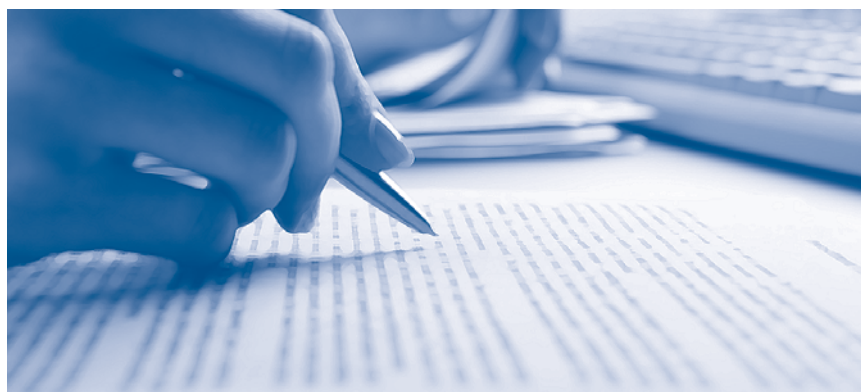


Banks, Asset Management or Consultancies' Inflation Forecasts: is there a better forecaster out there?

Tito Nícias Teixeira da Silva Filho

July, 2013

Working Papers



310

Working Paper Series	Brasília	n. 310	July	2013	p. 1-29
----------------------	----------	--------	------	------	---------

Working Paper Series

Edited by Research Department (Depep) – E-mail: workingpaper@bcb.gov.br

Editor: Benjamin Miranda Tabak – E-mail: benjamin.tabak@bcb.gov.br

Editorial Assistant: Jane Sofia Moita – E-mail: jane.sofia@bcb.gov.br

Head of Research Department: Eduardo José Araújo Lima – E-mail: eduardo.lima@bcb.gov.br

The Banco Central do Brasil Working Papers are all evaluated in double blind referee process.

Reproduction is permitted only if source is stated as follows: Working Paper n. 310.

Authorized by Carlos Hamilton Vasconcelos Araújo, Deputy Governor for Economic Policy.

General Control of Publications

Banco Central do Brasil

Comun/Dipiv/Coivi

SBS – Quadra 3 – Bloco B – Edifício-Sede – 14º andar

Caixa Postal 8.670

70074-900 Brasília – DF – Brazil

Phones: +55 (61) 3414-3710 and 3414-3565

Fax: +55 (61) 3414-1898

E-mail: editor@bcb.gov.br

The views expressed in this work are those of the authors and do not necessarily reflect those of the Banco Central or its members.

Although these Working Papers often represent preliminary work, citation of source is required when used or reproduced.

As opiniões expressas neste trabalho são exclusivamente do(s) autor(es) e não refletem, necessariamente, a visão do Banco Central do Brasil.

Ainda que este artigo represente trabalho preliminar, é requerida a citação da fonte, mesmo quando reproduzido parcialmente.

Citizen Service Division

Banco Central do Brasil

Deati/Diate

SBS – Quadra 3 – Bloco B – Edifício-Sede – 2º subsolo

70074-900 Brasília – DF – Brazil

Toll Free: 0800 9792345

Fax: +55 (61) 3414-2553

Internet: <<http://www.bcb.gov.br/?CONTACTUS>>

Banks, Asset Management or Consultancies' Inflation Forecasts: is there a better forecaster out there?*

Tito Nícias Teixeira da Silva Filho^{**}

The Working Papers should not be reported as representing the views of the Banco Central do Brasil. The views expressed in the papers are those of the author(s) and do not necessarily reflect those of the Banco Central do Brasil.

Abstract

The Focus Survey is a cunningly designed economic survey carried out by the Central Bank of Brazil. However, along its existence there have occasionally been some criticisms that its median forecast is “biased” due to allegedly strategic reasons. Although different groups of agents do have different strategic behaviours, one cannot take for granted that some groups predict differently from others. This paper tests if there are statistically significant differences in forecast accuracy between different groups of participants in the Focus Survey. Evidence shows some differences in forecasting ability among groups, but overall financial system forecasters have similar forecasting accuracy than consultancies.

Keywords: Survey of expectations, best forecaster, nonparametric tests, Friedman test and inflation forecasts.

JEL Classification: C12, C14, E37

* The author would like to thank Fabio Araujo, João Barata and Arnildo Correa as well as seminar participants at the Central Bank of Brazil for helpful comments.

** Central Bank of Brazil. E-mail: tito.nicias@bcb.gov.br.

“Currently (the Focus Survey) results have only **one vision**, which is normally more pessimistic than reality. That ends up influencing increases in the interest rate”

(Paulo Skaf, President of FIESP)¹

1 – Introduction

Since the onset of the rational expectations revolution agents’ expectations have begun to play a central role in macroeconomics. As a consequence, assessments of whether *actual* expectations are rational have become a key issue in empirical macro. Indeed, since then a massive number of papers have been written aiming at checking if actual expectations conform to the rational expectations paradigm.

The relevance of expectations has been boosted even further as of late 1980s, since when a large number of countries have decided to adopt inflation targeting frameworks. Ever since, not only the literature on the rationality of expectations has grown at an amazing pace but inflation expectation figures have left the pages of academic papers and have definitely entered into the daily economic debate.

A focal point of that literature has been to test whether inflation expectations are unbiased, although other requirements are also necessary to label a given series of expectations as rational. However, regardless of whether or not they are graded as rational, once one considers the existence and relevance of heterogeneous agents a myriad of interesting issues arise. For example, the degree of disagreement among forecasters can provide policymakers with valuable information about the inflation outlook [see Zarnowitz and Lambros (1987) and Lahiri *et al.* (1988)] and the overall state of the economy (Mankiw *et al.*, 2003).

Nonetheless, although surveys of expectations clearly show that heterogeneity is a key feature of agents’ expectations, some interesting issues have been overlooked by the literature. For example, are there some forecasters – or group of forecasters – that are consistently better than others? Or, do some groups of forecasters predict better in certain states of nature?

Such questions are particularly relevant given that several studies have shown that actual expectations bear “problematic” features [e.g. Nordhaus (1987), Bakhshi and

¹ Carta Capital (2011). Author’s translation, emphasis added. FIESP stands for São Paulo State’s Industry Federation. São Paulo is the most industrialized and rich state in the Brazilian federation.

Yates (1998), Thomas Jr. (1999) and Mankiw *et al.* (2003)]. Moreover – mainly in an attempt to *rationalize* those features – some have called to attention that agents' expectations can be affected by strategic behaviour and, therefore, that agents may not have the incentive to disclose their true expectations [e.g. Laster *et al.* (1999) and Pons-Novell (2003)]. Or yet, that agents might have different loss functions (e.g. Granger and Machina, 2006).²

A similar rationale has given rise to some criticism about the usefulness of the Focus Survey in Brazil [see Carta Capital (2011) and O Estado de São Paulo (2011)].³ The reasoning run as follows: a) the Survey is an “unbalanced” panel of forecasters since most of its participants come from the financial sector; b) financial sector participants have incentives to push upwards inflation expectations since higher interest rates are good for banking profits; c) thus, the median inflation expectation released by the Survey has an upward bias.⁴

In 2011, in order to enhance transparency, the Central Bank of Brazil has begun to report the median inflation expectation for three groups of respondents: banks, asset management firms and other institutions (see Central Bank of Brazil, 2011), so now agents have four median forecasts to choose from.

The aim of this paper is to test – using the rich database from the cunningly designed Focus Survey – whether there are some group of forecasters that predict better (or worse) than others. In order to test such a hypothesis forecasters from the following four groups of respondents were compared: commercial banks, investment banks, asset management firms and economic consultancy firms.

