanchard and Gali conclude that changes in oil prices have a larger impact when they coincide with other shocks in the economy, such as for example other commodity prices. They also find that the overall impact of oil price has decreased since the 1970's due to, for example, decreasing dependence on oil. However, there seems to be enough theoretical and empirical reasons to include the price of crude oil (*Oil*) as a variable in the model when predicting business cycles, especially since the results in the study of Blanchard and Gali are rather inconclusive when considering other countries than the US.

Another type of variables that should be considered for this kind of models is macroeconomic indicators. Macroeconomic variables have a proven track record when it comes to predicting business cycles (Estrella and Mishkin 1998).

One potentially useful macroeconomic variable is the confidence indicator of the Swedish manufacturing industry (*Confidence*) that is put together by the Swedish National Institute for Economic Research. This indicator index is based on monthly interviews with 4000-8500 companies where they comment on their current situation and their future expectations. Companies should be able to recognise business cycles early as a change in demand is likely to affect them with a very short delay.

Another possible macroeconomic variable is the GDP growth of countries outside Sweden. Being a small open economy like Sweden normally means a high dependence on trade with other countries. In the case of Sweden imports are as large as 45 percent and exports 55 percent of total Swedish GDP at market price in 2007 (Statistics Sweden).¹ Being a small open economy also makes it more likely that the Swedish business cycle follows the growth in the rest of the world rather than the other way around. In Sweden 75.8 percent of the value of exports and 85.1 percent of the value of imports comes from trade with other European countries (Statistics Sweden)² and it is therefore reasonable to assume that European growth is a sufficiently accurate proxy for Swedish dependence on the rest of the world. Another reason to use European growth is that the same factors that affect European growth also are likely to affect Sweden. Using world GDP growth would include countries with which Sweden has very little economical contact and countries that in other ways are very different to Sweden. We therefore use the total GDP for European countries as measured by OECD (*European GDP*).

The third variable in the macroeconomic category is the outlook for the Swedish building sector. One commonly used way to measure this is the number of residential buildings started during a specific time period. There are convincing empirical and theoretical reasons to believe that residential investment is an important predictor of the business cycle. The theoretical reason is quite intuitive; housing is often a highly leveraged investment which means that only a slight decline in asset prices will affect the whole housing market. Leamer (2007) considers measures of new housing permits or building starts to be the most reliable predictors for the business cycle. There is empirical research showing that eight recessions in the U.S. after World War II were preceded by “substantial problems in housing and consumer durables” (Leamer 2007, p.4).

¹ http://www.scb.se/templates/Product_22908.asp

² http://www.scb.se/templates/tableOrChart_51328.asp

When conducting a logistic regression, it is the sign of the coefficients obtained rather than the absolute value that is of interest. If one or several coefficients have a sign that contradict what is expected in theory, we can suspect that the model is mis-specified in some way. It is therefore important to have a clear picture of what sign to expect from the different variables. A negative sign means that a rise in this variable is associated with a lower probability of recession. Similarly, a positive sign implicates that a larger value of this variable means a higher probability for recession.

The following variables are expected to have a negative sign:

OMX - Increasing share prices should be associated with expectations of high dividends and hence strong economic activity.

Spread - In times of strong economic performance demand for long-term bonds, in relation to the demand for short-term bonds, decreases and there is therefore an increase in the long-term yield. A higher, positive, yield spread is thus associated with increased expectations of strong economic activity.

Confidence - Being a measure of the near future expectations of Swedish industrial companies a high value of this variable should be associated with strong economic activity.

Building - Following the cyclical demand and prices of real estate, expectations of higher economic activity leads to an increase in the number of newly started building projects.

Europe GDP - Improving economic conditions in the rest of Europe should affect Swedish economic activity positively due to the heavy dependence on exports.

The following variable is expected to have a positive sign:

Oil - In oil importing countries, such as Sweden, a higher oil price can be seen as a supply shock which affects the economic activity negatively.

Since the above mentioned variables are related to GDP in different ways and through different channels it is likely that the time horizons of their relationships are varying. We therefore use several lags of every variable in the regressions. When using quarterly data each variable is estimated with two different lags and for monthly data six lags are used, so that in both cases the lags cover a period of half a year.

In order to find out which time horizon should be used for the lags we conduct a simple test following the methodology used by Estrella and Mishkin (1998) where each independent variable is regressed against the dependent variable with three different lags; one month, six months and twelve months. Doing this, a pattern evolves showing us if the relationship is strongest with a shorter or a longer time horizon.

For all variables except oil price the test shows that a short time horizon gives the strongest relationship meaning that the use of lags one to six for monthly data and one to two for quarterly is to prefer. The argument for such a short time horizon is further strengthened by Haltmaier (2008) who states that the relationship between dependent and independent variable becomes less reliable with longer lags. However, for Oil price the test clearly shows that a somewhat longer time horizon means a stronger relationship and we therefore use monthly lags seven to twelve and quarterly lags three to four for this variable. Regarding European GDP we have to take into account that this data is not available immediately after each quarter but is released almost two months later. Hence, for the model to be practically useful we use lags three to eight in the monthly version and two to three in the quarterly version for this variable.

5. Results

5.1 Main Model with Quarterly Data

The result from running the logistic regression model for quarterly data can be seen in *Table 2* where the figures in brackets indicate the number of lags for each individual variable. After running a backward stepwise regression we get a model with six independent variables and one constant.