This paper presents two main heterogeneous features compared to others. First, it is the first paper to test such a hypothesis using the Focus Survey high quality data. Secondly, to my best knowledge the only other paper that have tested for differences among (more than two) *groups* of forecasters is that of Batchelor and Dua (1990).

However, the test here does not focus on possible accuracy differences due to forecasters' ideology or technique (as in Batchelor and Dua's paper) – which is difficult

² Strategic behaviour is certainly an issue in finance. For instance, forecasts from the sell side are certainly different from the buy side. However, it seems much less likely when one is predicting macroeconomic variables. Moreover, in latter case one can profit regardless of whether the variable being predicted goes up or won.

³ The Focus surveys agents' expectations regarding several economic variables (for a detailed account see Marques *et al.*, 2003).

⁴ Note that the mean would be much more affected by such reasoning than the median, which is the actual central tendency measure released.

(or even impossible) to define precisely since, for example, one can forecast using a combination of models based on different ideologies and techniques – but rather on the forecaster’s workplace. Not only this classification is easier (although not easy) to define, but it is more prone to uncover possible differences in forecast performance due to strategic behaviour, a commonly alleged reason to explain poor forecasting performance in practice. Indeed, Pons-Novell (2003) claims that forecasters “... behave differently, depending on their origin” and that those coming from nonfinancial sector tend to produce forecasts similar to the consensus.

Although the Focus Survey collects forecasts from several economic variables, CPI inflation forecasts are on the spotlight here, given their crucial role as a nominal anchor in the inflation targeting framework. Agents, therefore, have a strong incentive to get it right, as opposed to the GDP deflator, an indicator that central banks give less importance. Finally, differently from the GDP deflator – the inflation measure used by most of the studies listed in Section 2 – CPI inflation is not subjected to revisions, making inference more reliable.

The evidence presented here rejects the null hypothesis of equal forecasting accuracy in the following cases: Asset management firms produce better 6-month ahead forecasts than the other groups of agents: economic consultancy firms, commercial banks and investment banks. Moreover, they also seem to produce more accurate forecasts than investment banks for the 9 and 12-month ahead horizons. Therefore, overall the evidence suggests that if there is strategic behaviour either it is confined to a small number of agents or that its empirical relevance is of second order.

The paper is organised as follows: Section 2 carries out a literature review. Section 3 explains the data and the methodology used. Section 4 tests whether any of the above groups of agents predict better than others. Section 5 concludes the paper.

2 – Literature

Although the literature on expectation evaluation is enormous and heterogeneity has become an increasingly popular issue, there has been very little work assessing whether all forecasters are equally able or some predict better than others. When such (frequent) comparisons occur they involve either two types of agents or multiple *pairwise* assessments between individual forecasts and a (usually naive) forecast benchmark. One main focus of the first group of papers is to allow the possibility that

the central bank might have some informational advantage over the private sector, while the second strand of the literature aims at evaluating whether elaborated forecasts add value over simple forecasting methods.⁵

Stekler (1987) was one of the first to investigate the relative forecast accuracy between multiple forecasters. He investigated the performance of 24 professional U.S. respondents from the Blue Chip Survey, who had made predictions for nominal and real GDP and the GNP deflator, during the 1977–1982 period. He found evidence of different forecasting abilities between forecasters for all three variables.

However, Batchelor (1990) called to attention that Stekler had defined incorrectly the test statistic, and once the correct test was applied to Stekler's data it became insignificant in all cases, reversing his findings. Therefore, the null hypothesis that all forecasters had the same average forecasting ability could not be rejected.

Batchelor and Dua (1990) investigated whether forecasting accuracy differs according to the forecaster economic ideology and technique. They analysed 6, 12 and 18-month horizon forecasts for real growth, inflation (GNP deflator) and interest rate (3-month TBills) from 32 Blue Chip Survey participants, during the 1981.7–1986.12 period. They found scarce evidence of differences in forecast ability (i.e. only for 18-month growth forecasts). They assessed individual forecasters' performance as well and found statistically significant differences for inflation forecasts (in all horizons).⁶

Kolb and Stekler (1996) analysed interest rates forecasts expected to prevail 6 months ahead made by an unbalanced panel of 40 forecasters, whose predictions were published in the Wall Street Journal during the 1982–1990 period. They found no differences in their forecast ability regarding predictions of short-term interest rates (90-day T-Bill). However, they found significant differences in the average ability to predict long-term interest rates (30-year Treasury Bonds).

Ashiya (2010) assessed the performance of an unbalanced panel of 42 Japanese CPI forecasters during the 2004.4–2008.8 period. He analysed 0, 1, 2 and 5-month

⁵ For comparisons between central bank and the private sector forecasts see, for example, Romer and Romer (2000), Joutz and Stekler (2000) and Gavin and Mandal (2001). Other papers compare central bank's forecasts with those from naive predictions (e.g. Gavin and Mandal, 2003) and econometric forecasts [e.g. Joutz and Stekler, (2000), Sims (2002) and Faust and Wright (2009)]. Others assess forecasts from multilateral institutions (e.g. Abreu, 2011). Another strand of the literature compares market forecasts with naive predictions (e.g.). In the above cases only pairwise comparisons are being carried out.

⁶ These conclusions stem from the non-parametric evidence. However, they presented both parametric and non-parametric evidence, even though they expressed their clear preference for the latter.

ahead CPI forecasts and found evidence of different forecast performance among forecasters.

More recently, D'Agostino *et al.* (2011) analysed the performance of unbalanced panels up to around 300 SPF respondents' forecasts for inflation (GNP/GDP deflator) and real output during the 1968–2009 period. Two horizons were analysed: current quarter and one year ahead. Although their conclusion is not quite as follows, the evidence presented in the paper do suggests heterogeneity between **groups of forecasters**. The evidence also revealed the existence of really bad forecasters.

2.1 – Some Brief Comments

Despite the scant literature available, when one carefully reads it some impressions emerge. First, and most importantly, not only the qualitative results seem to depend on whether or not one uses balanced or unbalanced panels, but as D'Agostino *et al.*'s results show, accuracy tests could be very sensitive to the cut-off point (i.e. how many respondents to include in the sample). This is worrisome since selection bias is a real concern here.

Moreover, unbalanced panels typically include agents that have made forecasts for just a few periods. For example, their paper include agents that have made as few as ten forecasts in a 30-year sample, while Kolb and Stekler's considered agents that have made as few as 5 forecasts out of 17. It is really a relevant issue to what extent these uncommitted (or irregular) forecasters truly add useful information rather than just noise when inference is being made.

Of course, other factors might also have influenced the results above such as apparently too small samples, like the one available to Ashiya.⁷ Indeed, an unexpected result from Ashiya's paper is the evidence of heterogeneity in horizons as short as 0 and 1 month ahead, when forecasting is supposed to be *relatively* easier. In this sense the heterogeneity found by Kolb and Stekler's makes more sense. On the other hand, the very long span considered by D'Agostino *et al.* bring other concerns, since the SPF survey underwent major changes (e.g. Croushore, 1993) during that period, a fact that

⁷ Moreover, the database included the very first months of the survey, when noisy information (e.g. reporting errors) is likely to be relevant.

could have improperly influenced the results. Indeed, the survey actually ended for a brief period in early 1990s.

Therefore, taking into account the apparent sensitivity of unbalanced panels results, and given the one decade of high quality data available from the Focus Survey, all efforts will be made to compile a balanced panel of forecasters as large as possible.

3 – Data and Methodology

The Focus Survey is a real time market expectations survey carried out by the Central Bank of Brazil in which participants have access to an internet link and can feed the survey database as they wish.⁸ For example, it is possible to update forecasts in a daily basis.⁹ Moreover, the Survey has cross checks procedures to avoid the input of (some types of) inconsistent data (e.g. monthly forecasts need to “add up” to corresponding annual figures) and to mitigate reporting errors.¹⁰

It has begun in May 1999, helping setting the terrain for the adoption of the inflation targeting framework in Brazil few months later (July 1999). With that aim the Survey was cunningly designed to be a flexible, precise and high frequency tool for gathering market expectations for a broad range of economic variables.¹¹

One of the distinctive features of the Focus Survey has been the increasing number of participants along time, even though participation frequency and turnover have been higher than desirable, making reliable inferences more difficult to get.

The Survey is open for contributions from a wide range of economic agents.¹² Those that wish to participate and meet the requirements (see Marques *et al.*, 2003) need only to get in touch with the central bank in order to get registered and obtain a login and a password. The BCB has made continuous efforts to increase the number and

⁸ The survey has become online as of November 2001. The system upon which the survey is carried out generates daily reports to the members of the monetary policy committee with the most updated information. The BCB publishes a report every Monday to the public.

⁹ There are some restrictions, however, in the way agents feed the database. For example, if agents do not confirm their forecasts in the system within a 30-day period since they last updated them, the forecasts are deleted from the database, even if they have actually kept them unchanged.

¹⁰ The very fact that agents do not have to fill a questionnaire, whose data would have to be later compiled to build/update the database, minimizes reporting errors.

¹¹ Currently, forecasts for the current and each one of the next 18 months, as well as those for the current and the next five years are investigated.

¹² More precisely: “In principle, any institution (commercial banks and other financial institutions, non financial firms, consultancies, class associations, universities, etc) can be a member of the Survey, being the only requirement to deliver regular and robust forecasts” (Marques *et al.*, 2003). Author’s translation.

types of respondents. However, as common to other surveys, most respondents come from the financial sector. In December 2011, for example, around 85% of participants came from the financial sector.

Forecasts are more useful to policymakers the earlier they are produced, so as to allow pre-emptive policies to be implemented. Estimates regarding the monetary policy transmission mechanism in Brazil have indicated that moves in the policy rate take between two to five quarters to work themselves through prices (Inflation Report 2009 and 2012). Therefore, the forecasting horizons investigated – 6, 9 and 12 month-ahead – are particularly useful to monetary policy. Note also that one-year ahead forecasts is the most popular forecast horizon analysed in the literature, so the qualitative results obtained here can be easily compared to other countries' evidence.

Agent's n inflation expectation at time t regarding i month-ahead inflation measured over the last k periods is indicated by

$$E_t^n \pi_{t+i}^k \quad (1)$$

where $\pi_t^k = p_t - p_{t-k}$, $p_t = \ln P_t$ and $n = 1, 2, \dots, N$ indexes each forecaster. Moreover, in order to avoid inference problems caused by overlapping observations $i = k \in \{6, 9 \text{ and } 12\}$.¹³ Hence, agent's n forecast error regarding its inflation forecast over, for example, the next 12 months is given by

$$e_{t,t+12}^n = E_t^n \pi_{t+12} - \pi_{t+12} \quad (2)$$

The sample used in the paper begins in January 2002 and ends in June 2012, which means that one is able to get ten, fourteen and twenty-one annual, nine-month and semi-annual forecasts, respectively.

The distributional properties of forecast errors are crucial when one wishes to evaluate forecast performance. Traditional tests usually assume that forecast errors are normally distributed, among other “well-behaved” assumptions. However, as Harvey and Newbold (2003) show, there is strong evidence of non-normality in forecast errors from many economic variables. Indeed, it is quite common that some of the assumptions behind forecast accuracy tests are not met in empirical studies. Moreover,

¹³ Therefore, one of the indexes can be dropped.

parametric tests can be severely affected by outliers. Finally, small samples also pose challenges for parametric testing.

Taking concerns such as these into consideration this paper uses a variation of the Friedman test (Friedman, 1937 and 1940) in order to test the hypothesis of equal forecast accuracy among groups of agents. The Friedman test is a non-parametric test based on rankings that is used to make multiple comparisons between non-normal and dependent data.

Accuracy rankings are built according to the size of the absolute individual forecast errors. For every period a score is given to each forecaster according to his place in the ranking. More precisely, the top forecaster (i.e. the one with the smallest error) is given the score one, the runner up gets a score of two and so on until the worst forecaster gets a score equal to N , the number of forecasters at that period. If a tie occurs each forecaster involved in the tie gets the average rank that would occur if there were no tie.

Beginning at the “micro” level, the rationale of the test is the following: the average rank ($\bar{r}_t = N^{-1} \sum_{n=1}^N r_{n,t}$) across N forecasters ($n = 1, 2, \dots, N$) in each time period ($t = 1, 2, \dots, T$) equals $0.5(N + 1)$. If ranks were randomly given (i.e. all forecasters had equal average forecasting ability) then the rank sum over t for each forecaster ($R_n = \sum_{t=1}^T r_{n,t}$) would have the same expected value $[0.5T(N + 1)]$.¹⁴ Therefore, it is possible to test for equal forecasting ability among respondents by comparing how close their rank sums are to their expected value.

Similarly, at the “macro” level, when the N forecasters are divided among M groups, $M = (g_1, g_2, \dots, g_M)$ of sizes equal to $N^G = (n_1, n_2, \dots, n_M)$ respectively, agents’ rank sums within each group equal $R_m = \sum_{t=1}^T \sum_{n=1}^{n_m} r_{n,t}$, while the expected rank sum among each group equals $0.5n_jT(N + 1)$. The variance for each individual rank is $N(N + 1)/12$, hence the variance of the sum of independent ranks within each group is $n_jTN(N + 1)/12$. The test statistic for equal average accuracy among groups can be constructed as follows (see Batchelor and Dua, 1990).

$$\chi_{F(M-1)}^2 = \frac{12}{TN(N+1)} \sum_{j=1}^M n_j^{-1} [R_m - 0.5n_jT(N + 1)]^2 \quad (3)$$

¹⁴ Note that T varies according to the forecast horizon considered.

Under the null hypothesis of no difference in average forecasting ability among groups the Friedman statistic (χ^2_F) can be approximated by a chi-square distribution with M-1 degrees of freedom.

The null and alternative hypothesis are stated as follows

$$H_0: \mu_1 = \mu_2 = \dots = \mu_m \quad (4)$$

$$H_1: \text{not all group means are equal} \quad (5)$$

3.1 – Data Problems and Limitations

Despite the appealing features of the Focus Survey it has also been affected by a pervasive problem in surveys worldwide: the irregularity of respondents' participation. The erratic participation frequency – which basically stems from three sources – entails data problems and affects inference. While some problems are manageable or even possible to circumvent others are not.

For example, during the analysed period (2002.1–2012.6) the Brazilian financial system witnessed major changes. Many financial institutions went bankrupt, while others were involved in takeovers and mergers. Therefore, through time many respondents simply ceased to exist and, of course, stopped providing forecasts. Such type of irregularity faced by the Survey is inevitable and there is nothing one can do to avoid them or deal with their consequences.

Some representative examples are: Banco Itaú – one of the main Brazilian commercial banks – which had previously bought all Bank of Boston's Brazilian assets in 2006, took control of Unibanco – another major commercial bank player – in 2009. The ABN AMRO Bank was bought by Santander in 2007. Nossa Caixa – one of the most important state controlled banks of Brazil – was incorporated in 2009 by Banco do Brasil, one of the biggest (state) commercial banks. In all those cases, assiduous participants unavoidably and regrettably disappeared from the Survey.