Table 2: Variables in the Equation with Quarterly Data

Variables	B	S.E.	Sig.
Confidence(1)	-0.048	0.027	0.080
Europe GDP(3)	-54.635	30.746	0.076
Oil(4)	2.988	1.225	0.015*
OMX(1)	-4.296	1.472	0.004*
Spread(1)	-114.972	41.945	0.006*
Constant	2.698	0.981	0.006*

* Significant at the 5 percent level

Obviously all coefficients are significant at the ten percent level and important to notice is also that all coefficients have the expected sign.

When building a model for forecasting and prediction purposes, an important consideration when evaluating the model is how well the predicted values match the actual observed values of the dependent variable during the sample period. This is called an in-sample evaluation and the results of it can be seen in *Table 3* and graphically in *Graph A2*.

Table 3: Classification Table, Quarterly Data

Observed	Predicted		
	Expansion	Recession	Percentage Correct
Expansion	41	9	82.0
Recession	10	30	75.0
Overall Percentage			78.9

In *Table 3*, the columns represent the number of predicted classifications of the dependent variable for all time periods in the sample and the rows are the actual observed classifications of the dependent variable from our sample. This means that we observe 50 cases of expansion and 40 cases of recession in our time series. Of the 50 observed expansions our model correctly predicts 41 and out of the 40 observed recessions the model correctly predicts 30. We therefore conclude that periods of expansion are predicted correctly in 82.0 percent of the cases and recessionary periods are predicted correctly in 75.0 percent of the cases. This implies an overall percentage of correctly predicted observations of 78.9 percent. This can be compared to the result of a model with only the intercept which succeeds at predicting only 55.6 percent correctly thus indicating that our explanatory variables add predictive power.

As mentioned above, the focus when evaluating a prediction model lies in the comparison of the predicted to the observed classifications of the dependent variable, i.e. the predictive efficiency. However, this measure does not take into account the accurateness of the probabilities used to reach the classification i.e. the fit of the model. A model could thus have a high proportion of correctly classified predictions but if the probabilities all lie around the threshold of for example 50 percent, the goodness-of-fit will be worse than for a model which produces probabilities close to “0” or “1”. Even though the predictive efficiency is the more important measure for our purposes it is still valuable to look at the goodness-of-fit of the model as a way to test its appropriateness. Two measures of the goodness-of-fit are summarized in *Table 4*.

Table 4: Goodness-of-Fit, Quarterly Data

- 2 Log Likelihood	Hosmer and Lemeshow Test		
	Chi-square	Df	Sig.
68.185	1.98	8	0.982

The -2 Log Likelihood statistics can be seen as a badness-of-fit statistics. We would therefore want the coefficient to be as low as possible to show a good fit between the model and the observed data. A rule of thumb is that a value below 100 indicates a good fit and a value below 20 a very good fit (University of Colorado Denver 2006); hence our obtained value of 68.185 indicates the former.

With the Hosmer and Lemeshow Goodness-of-Fit we test the null hypothesis that there is no difference between observed and model-predicted values. We cannot reject this null hypothesis since the level of significance in our model is 0.982 which is well above the critical 0.05 level. Hence, the Hosmer-Lemeshow test indicates a good fit for the model.

Overall, the different goodness-of-fit tests all imply that our model has a sufficiently good fit when using quarterly data.

5.2 Complementary Model with Monthly Data

As mentioned before we also perform a complementary regression using monthly data to assess whether the use of another periodicity would alter the results of the main model. After running the backward-stepwise elimination we get the variables in *Table 5*. As above the number in the brackets indicate the number of lags but this time in months.

Table5: Variables in the Equation with Monthly Data

Variables	B	S.E.	Sig.
Building(1)	-1.735	0.684	0.011*
Confidence(3)	-0.070	0.016	0.000*
Europe GDP(8)	-64.828	19.488	0.001*
Oil(12)	2.737	0.711	0.000*
OMX(1)	-4.289	0.824	0.000*
Spread(1)	-102.884	24.140	0.000*
Constant	3.011	0.658	0.000*

* Significant at the 5 percent level

Just as for the main model all coefficients are significant at the ten percent level and all of the coefficients have the expected sign.

The results of the in-sample evaluation can be seen in *Table 6* and graphically in *Graph A3*.

Table 6: Classification Table, Monthly Data

Observed	Predicted		
	Expansion	Recession	Percentage Correct
Expansion	129	21	86.0
Recession	26	94	78.3
Overall Percentage			82.6

It seems like the use of monthly data actually gives a higher percentage of cases that are correctly predicted than our main model with quarterly data.

The corresponding goodness of fit of the model using monthly data can be seen in *Table 7*.

Table 7: Goodness-of-Fit, Monthly Data

- 2 Log Likelihood	Hosmer and Lemeshow Test		
	Chi-square	Df	Sig.
186.986	11.582	8	0.171

Using the same reasoning as with quarterly data it appears as if the use of monthly data gives a considerably worse fit for the -2 Log Likelihood and the Hosmer and Lemeshow test.

5.3 Sensitivity Analysis

In order to check the robustness of a model a sensitivity analysis is often conducted. The aim of such an analysis is to see to what extent the results and implications of the model change when altering different factors and underlying assumptions.