In two other situations, frequent respondents simply either stopped providing forecasts during some time periods, creating gaps in the database, or did not provide forecasts for all the horizons investigated, creating gaps for specific horizons. For example, in period “t” a given respondent might have provided inflation forecasts for each of the next seven months but not for the eighth to the twelfth months. Thus, in this

case one is unable to know its expectation for, say, nine and twelve-month ahead inflation. However, differently from the former case, in these two situations the effects on the database could sometimes be mitigated or even offset.

Indeed, when the database was being built some rules were set to handle them. First, in those cases where a specific respondent did not participate in month “ t ”, the (equivalent) forecast provided in month “ $t-1$ ” was used instead, whenever available. When the previous month forecast was also unavailable, the average forecast produced by the other respondents in month “ t ” was used. In this manner, one avoids using a possibly too outdated information set (e.g. forecasts reported at “ $t-2$ ”), which would probably have put that specific forecaster into a disadvantageous position relatively to others. Moreover, since the average forecast is “neutral” [see equation (3)], using it makes any rejection of H_0 more meaningful.

In those cases where respondents did not provide forecasts for all the surveyed horizons a simple and reliable solution was used, whenever possible.¹⁵ Sometimes not all monthly forecasts for the next twelve months are provided by the respondents, but just those for the immediately following months. It turns out, however, that in many cases the calendar year forecast is also provided. Therefore, one is able to get longer forecasts for that respondent by using the pro-rata monthly rates for the missing months, taking into account the monthly forecasts already informed. Since the shortest forecast horizon being analysed here is semi-annual, seasonality issues are not likely to be important and, therefore, approximations should be good enough.

Note, however, that the above procedures should be used with parsimony, otherwise one could simply start applying them to “fill all the gaps” in the database. Indeed, while using them allows one to “save” valuable data, increasing the sample size and improving inference, using them “too frequently” will make each forecaster data less representative, hindering inference.¹⁶

After a long and detailed filtering procedure, three databases were built: one for semi-annual forecasts, other for the three-quarter ahead horizon and another for annual forecasts. In the first case 29 participants were selected, in the second 28 and in the latter 32 respondents. Notice that these are large sized balanced panels, not only

¹⁵ This solution applies to the 6 and 9-month ahead forecasts, since the 12-month ahead forecasts are actually calendar-year forecasts.

¹⁶ Indeed, in many cases entire series of forecasts would be discarded just because of one or two missing observations.

compared to what one finds in the literature, but relative to the available sample size. Just to give an idea, in the second quarter of 2012 the total number of forecasters in the SPF in the U.S. totalled 39. In the Focus Survey, during the period analysed, the number of forecasters providing one-year ahead inflation forecasts ranged from a minimum of 35 to a maximum of 91. The median value was 57.

For the semi-annual forecasts, the above procedures were used for 13 out of 29 forecasters. Within those 13 cases, the “last month solution” alone was applied 7 times: once for 6 forecasters and twice for one forecaster. In the remaining 6 cases some combination of the three procedures above were used: they were used to recover four observations twice, three observations once and two observations three times.

For the three-quarter ahead horizon, the above procedures were also used 13 times. The “pro-rata solution” alone was applied 7 times, usually to get the last three months of the forecasting horizon. In the remaining 6 cases the last month procedure alone was used in four cases.

For the annual forecasts the above procedures were used for 8 out of 32 forecasters. Within those 8 cases, the “last month solution” was applied 6 times: once for 4 forecasters and twice for other 2 forecasters. Only for two forecasters the average forecast was (once) used.

These numbers – together with the likely good approximations provided by the procedures listed above – seem small enough to produce a clear positive effect on inference. Finally, just for completeness, there is also another obviously source for “infrequent” participation: the arrival of new respondents to the Survey database. This has been an increasingly important factor, since the number of respondents has grown substantially over time.

4 – Empirical Results

After a careful analysis and data compilation, as described above, agents were pooled into different groups according to their workplace, in order to test if the forecast performance varies with that characteristic. How many groups should one analyse to test such a hypothesis? Unfortunately, this decision was constrained by respondents’ participation frequency, since the erratic participation prevented some groups from being formed.

For example, during the Survey's period of existence very few respondents came from the real sector, and among those none have been a frequent participant, which means that it was not possible to select even one regular real sector forecaster. Thus, this theoretically relevant group is unfortunately absent from the analysis. Note, however, that its absence *per se* could be informative to our purposes. More specifically, why is it so difficult to find real sector agents willing (or able) to participate in the Survey?¹⁷ The absence might reflect the fact that rather than producing their own *macro* forecasts, real sector firms usually get them from other groups of agents such as banks or economic consultancies firms.

If this is the main reason the objections regarding the “unbalanced” nature of the survey becomes less convincing. Moreover, that absence turns consultancies into an even more theoretically relevant group. Not only they also come from the non-financial sector, but they are a very good control group that could “replace” real sector firms. Indeed, those firms are paid to provide as accurate as possible forecasts, so it is extremely unlikely that they act strategically. Moreover, their forecasts are “consumed” by different kinds of agents, with different interests, including both banks and real sector firms.

At the end, besides the economic consultancy firms the following three groups of agents were thought to bear sufficient different characteristics (and incentives) so as to be analysed separately: commercial banks, investment banks and asset management firms. While the former group is engaged in the typical banking business, having its funding coming mostly from sight deposits and its core business largely related to lending to individuals, investment banks are not allowed to offer checking accounts and are involved in many other business, such as mergers and acquisitions, financial advisory work, equity and bond financing, project finance, etc... Also, while traditional banking is a more regulated activity, investment banking faces additional degrees of freedom to carry out their businesses. Finally, asset management firms are financial institutions whose main business is to manage other peoples' money and, in principle, should rely more heavily on being accurate in order to survive. Table A in the Appendix provides the list of participants within each group, for each forecasting horizon. Given the reasons laid out above, it would be interesting to check how these three groups fare

¹⁷ See footnote 12.

against the economic consultancy group regarding their relative forecasting performance.

It should be noted that although the number of participants in the economic consultancy group is smaller than in the other groups, as Table 1 shows, those firms are by far the most important ones in that segment. Moreover, similarly to the banking sector that segment is highly concentrated. In other words, despite the smaller number of respondents they are likely to represent a similar share of their industry as the other groups' participants.

Table 1
IPCA Inflation Forecast Groups' Errors: Summary Statistics and Normality Tests

Groups	Summary Statistics				Normality Tests			
	Forecasters	Total Forecasts	Mean	SD	Skewness	Excess Kurtosis	Chi ² Test	Individual Rejections ¹
6-Month Ahead								
Commercial Banks	9	180	-0.48	1.47	-3.11	10.16	41.62***	100%***
Investment Banks	9	180	-0.50	1.44	-3.14	10.21	43.98***	100%***
Asset Management	7	140	-0.45	1.43	-3.22	10.54	49.93***	100%***
Consultancies	4	80	-0.48	1.44	-3.11	10.12	41.64***	100%***
9-Month Ahead								
Commercial Banks	9	126	-0.77	1.98	-2.17	4.52	13.83***	78%***
Investment Banks	9	126	-0.84	1.96	-2.18	4.59	13.87***	67%***
Asset Management	6	84	-0.70	2.10	-2.18	4.69	13.47***	100%***
Consultancies	4	56	-0.78	1.96	-2.28	4.92	15.99***	100%***
12-Month Ahead								
Commercial Banks	9	90	-0.91	2.74	-1.52	2.18	7.17**	78%**
Investment Banks	10	100	-0.99	2.74	-1.57	2.24	7.23**	100%**
Asset Management	8	80	-0.81	2.79	-2.02	3.25	11.00***	70%**
Consultancies	5	50	-0.90	2.70	-1.57	2.18	7.10**	50%**

*, ** and *** indicate significance at 10%, 5% and 1%, respectively.

(1) Percentage of rejections within each group of forecasters according to the significance level indicated by the number of stars. For the 9-month ahead horizon the normality tests that were not rejected at 1% were significant at 10% in two occasions and insignificant in 3 occasions. For the 12-month ahead horizon all the normality tests that were not rejected at 1% were significant at 5%.

Table 1 presents summary statistics for each group's forecast errors, as well as associated normality tests. As it can be seen, the mean forecast errors are reasonably similar across groups, although the errors from the asset management firms have been smaller than the others. Note also that inflation has been under predicted by an economically relevant amount by all groups during the sample analysed.¹⁸ Finally, the size of the forecast errors grows with the forecasting horizon as expected.

¹⁸ Average semi-annual, nine months and annual inflation rates during the sample period were, respectively, 3.2%, 4.8% and 6.5%.

The normality tests show that the widely used normality assumption is clearly inadequate in all cases (i.e. all groups and horizons). Errors are skewed to the left and leptokurtic (see Figure 1 in the Appendix). Moreover, as the last column shows this is not an outcome of aggregation. When individual errors are tested normality is clearly rejected once again. This evidence shows the importance of using a non-parametric approach.

Table 2 shows that the data is also highly correlated, both within groups and between groups, and that feature should also be taken into account when inferences are being made. This should come as a surprise, since not only agents share very similar information sets, but also interact a lot with each other when making forecasts.

Table 2
IPCA Inflation Forecast Errors: Correlation

Groups	Within Groups		Between Groups			
	Min	Max	Commercial Banks	Investment Banks	Asset Management	Consultancies
6-Month Ahead						
Commercial Banks	0.913	0.985	1.000			
Investment Banks	0.908	0.988	0.998	1.000		
Asset Management	0.876	0.989	0.998	0.998	1.000	
Consultancies	0.877	0.972	0.994	0.996	0.995	1.000
9-Month Ahead						
Commercial Banks	0.848	0.991	1.000			
Investment Banks	0.904	0.983	0.998	1.000		
Asset Management	0.913	0.984	0.998	0.997	1.000	
Consultancies	0.961	0.990	0.994	0.996	0.995	1.000
12-Month Ahead						
Commercial Banks	0.872	0.995	1.000			
Investment Banks	0.880	0.998	0.999	1.000		
Asset Management	0.501	0.995	0.986	0.989	1.000	
Consultancies	0.900	0.993	0.999	0.999	0.989	1.000

Before proceeding to the results of the Friedman test it should be noted that the literature recommends (e.g. Seskin, 2000) that when there are tied ranks the Friedman statistic should be corrected by a factor that take those ties into consideration. The correction increases the power of the test. The procedure is shown by equations (6) and (7).

$$\chi_{F^*}^2 = \frac{\chi_F^2}{c} \quad (6)$$

where

$$C = 1 - \frac{\sum_{i=1}^S (t_i^3 - t_i)}{T(m^3 - m)} \quad (7)$$

S is the total number of tied series, and t_i is the number of tied rankings in the i_{th} series of ties, while m is the number of treatments (in our case the number of groups) being compared and T the number of periods.

In addition, Iman and Davenport (1980) argues that the chi-square approximation of the Friedman test is undesirable conservative and recommends the use of the following F distribution approximation instead

$$F_{ID} = \frac{(T-1)\chi_F^2}{T(m-1)-\chi_F^2} \quad (8)$$

where, in our case, m is the number of groups of forecasters and T the number of periods. This variation of the Friedman statistic is F-distributed with $(m - 1)$ and $(T - 1)(m - 1)$ degrees of freedom. Table 3 shows the results of the three test statistics for all three horizons.

Table 3
Friedman Test Results

Tests	χ_F^2	χ_F^{2*}	Critical Value ¹	F_{ID}^2	Critical Value ³
6-Month Ahead	5.95 (0.114)	6.03 (0.111)	7.82	2.08 (0.107)	2.76
9-Month Ahead	4.79 (0.188)	5.16 (0.161)	7.89	1.67 (0.159)	2.85
12-Month Ahead	3.65 (0.302)	4.10 (0.258)	7.80	1.25 (0.251)	2.96

The figures between parentheses are p-values.

(1) All critical values refer to the 5% significance level. Critical values for 9 and 12-month ahead horizons are exact values from Martin *et al.* (1993). The critical value for the 6-month ahead horizon is from the standard chi-square table.

(2) The tie-corrected Friedman statistic was used in equation (8).

(3) The critical value for the Iman & Davenport test is from the standard F table.

The evidence is similar regardless of the test statistic used. More specifically, the null was not rejected at the traditional 5% or 10% significance levels in any of the three forecasting horizons investigated. However, the value of the Friedman statistic for the 6-month ahead forecasts is barely significant at 10%. Given that non parametric tests are not the most powerful ones, this evidence might indeed be signaling real differences in the forecast ability among groups and deserves further investigation. Note also that the tie-corrected Friedman statistic is more powerful than the original one, as expected. This

is especially true for the one year ahead forecasts, a horizon when ties occur more frequently.¹⁹

Given the relatively low p-value associated with the 6-month ahead forecasts, it seems worthwhile to perform *post hoc* tests. Note that these tests will also be carried out for the other two forecast horizons, despite the larger p-values involved. This procedure is recommended by statisticians such as Hsu (1996, page 177) who calls to attention that “An unfortunate common practice is to pursue multiple comparisons only when the null hypothesis of homogeneity is rejected.”, and Zar (2010, page 226) who warns that “Indeed, power may be lost if a multiple-comparison test is performed only if the ANOVA concludes a significant difference among means.”

The next issue regards the significance levels to be used in those tests. Two factors should be considered here. First, given the early established importance of using the economic consultancy firms as the control group, there is no need to correct the familywise error rate. In this regard Sheskin (2004, Chapter 22, original emphasis) argues that “When a limited number of comparisons are planned prior to collecting the data, most sources take the position that a researcher is not obliged to control the value of α_{FW} . In such a case, the per comparison Type I error rate (α_{PC}) will be equal to the prespecified value of alpha.”

Table 4
Groups’ Mean Ranks for the Friedman Statistic¹

	Commercial Banks	Investment Banks	Asset Management	Economic Consultancies
6-Month Ahead	15.46	15.37	13.52	15.71
9-Month Ahead	14.88	15.42	13.15	13.59
12-Month Ahead	16.72	17.53	14.88	16.64

(1) The expected rank varies in the three horizons due to differences in the number of participants.

Second, as Table 4 shows, the difference in mean ranks between the asset management and investment bank groups seem too big to be a random outcome. Thus, it is worthwhile to test if those differences are statistically significant. In such cases (i.e. unplanned comparisons) the literature argues for the use of corrected p-values, in order to avoid inflating the Type I error rate. Hence, two sets of tests are carried out next: the

¹⁹ More specifically, annual forecasts are calendar year forecasts, and this creates a clear “focal period” increasing the probability of ties.

first has consultancies firms as the (planned) control group and the second considers the asset management firms as the (unplanned) control group.