Since the choice of peaks and troughs inevitably is dependent on the used classification method, we want to test how our results are affected by changes in the dependent variable. This is done by removing one of

our defined peaks, 1995-Q2, which, although being preceded by an expansion, occurred in a period with negative output gap. This removal means that a series of observations of the dependent variable that originally were classified as a contraction (1995-Q2 until 1996-Q4) now are considered to be a part of a longer expansion from 1993-Q4 until 2000-Q3, altering the values of the dependent variable. Besides that everything else is the same as in the models presented above.

The results from running the main model on the altered dependent variable can be seen in *Table 8*.

Table 8: Variables in the Equation with Quarterly Data

Variables	B	S.E.	Sig.
Europe GDP(3)	-95.045	38.766	0.014*
Oil(4)	5.344	1.864	0.004*
OMX(1)	-11.360	3.494	0.001*
OMX(2)	5.107	2.954	0.084
Spread(1)	-250.457	71.781	0.000*
Constant	4.455	1.380	0.001*

* Significant at the 5 percent level

Here one more lag of the OMX variable has been found significant and this lag, OMX(2) does not have the expected sign. Even though this result is not significant at the 5 percent level it could be a sign of warning that the model is mis-specified, i.e. that there either might be some independent variable missing or that the dependent variable is not correctly classified. Except for the added OMX variable the only difference between this new and the original main model is that the confidence variable is not significant for any lag and thus has been eliminated from the model.

The corresponding coefficients with monthly data are presented in *Table 9*.

Table 9: Variables in the Equation with Monthly Data

Variables	B	S.E.	Sig.
Confidence(2)	0.055	0.019	0.003*
Europe GDP(3)	-105.393	27.220	0.000*
Oil(10)	4.611	0.936	0.000*
OMX(1)	-7.586	1.383	0.000*
OMX(5)	-4.570	2.611	0.080
OMX(6)	4.774	2.436	0.050
Spread(1)	-213.563	54.535	0.000*
Spread(4)	-98.173	44.077	0.026*
Constant	5.473	0.961	0.000*

* Significant at the 5 percent level

A couple of things are worth noticing with the results from this new version of the monthly model. First of all, we have a greater number of variables than before. There are especially an increased number of OMX variables, which now occurs with three different significant lags. Secondly, the Building variable has disappeared completely and thirdly, among the coefficients there are two, OMX(6) and Confidence(2), which do not have the expected sign.

Overall, the sensitivity analysis shows that our original main model is fairly robust. Using a monthly periodicity on the other hand, causes the model to become more sensitive to changes in the dependent variable.

5.4 Out-of-Sample Evaluation

To test our model's ability to predict outside the sample, which is of course the ultimate purpose when forecasting, we use the same methodology as Haltmaier (2008) and re-estimate the chosen variables with a sample consisting only of observations up to 1999-Q4. This model with the re-estimated coefficients is then used to forecast the remaining period up until 2008-Q2. Since the re-estimation is done on the same variables and the same lags as in the full sample version it is not strictly speaking an out-of-sample prediction but it allows us compare the results with those obtained in the in-sample evaluation. It means however that some of the variables might not be significant when using the shorter sample. The new coefficients obtained through the re-estimation can be seen in *Tables A2* and *A3* in the Appendix.

The results of the out-of-sample prediction using both a quarterly and a monthly periodicity can be seen in *Table 10* and graphically in *Graph A4* and *Graph A5*.

Table 10: Out-of-Sample Prediction

Periodicity	% of total observations correctly categorized	% of recessionary periods correctly categorized in prediction	% of predicted recessionary periods being false alarms
Quarterly, from Q1 2000 – Q2 2008	55.9	55.6	41.2
Monthly, from January 2000 – June 2008	76.5	75.9	21.2

6. Analysis

The evaluation of the models used in this thesis can be divided into in-sample evaluation and out-of-sample evaluation. In the in-sample evaluation a time series of predictions obtained through the model is compared to the actual observations for the same period. In order to decide whether our results can be considered satisfactory or not, we compare them to those of Haltmaier (2008). We can show that the number of correctly categorized observations using our main model (78.9 percent) is slightly below those reached by Haltmaier whose results vary between 75.4 and 92.4 percent with a large majority lying above 80 percent. This could be attributed to the vast difference in size between our sample and the sample used by Haltmaier whose time series starts in the early 1970's. More observations should allow a more accurate regression thus resulting in a better model performance. One can therefore expect that a re-modelling with a longer series of Swedish data should improve the in-sample performance to a level comparable to that obtained by Haltmaier. This hypothesis is further strengthened by the fact that our complementary model based on monthly data, with a sample three times as large as the one for the quarterly version, actually succeeds to predict 82.6 percent of the observations correctly.

The importance of the sample size becomes even more apparent when looking at the out-of-sample forecast. Here the quarterly results can be regarded as rather weak, compared both to the in-sample results and the results obtained by Haltmaier (2008) who gets at least 69.0 percent of total observations correctly categorized. The results of using monthly data are somewhat better but still slightly inferior to the in-sample results. Looking at *Graphs A4-5* in the Appendix it seems like the biggest problem with our out-of-sample prediction is that we miss the recession starting in the third quarter of 2006 by approximately a year.

We believe that the relatively poor out-of-sample performance of especially the quarterly model can be attributed to the reduction in sample size when using only observations up until year 2000. The deterioration in accurateness that could be the result of using such a small sample size might be enough to worsen the predictive performance quite substantially. This view is supported by the fact that the monthly model with its larger number of observations performs reasonably well also out-of-sample.

From this perspective one could draw the conclusion that a monthly periodicity is to prefer to a quarterly. However, as indicated by our goodness-of-fit test, the use of a monthly periodicity means a worse fit in the model which could be attributed to the higher volatility in this type of data. It is easy to see how a model with a bad fit i.e. with the estimated probabilities for many periods lying close to the threshold, should be less robust than a model with a good fit. In a model with a poor fit only a small change in the

probability could alter the prediction in a given time period from an “expansion” to a “contraction” or vice versa. This could explain the results of the sensitivity analysis where the monthly model becomes rather distorted when one of our peaks is removed from the data set. Hence, the poor fit of the monthly model makes it less robust with respect to small changes in the sample thereby making it more sensitive to the classification and definition of variables. The quarterly model is also distorted in the sensitivity analysis but to a much lesser extent, which indicates that the quarterly model is less sensitive to changes in the dependent variable.

Looking at the specification of our main model, it is to a high degree in line with what we had expected, with a big exception being the exclusion of the Building variable. As can be seen in *Table 2*, all other variables are significant at the 10 percent level and, most importantly, they all have the expected sign. The Building variable not being significant could imply that the building of residential housing is not as cyclical as we had expected beforehand. Another reason could be that building a new house is a rather long-term decision that usually has to be planned several years in advance before the actual building begins. This would imply that building starts is a lagging variable in relation to the business cycle rather than a leading one. A better leading variable might be the change in the number of applied or granted building permits, since this application procedure is in an earlier stage of the building process. Not applying for more permits should be less costly than to abort the start of the actual building.

Using monthly data as shown in *Table 5*, confirms the results of the main model except that Building in this case actually is significant. Worth considering regarding the Building variable in the monthly model is that it is based on underlying quarterly data and this approximation might be slightly over-simplistic. One should thus be careful when drawing any conclusions on the viability of the Building variable from this.

In general the results of our main model are in line with the results in the previous papers described in the Literature review. The share price index and the yield curve are used as reliable variables in the models of most other papers and they also seem to be the most significant ones in our model. There are however some differences that are worth noticing between our work and previous research. Notably, the oil price variable is a more reliable predictor in our model than in Haltmaier (2008) where it is found insignificant for several countries. We believe that this is due to the fact that she uses only lags with a rather short time horizon in her model while our findings show that the oil price is more strongly related to GDP on a longer time horizon. The rationale behind our findings could be that the oil price, as opposed to some of the other variables, affects the Swedish business cycle directly instead of being merely a reflection of what society or the market thinks about the future as is the case with for example stock prices. It is then

plausible that variables that serve more as indicators should have shorter lags than the oil price which needs a longer time to have a direct impact on the economy.

Furthermore, none of the other authors discuss the influence of other countries on the domestic business cycle. This is mostly due to the fact that the majority of them evaluate models on the U.S. economy and not on a small open economy like we do. Our results show that the external influence of GDP developments in other countries, in our case the rest of Europe, is an important indicator of the business cycle developments in Sweden.

7. Conclusion

The research done in this thesis supports the view that a logistic regression model can be a helpful tool when forecasting turning points in the business cycle. The main model succeeds in predicting more than three quarters of all observations correctly and with the adding of more data we expect the predictive accuracy to become better. The selection of the wide variety of different variables seems successful with all except one being significant and having the correct sign in our main model. However, our work also shows that this kind of model is rather sensitive to which underlying data is used. Choices of periodicity, classification rules for the dependant variable and use of lags all have a significant impact on the results obtained and have to be thought through properly when setting up the model. The size of the sample seems to be especially important for the performance of a logistic regression model. In order to get a reasonably accurate model a sample size at least comparable to our full sample of quarterly data seems to be necessary.

We also conclude that the optimal choice of explanatory variables might differ between countries with different characteristics and environments. Adding more country specific variables, such as GDP for neighbouring countries, can improve the performance of the model and make it reflect local conditions better.

To sum up, a logistic regression model with quarterly data has several appealing features and can be helpful to policymakers and others when identifying early signs of business cycle turning points. As all forecasting models it is however far from perfect and its results should hence be considered more as guidance than an absolute truth.

7.1 Discussion

The purpose of a forecasting model is to help predict events in the future. However, all forecasting has inevitably to be based on experiences of the past and all evaluations of such a model have to be done looking back at events that already have occurred. What makes things problematic is the fact that we are surrounded by a great deal of uncertainty, thus making it impossible to build a model based on past experiences that predicts the future perfectly.

With this said, it should be clear that our models, just like all other models, can be seen merely as guidance tools and not surprisingly they are subject to several issues worth discussing.

One such concern is the fact that the last recession in our time series, which begins in the end of 2006, is predicted almost a year late both using in-sample and out-of-sample forecasting. This shows that even though we use a wide spectrum of different variables, sometimes the rationales behind a turning point might be so extraordinary and complex that a model like ours may not give up a warning until it is too late. A slightly more optimistic way of looking at it, is that the classification of the business cycle in these kinds of models does not take the magnitude of recessions and expansions into account. That the forecast misses out on some of the early periods in a recession might therefore be due to the downturn not being particularly strong in the beginning. In that case a seemingly late prediction might actually be of help for policymakers if it arrives before a more dramatic development in the cycle occurs.