The test statistic is given by equation (9), which takes into account the different number of respondents in each group.

$$Z = (R_c - R_i) / \sqrt{\left(\frac{1}{m_c} + \frac{1}{m_i}\right) \frac{m(m+1)}{12T}} \quad (9)$$

Where R_c and R_i are, respectively, the average rankings from the control and non-control groups, where i indexes the specific non-control group under comparison. Similarly, m_c and m_i are the number of respondents in each group.

Given that there is no theoretical reason for strategic behavior by the consultancy group, the null is that forecasts from that group should be more accurate than the other groups. Hence, the null and alternative hypothesis are stated as follows

$$H_0: \mu_{cons} < \mu_i \quad (10)$$

$$H_1: \mu_{cons} \geq \mu_i \quad (11)$$

where i = comercial banks, investment banks or asset management firms.

Similarly, given the evidence in Table 4, the null and alternative hypotheses when the control group are the asset management firms are:

$$H_0: \mu_{asset} < \mu_i \quad (12)$$

$$H_1: \mu_{asset} \geq \mu_i \quad (13)$$

Where i = comercial banks, investment banks or consultancies.

Table 5 provides the results when the economic consultancies are the control group. As it can be seen the only significant difference is for the 6-month horizon when the comparison is made against the asset management group. More specifically, the asset management group seems to provide better forecasts than the consultancy group when predicting inflation 6-month ahead. This result is not supportive of the hypothesis that the unbalanced nature of the Focus Survey produces biased (i.e. less accurate)

median forecasts. On the contrary, there is evidence, although only for the shorter horizon, that one group of agents from the financial system produces better forecaster than the non-financial control – group.

Table 5
Pairwise Tests Having Economic Consultancies Firms as the Control Group ¹

Groups	Mean Rank Difference	Standard Deviation	Z Statistic	P Value	Alfa ($\alpha = 0.05$)
6-Month Ahead					
Commercial Banks	0.25	1.12	0.23	0.41	0.05
Investment Banks	0.34	1.12	0.31	0.38	0.05
Asset Management	2.19	1.16	1.88	0.03	0.05
9-Month Ahead					
Commercial Banks	-1.29	1.32	-0.97	0.84	0.05
Investment Banks	-1.84	1.32	-1.39	0.92	0.05
Asset Management	0.43	1.42	0.31	0.38	0.05
12-Month Ahead					
Commercial Banks	-0.08	1.65	-0.05	0.52	0.05
Investment Banks	-0.89	1.62	-0.55	0.71	0.05
Asset Management	1.76	1.69	1.04	0.15	0.05

Regarding the second set of comparisons (i.e. asset management firms as the control group) the Hochberg's procedure (Hochberg, 1988) is used. This test – considered to be more powerful than the Bonferroni-Dunn test – is a procedure that sequentially tests the hypotheses, which are ordered by their (decreasing) level of significance, while adjusting the value of alfa in a step-up way.

It works as follows: let $p_1, p_2, \dots, p_{K-1}$ be the p -values ordered from the smallest to largest, so that $p_1 \leq p_2 \leq \dots \leq p_{K-1}$, where K is the number of groups, and $H_1, H_2, \dots, H_{K-1}$ be the associated hypotheses. The Hochberg step-up procedure compares each p_i , beginning with the least significant test statistic, with $\alpha/(K - i)$. The comparison continues until the test rejects the null hypothesis. That is, as long as a given test fails to reject the null one proceeds to the next comparison. As soon as the null hypothesis is rejected, all the remaining hypotheses – with smaller p -values – are automatically rejected as well. Table 6 provides the results.²⁰

²⁰ Instead of correcting the critical values one can equivalently correct the estimated p -values in the following way: $p_i^* = (K - i)p_i$ where p_i^* is the Hochberg corrected p -value. This makes comparisons easier to interpret and was the strategy followed in Table 6.

Table 6
Pairwise Tests Having Asset Management Firms as the Control Group¹

i	Groups	Mean Rank Difference ¹	Standard Deviation	Z Statistic	P Value	Hochberg's Adjusted P-Value ¹
6-Month Ahead						
1	Commercial Banks	-1.94	0.94	-2.07	0.019	0.057
2	Investment Banks	-1.85	0.94	-1.98	0.024	0.048
3	Consultancies	-2.19	1.16	-1.88	0.030	0.030
9-Month Ahead						
1	Investment Banks	-2.27	1.16	-1.96	0.025	0.075
2	Commercial Banks	-1.72	1.16	-1.49	0.069	0.137
3	Consultancies	-0.43	1.42	-0.31	0.380	0.380
12-Month Ahead						
1	Investment Banks	-2.65	1.41	-1.88	0.030	0.090
2	Commercial Banks	-1.84	1.44	-1.27	0.101	0.203
3	Consultancies	-1.76	1.69	-1.04	0.149	0.149

As it can be seen all nulls are rejected in the 6-month ahead horizon. That is, as Table 5 had already suggested the forecasting performance of the asset management group is significantly better than the other groups at a significance level of 5%. As to the other two horizons, there is evidence of different forecast ability only between the control group and the investment banks group at significance level lower than 10%. More specifically, the investment banks group seems to provide worse forecasts than the control group for the 9 and 12-month ahead horizons.

5 – Conclusion

In recent years some agents have expressed doubts about the usefulness of the median inflation forecast collected by the Focus Survey. The reasoning is based on an alleged strategic behaviour. That hypothesis was tested using a non parametric approach based on rankings by pooling suitable respondents into four different groups according to their workplace, and testing whether their average forecast accuracy has been significantly different from one another. One of those groups is economic consultancies firms, which are certainly paid to be as accurate as possible, and provide forecasts to several kinds of agents. Therefore, this group acted as an important control group.

The null hypothesis of equal forecasting accuracy – measured by ranks – is rejected in the following cases. Asset management firms produce better 6-month ahead forecasts than the other groups of agents: economic consultancy firms, commercial banks and investment banks. Moreover, they also seem to produce more accurate forecasts than investment banks for the 9 and 12-month ahead horizons. Finally, there

was no evidence that commercial and investment banks differ in their forecasting abilities from economic consultancies firms.

The favourable evidence on the asset management firms is consistent with the fact that the vast majority of revenues of those firms are tightly linked to the quality of their forecasts, since they do not provide many of the services provided by commercial and investment banks.

Consequently, claims that the Focus Survey median forecast is less reliable due to the “unbalanced” nature of the survey are not supported by the data. Moreover, the widespread negative bias of forecast errors in all analysed horizons is inconsistent with the alleged hawkish behaviour of financial institutions.

Finally, although the evidence presented here does not necessarily rule out strategic behaviour, it suggests that if it does exist it is confined to a small number of agents or its empirical relevance is of second order.

References

- Abreu, Ildeberta (2011).** “International Organisations’ vs. Private Analysts’ Forecasts: An Evaluation”, Banco de Portugal Working Paper, N° 20 (July).
- Batchelor, R. A. (1990).** “All Forecasters Are Equal”, *Journal of Business & Economic Statistics*, Vol. 8, N° 1 (January).
- Bakhshi, H. and A. Yates (1998).** “Are UK Inflation Expectations Rational?”, Bank of England Working Paper No 81.
- Carta Capital (2011).** “Sobre pombos e raposas”, March 23.
- Central Bank of Brazil (2012).** “Mecanismos de Transmissão da Política Monetária nos Modelos do Banco Central”. *Inflation Report*, March 2012.
- _____ (2011). *Inflation Report*, March 2011.
- _____ (2009). “The Lags in the transmission of Monetary Policy to Prices”. *Inflation Report*, June 2012.
- Croushore, Dean (1993).** “Introducing: The Survey of Professional Forecasters”. Federal Reserve Bank of Philadelphia Business Review, November/December.
- D’Agostino, A., K. McQuinn and K. Whelan (2012).** “Are Some Forecasters Really Better Than Others?”, *Journal of Money Credit and Banking*, Vol. 44, N° 4.
- Faust, J., and J. H. Wright (2009).** “Comparing Greenbook and Reduced Form Forecasts using a Large Realtime Dataset”, *Journal of Business and Economic Statistics*, Vol. 27, N° 4.
- Friedman, M. (1937).** “The Use of Ranks to Avoid the Assumption of Normality Implicit in the Analysis of Variance”, *Journal of the American Statistical Association*, 32.
- _____ (1940). “A Comparison of Alternative Tests of Significance for the Problem of m Rankings”, *The Annals of Mathematical Statistics*, 11.
- Gavin, W. and R. Mandal (2001).** “Forecasting Inflation and Growth: Do Private Forecasts Match Those of Policymakers?” *Review*, Federal Reserve Bank of St. Louis, May/June.
- _____ (2003). “Evaluating FOMC Forecasts”, *International Journal of Forecasting* 19.
- Granger, Clive W. J. and Mark J. Machina (2006).** “Forecasting and Decision Theory?” *Handbook of Economic Forecasting*, Chapter 2.
- Joutz, F. and H.O. Stekler, (2000).** “An Evaluation of the Predictions of the Federal Reserve,” *International Journal of Forecasting*, 16.
- Harvey, D. I. and P. Newbold (2003).** “The non-normality of some macroeconomic forecast errors”, *International Journal of Forecasting*, 19.
- Hochberg, Y. (1988).** “A Sharper Bonferroni Procedure for Multiple Tests of Significance”, *Biometrika*, Vol. 75.
- Hsu, J. C. (1996).** *Multiple Comparisons: Theory Methods*, Chapman and Hall, London.

- Iman, R. L. and J. M. Davenport (1980).** “Approximations of the Critical Region of the Friedman Statistic”, *Communication in Statistics – Theory and Methods* A9(6).
- Lahiri, K., C. Teigland. and M. Zaporowski (1988).** “Interest Rates and the Subjective Probability Distribution of Inflation Forecasts”. *Journal of Money Credit, and Banking*, Vol. 20, No 2.
- Laster, D., Paul Bennett and In S. Geoum (1999).** “Rational Bias in Macroeconomic Forecasts”. *The Quarterly Journal of Economics*, February.
- Marques, A. B. C., P. Fachada and D. C. Cavalcanti (2003).** “Sistema Banco Central de Expectativas de Mercado”, Banco Central do Brasil, Nota Técnica Nº 36, Junho.
- Martin, Louise, R. Leblanc and N. Ky Toan (1993).** “Tables for the Friedman Rank Test”, *The Canadian Journal of Statistics*, Vol. 21, Nº 1.
- OESP (2011).** “Focus muda para BC não ficar refém de visão financeira”. O Estado de São Paulo, 9 de março.
- Pons-Novell, Jordi (2003).** “Strategic Bias, Herding Behavior and Economic Forecasts,” *Journal of Forecasting*, Vol. 22.
- Romer, C. and D. Romer (2000).** “Federal Reserve Information and the Behavior of Interest Rates,” *American Economic Review*, **90** (2000) 429-57.
- Sheskin, D. J. (2004).** *Handbook of Parametric and Nonparametric Statistical Procedures*, 3rd Edition, Chapman & Hall/CRC.
- Sims, C. (2002).** “The Role of Models and Probabilities in the Monetary Policy Process”, *Brookings Papers on Economic Activity*, 2.
- Stekler, H. O. (1987).** “Who Forecasts Better?”, *Journal of Business & Economic Statistics*, Vol. 5, Nº 1 (January).
- Thomas Jr., Lloyd B (1999).** “Survey Measures of Expected U.S. Inflation”. *Journal of Economic Perspectives*, Vol. 13(4), Fall.
- Zar J. H. (2010).** *Biostatistical Analysis*, 5th Edition, Prentice Hall.
- Zarnowitz, Victor and Louis A. Lambros (1987).** “Consensus and Uncertainty in Economic Prediction”, *Journal of Political Economy*, Vol. 95(3).

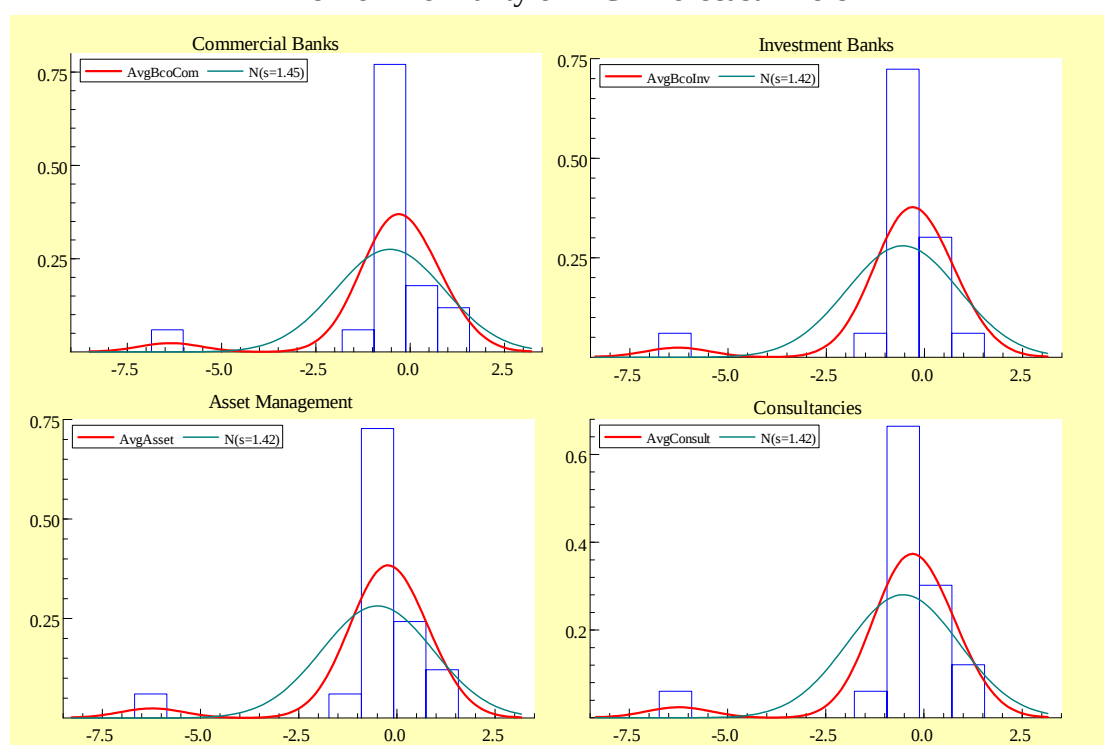
Appendix

Table A
Group's Respondents for Each Horizon (2002.1–2012.6)*

Commercial Banks	Investment Banks	Asset Management	Consultancies
ABC (6, 9, 12)	BBM (6, 9, 12)	Brascan (6, 9, 12)	LCA (6, 9, 12)
Bradesco (6, 9, 12)	BES (6, 9, 12)	Gap (6, 12)	MBA (12)
Banco do Brasil (6, 9, 12)	Credit Suisse (6, 9, 12)	HSBC (6, 9, 12)	MCM (6, 9, 12)
Citibank (6, 9, 12)	Deutsch (6, 9, 12)	Icatu (6, 9, 12)	Rosenberg (6, 9, 12)
Cruzeiro do Sul (6, 9, 12)	Fibra (6, 9, 12)	Itaú (6, 9, 12)	Tendências (6, 9, 12)
HSBC (6, 9, 12)	ING (12)	Nobel (6, 9, 12)	
Itaú (6, 9, 12)	JP Morgan (9)	Opportunity (6, 12)	
Safra (6, 9, 12)	Pactual (6, 12)	Votorantim (9, 12)	
Santander (6, 9, 12)	Paribas (6, 9, 12)		
	Schahin (6, 9, 12)		
	Votorantim (6, 9, 12)		

(*) Numbers between parentheses indicate in what horizons that respondent's forecasts are considered.