The discussion above also highlights the importance of correctly classifying the dependant variable. This is not a problem when the binary nature of the subject measured is straightforward, e.g. you either own a car or you do not, but in the case of business cycles it becomes more complex. To decide whether or not you are in a contraction depends on how you define peaks and troughs. Even though we use a systematic approach similar to the one used by Chin, Geweke and Miller (2000) when doing this, it has to involve a certain degree of subjectivity. To minimize the subjectivity we have followed the definitions of the Swedish National Institute of Economic Research, which can be considered as the foremost authority on this kind of research in Sweden, as close as possible.

Another subject worth discussing is the time horizons of lags used. Per definition a model cannot predict any further into the future than its shortest lag. In this thesis the shortest lags used are one quarter and one month which means that our models cannot be used to predict events further ahead than that. It could then be debated whether or not a warning on such a short notice is of any practical use for policymakers. Worth noticing however, is that GDP figures are usually delayed a couple of months which means that a prediction indicating a turning point might actually arrive at least three months before the actual GDP data has been released. Such a time frame should at least make it possible for policymakers to prepare their actions and thereby shorten the response time once a turning point has actually been confirmed.

The fact that GDP figures and other macro variables are quite difficult and time-consuming to calculate accurately, even for agencies such as OECD, is however a problem when using such variables as independents in a forecasting model. The lack of accurateness in the calculations means that they are often updated several times which implies that the original figures plugged into a model might actually not be a true reflection of reality. Hence, there is a trade off between waiting for revised numbers, a process that could take several months, and producing useful forecasts that arrive in time for policy makers to take advantage of them. Since the last update of the figures might not arrive until a year after

the initial release, choosing longer lags is not really a viable option either, as this would significantly deteriorate the predictive power. We obviously recognise this problem but at the same time we want to emphasize that this is a general problem for most existing forecasting models when including these variables. Furthermore, even though the original numbers might not be exactly the same as the final version, this deviation is usually quite small and should not alter the results substantially.

Another issue related to the nature of the data is the availability of reliable time series. As mentioned in the analysis the size of the sample has a rather strong impact on the accurateness and performance of the model. Even though Sweden is a country in which economic statistics have been collected for a relatively long time, the lack of reliable and comparable data makes it difficult to use observations further back in time than the 1980's. Particularly, having only so many observations, limits the possibility to perform a useful out-of-sample evaluation since all shrinkage of the sample means a relatively large proportional decrease in the number of observations left to use. In the case of Sweden, sample size is thus mostly an issue when trying to evaluate the model through the use of out-of-sample forecasting and is not as problematic for practical use. However, in less developed countries where available time series might be even shorter than in Sweden, this might actually put a limit on the predictive performance of the model. One way to increase the numbers of observations is to transform underlying quarterly data to monthly observations as is done in this thesis for some variables in the monthly models. However, as mentioned above this method is not very refined and rests on the assumption that the value of each month is a third of that of the corresponding quarter. Since the data made available using this method is merely a rough approximation it is clear that the results derived from it should be considered carefully and only be used when absolutely necessary.

The issues mentioned above are meant to highlight the different difficulties associated with forecasting based on logistic regression models. This discussion should however not be seen as an argument for not using such models but is rather meant to emphasize the importance of not relying on just one model. For best possible performance, forecasting should rely on the use of a variety of different models and not least the judgement of the forecasters.

8. References

- Blanchard, O.J. and Gali, J. (2007), “The Macroeconomic Effects of Oil Shocks: Why are the 2000s so Different from the 1970s?” Working Paper 13368. National Bureau of Economic Research, Cambridge.
- Bloomberg.
- Bodie, Z., Kane, A. and Marcus, A.J. (2005), *Investments*, Boston: McGraw-Hill.
- Chin, D., Geweke, J. and Miller, P. (2000), “Predicting Turning Points.” Federal Reserve Bank of Minneapolis Research Department Staff Report 267.
- Datastream.
- Estrella, A. and Hardouvelis, G.A. (1991), “The Term Structure as a Predictor of Real Economic Activity.” *The Journal of Finance*, Vol. 46, No.2, 555-576.
- Estrella, A. and Mishkin, F. (1998), “Predicting U.S. Recessions: Financial Variables as Leading Indicators.” *The Review of Economics and Statistics*, Vol. 80, No. 1, 45-61.
- Fabozzi, F.J. (2007), *Bond Markets, Analysis, and Strategies*, Upper Saddle River: Pearson Prentice Hall.
- Filardo, A.J. (2004), “The 2001 US recession: what did recession prediction models tell us?” BIS Working Papers 148, Monetary and Economic Department of the Bank for International Settlements.
- Gaudreault, C. and Lamy, R. (2001), “Forecasting a One Quarter Decline in U.S Real GDP with Probit Models.” Working Paper 2002-05. Department of Finance, Economic and Fiscal Policy Branch.
- Gujarati, D.N. (2003), *Basic Econometrics*, Boston: McGraw-Hill.
- Haltmaier, J. (2008), “Predicting Cycles in Economic Activity.” International Finance Discussion Paper Number 926. Board of Governors of the Federal Reserve System.
- Leamer, E.E. (2007), “Housing is the Business Cycle.” Working Paper 13428. National Bureau of Economic Research, Cambridge.
- Kling, J.L. (1987), “Predicting the Turning Points of Business and Economic Time Series.” *The Journal of Business*, Vol. 60, No. 2, 201-238.
- Menard, S. (1995), *Applied Logistic Regression Analysis*, Thousand Oaks: Sage Publications.

NasdaqOMXNordic, <http://www.nasdaqomxnordic.com>

National Institute of Economic Research (Konjunkturinstitutet), <http://www.konj.se>

National Institute of Economic Research (2005), *Konjunkturläget*, Available [online]: <http://www.konj.se/download/18.328c401048bb6f3da80004044/Konjunkturl%C3%A4get+10ny.pdf> [2008-12-16].

Riksbank, <http://www.riksbank.se>

Statistics Sweden (SCB), <http://www.scb.se>

University of Colorado Denver (2006), *HowToReg*. Available [online]: <http://www.nursing.ucdenver.edu/pdf3/LogRegHowTo.pdf> [2008-12-09].

Zarnowitz, V. and Moore, G.H. (1982), "Sequential Signals of Recession and Recovery." *The Journal of Business*, Vol. 55, No. 1, 57-85.

Appendices

Table A1: Variables and Sources

Variable	Name	Source
Stock Price Index	OMX Stockholm PI	OMX (http://www.nasdaqomxnordic.com/index/)
Oil Price	Brent Crude Spot Rate	Bloomberg (Ticker: USCRWTIC Comdty)
Long-Term Interest Rate	SWEDEN BOND YIELD GOVT.10 YR(ECON) - MIDDLE RATE	Datastream (Code: SWEGLTB)
Short-Term Interest Rate	SSVX 3M	The Riksbank (http://www.riksbank.se/templates/stat.aspx?id=16739)
Confidence Indicator	Confidence Indicator of the Swedish Industry	National Institute of Economic Research (http://www.konj.se/statistik/konjunkturbarmetern/foretag/tillverkningsindustrin.4.4756f14e114cfa17dd780001419.html)
Building Industry	Started Residential Building, Volume	Datastream (Code: SDHOUSE.P)
External Influence	All Europe GDP	Datastream (Code: EUOCFGDPG)

Table A2: Variables from Out-of-Sample Evaluation with Quarterly Data Q1 2000 - Q2 2008

Variables	B	S.E.	Sig.
Confidence(1)	-0.401	0.226	0.076
Europe GDP(3)	-525.346	275.843	0.057
Oil(4)	25.277	13.751	0.066
OMX(1)	-12.057	12.208	0.323
Spread(1)	-138.960	73.732	0.059
Constant	14.228	7.741	0.066

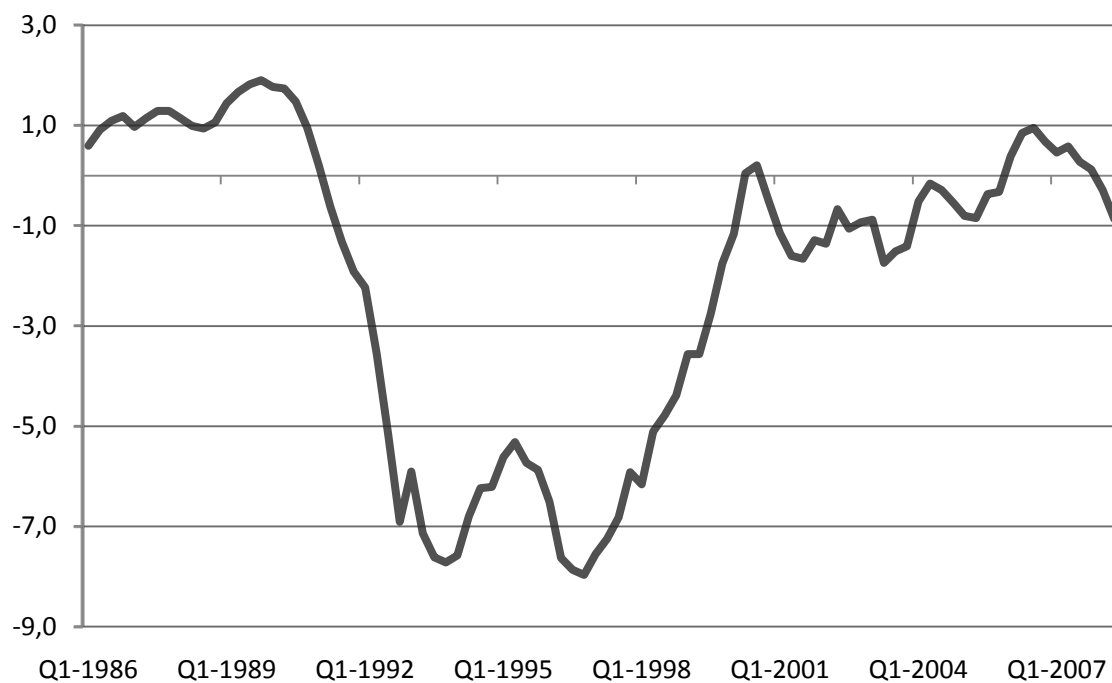
* Significant at the 5 percent level

Table A3: Variables from Out-of-Sample Evaluation with Monthly Data Jan. 2000 – Jul. 2008

Variables	B	S.E.	Sig.
Building(1)	-0.458	1.228	0.710
Confidence(3)	-0.304	0.073	0.000*
Europe GDP(8)	-402.485	99.993	0.000*
OMX(1)	-4.115	2.322	0.076
Oil(12)	14.252	3.528	0.000*
Spread(1)	-160.195	46.902	0.001*
Constant	10.901	2.784	0.000*