Figure 1
The Non-Normality of IPCA Forecast Errors



Banco Central do Brasil

Trabalhos para Discussão

Os Trabalhos para Discussão do Banco Central do Brasil estão disponíveis para download no website
<http://www.bcb.gov.br/?TRABDISCLISTA>

Working Paper Series

The Working Paper Series of the Central Bank of Brazil are available for download at
<http://www.bcb.gov.br/?WORKINGPAPERS>

- | | | |
|------------|---|----------|
| 284 | On the Welfare Costs of Business-Cycle Fluctuations and Economic-Growth Variation in the 20th Century
<i>Osmani Teixeira de Carvalho Guillén, João Victor Issler and Afonso Arinos de Mello Franco-Neto</i> | Jul/2012 |
| 285 | Asset Prices and Monetary Policy – A Sticky-Dispersed Information Model
<i>Marta Areosa and Waldyr Areosa</i> | Jul/2012 |
| 286 | Information (in) Chains: information transmission through production chains
<i>Waldyr Areosa and Marta Areosa</i> | Jul/2012 |
| 287 | Some Financial Stability Indicators for Brazil
<i>Adriana Soares Sales, Waldyr D. Areosa and Marta B. M. Areosa</i> | Jul/2012 |
| 288 | Forecasting Bond Yields with Segmented Term Structure Models
<i>Caio Almeida, Axel Simonsen and José Vicente</i> | Jul/2012 |
| 289 | Financial Stability in Brazil
<i>Luiz A. Pereira da Silva, Adriana Soares Sales and Wagner Piazza Gaglianone</i> | Aug/2012 |
| 290 | Sailing through the Global Financial Storm: Brazil's recent experience with monetary and macroprudential policies to lean against the financial cycle and deal with systemic risks
<i>Luiz Awazu Pereira da Silva and Ricardo Eyer Harris</i> | Aug/2012 |
| 291 | O Desempenho Recente da Política Monetária Brasileira sob a Ótica da Modelagem DSGE
<i>Bruno Freitas Boynard de Vasconcelos e José Angelo Divino</i> | Set/2012 |
| 292 | Coping with a Complex Global Environment: a Brazilian perspective on emerging market issues
<i>Adriana Soares Sales and João Barata Ribeiro Blanco Barroso</i> | Oct/2012 |
| 293 | Contagion in CDS, Banking and Equity Markets
<i>Rodrigo César de Castro Miranda, Benjamin Miranda Tabak and Mauricio Medeiros Junior</i> | Oct/2012 |
| 293 | Contágio nos Mercados de CDS, Bancário e de Ações
<i>Rodrigo César de Castro Miranda, Benjamin Miranda Tabak e Mauricio Medeiros Junior</i> | Out/2012 |

294	Pesquisa de Estabilidade Financeira do Banco Central do Brasil <i>Solange Maria Guerra, Benjamin Miranda Tabak e Rodrigo César de Castro Miranda</i>	Out/2012
295	The External Finance Premium in Brazil: empirical analyses using state space models <i>Fernando Nascimento de Oliveira</i>	Oct/2012
296	Uma Avaliação dos Recolhimentos Compulsórios <i>Leonardo S. Alencar, Tony Takeda, Bruno S. Martins e Paulo Evandro Dawid</i>	Out/2012
297	Avaliando a Volatilidade Diária dos Ativos: a hora da negociação importa? <i>José Valentim Machado Vicente, Gustavo Silva Araújo, Paula Baião Fisher de Castro e Felipe Noronha Tavares</i>	Nov/2012
298	Atuação de Bancos Estrangeiros no Brasil: mercado de crédito e de derivativos de 2005 a 2011 <i>Raquel de Freitas Oliveira, Rafael Felipe Schiozer e Sérgio Leão</i>	Nov/2012
299	Local Market Structure and Bank Competition: evidence from the Brazilian auto loan market <i>Bruno Martins</i>	Nov/2012
299	Estrutura de Mercado Local e Competição Bancária: evidências no mercado de financiamento de veículos <i>Bruno Martins</i>	Nov/2012
300	Conectividade e Risco Sistêmico no Sistema de Pagamentos Brasileiro <i>Benjamin Miranda Tabak, Rodrigo César de Castro Miranda e Sergio Rubens Stancato de Souza</i>	Nov/2012
300	Connectivity and Systemic Risk in the Brazilian National Payments System <i>Benjamin Miranda Tabak, Rodrigo César de Castro Miranda and Sergio Rubens Stancato de Souza</i>	Nov/2012
301	Determinantes da Captação Líquida dos Depósitos de Poupança <i>Clodoaldo Aparecido Annibal</i>	Dez/2012
302	Stress Testing Liquidity Risk: the case of the Brazilian Banking System <i>Benjamin M. Tabak, Solange M. Guerra, Rodrigo C. Miranda and Sergio Rubens S. de Souza</i>	Dec/2012
303	Using a DSGE Model to Assess the Macroeconomic Effects of Reserve Requirements in Brazil <i>Waldyr Dutra Areosa and Christiano Arrigoni Coelho</i>	Jan/2013
303	Utilizando um Modelo DSGE para Avaliar os Efeitos Macroeconômicos dos Recolhimentos Compulsórios no Brasil <i>Waldyr Dutra Areosa e Christiano Arrigoni Coelho</i>	Jan/2013
304	Credit Default and Business Cycles: an investigation of this relationship in the Brazilian corporate credit market <i>Jaqueline Terra Moura Marins and Myrian Beatriz Eiras das Neves</i>	Mar/2013

304	Inadimplência de Crédito e Ciclo Econômico: um exame da relação no mercado brasileiro de crédito corporativo <i>Jaqueline Terra Moura Marins e Myrian Beatriz Eiras das Neves</i>	Mar/2013
305	Preços Administrados: projeção e repasse cambial <i>Paulo Roberto de Sampaio Alves, Francisco Marcos Rodrigues Figueiredo, Antonio Negromonte Nascimento Junior e Leonardo Pio Perez</i>	Mar/2013
306	Complex Networks and Banking Systems Supervision <i>Theophilos Papadimitriou, Periklis Gogas and Benjamin M. Tabak</i>	May/2013
306	Redes Complexas e Supervisão de Sistemas Bancários <i>Theophilos Papadimitriou, Periklis Gogas e Benjamin M. Tabak</i>	Mai/2013
307	Risco Sistêmico no Mercado Bancário Brasileiro – Uma abordagem pelo método CoVaR <i>Gustavo Silva Araújo e Sérgio Leão</i>	Jul/2013
308	Transmissão da Política Monetária pelos Canais de Tomada de Risco e de Crédito: uma análise considerando os seguros contratados pelos bancos e o spread de crédito no Brasil <i>Debora Pereira Tavares, Gabriel Caldas Montes e Osmani Teixeira de Carvalho Guillén</i>	Jul/2013
309	Converting the NPL Ratio into a Comparable Long Term Metric <i>Rodrigo Lara Pinto Coelho and Gilneu Francisco Astolfi Vivan</i>	Jul/2013