* Significant at the 5 percent level

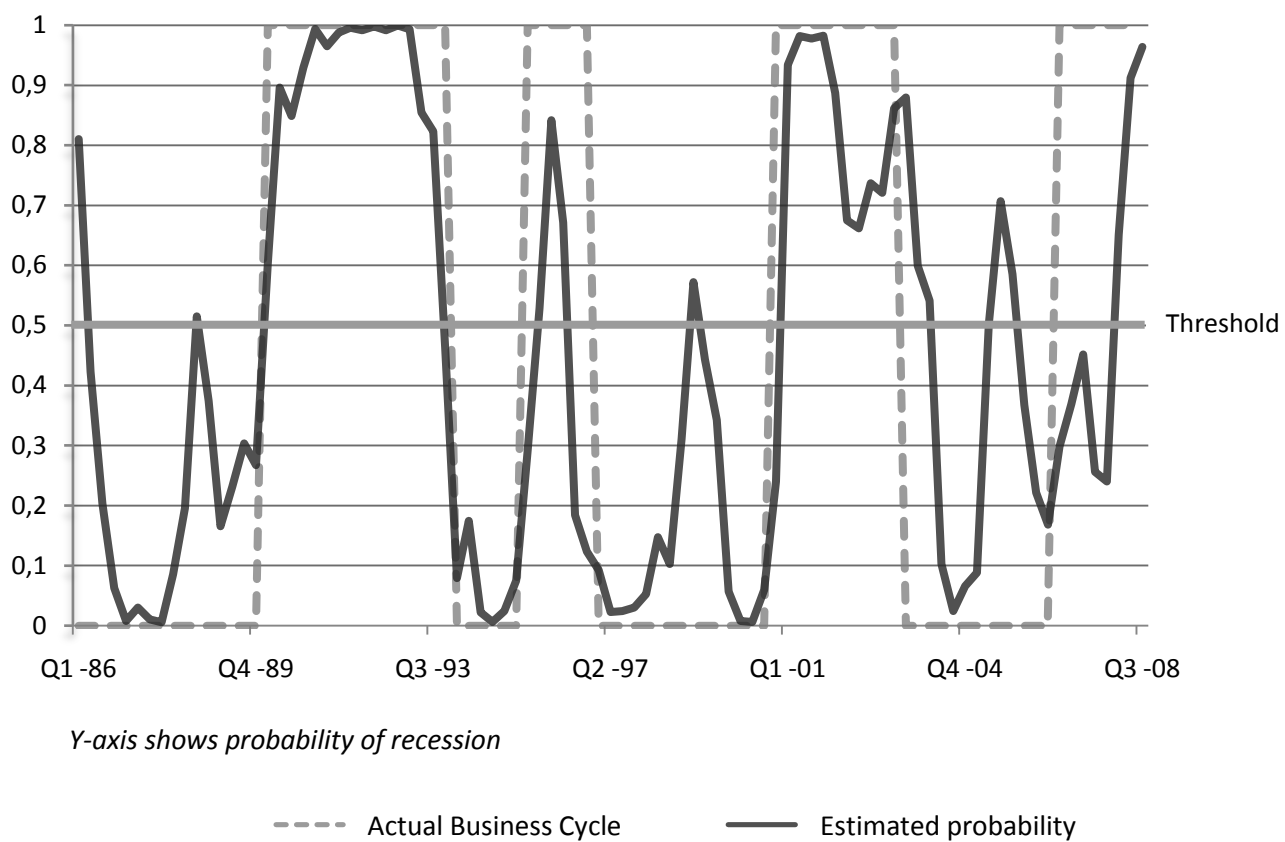
Graph A1: GDP Gap, in Percentage of Potential GDP



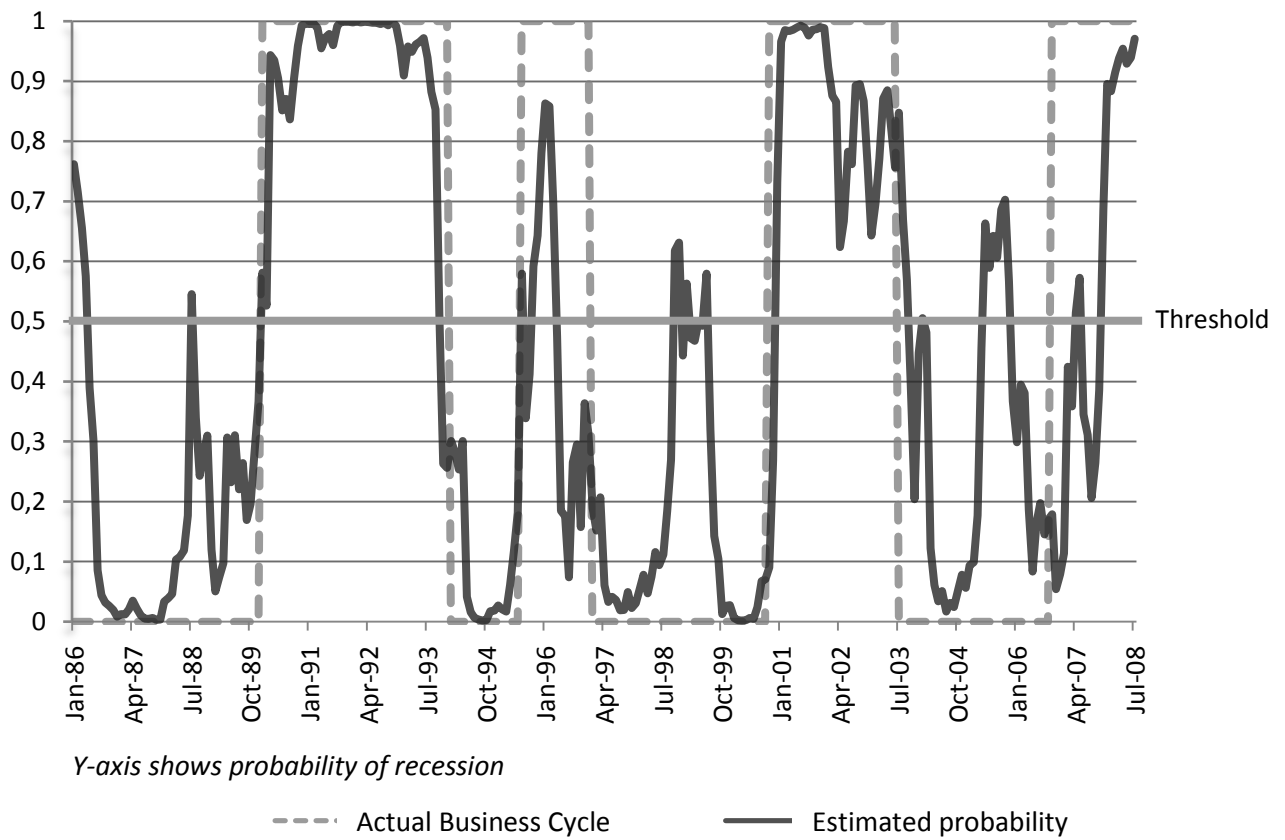
Y-axis shows percentage of Potential GDP

— GDP Gap, in Percentage of Potential GDP

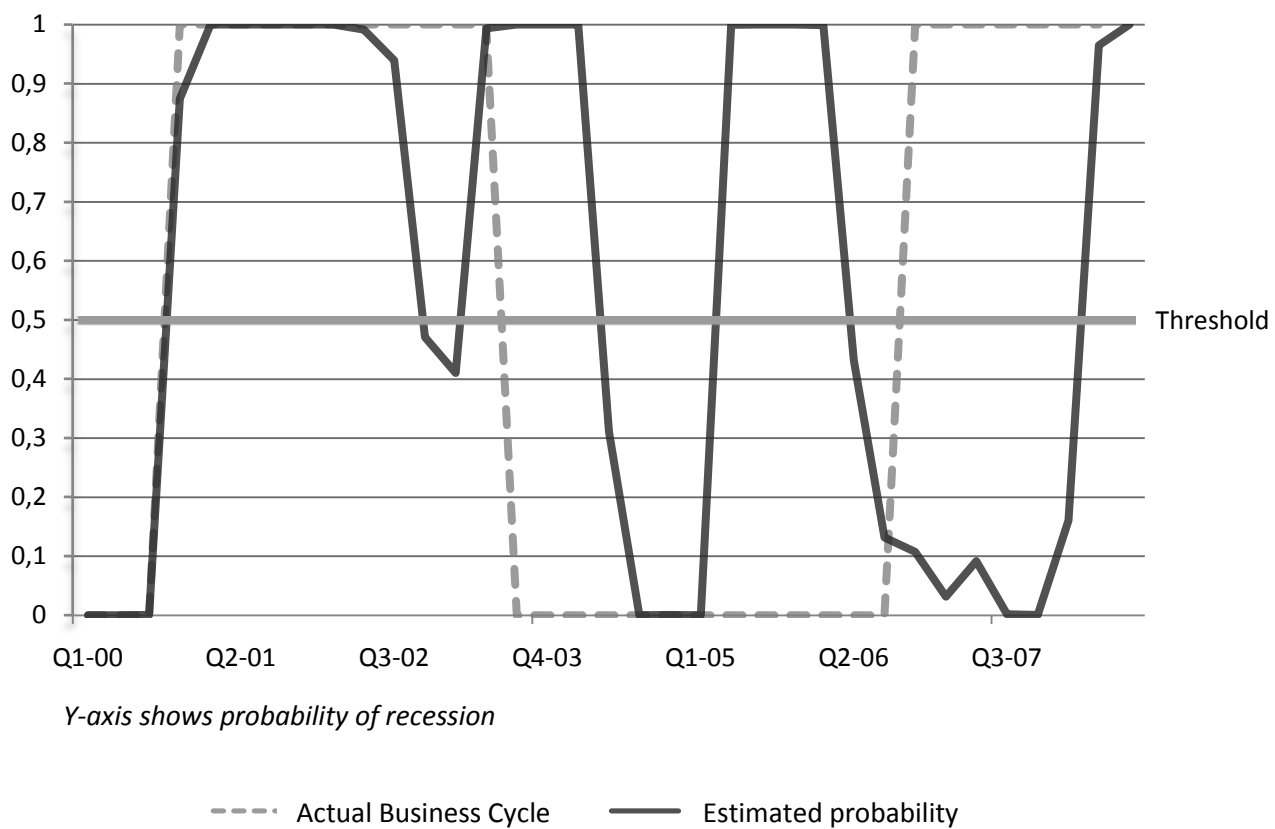
Graph A2: In-Sample Quarterly Data



Graph A3: In-Sample Monthly Data



Graph A4: Out-of-Sample Quarterly Data



Graph A5: Out-of-Sample Monthly Data

