

Forecasting presidential elections: When to change the model

Michael S. Lewis-Beck^{a,*}, Charles Tien^{b,1}

^a Department of Political Science, University of Iowa, 341 Schaeffer Hall, Iowa City, Iowa 52242, United States

^b Department of Political Science, Hunter College & the Graduate Center, CUNY, 695 Park Avenue, New York, NY 10065, United States

Abstract

Here, we address the issue of forecasting from statistical models, and how they might be improved. Our real-world example is the forecasting of US presidential elections. First, we ask whether a model should be changed. To illustrate problems and opportunities, we examine the forecasting history of different models, in particular our own, which has tried to foresee presidential selection since 1984. We apply what we learn to the question of whether our Jobs model, which offered an accurate ex ante point estimate for 2004, should be changed for 2008. We conclude there is room for judicious, theory-driven adjustment, but also raise a caution about inadvertent curve-fitting. Some evidence is offered that simple core models, based on strong theory, may perform almost as well as more stretched models.

© 2008 International Institute of Forecasters. Published by Elsevier B.V. All rights reserved.

Keywords: Adjusting forecasts; Econometric models; Evaluating forecasts; Model selection; Regression

1. Introduction

Traditionally, political science has focused on explanation, rather than forecasting. Some scholars still argue that forecasting lies outside the province of proper political science (Colomer, 2007; Van Der Eijk, 2005). Be that as may be, forecasting has become an important sub-field in the discipline, at home and abroad (see the reviews in Lewis-Beck & Rice, 1992; Lewis-Beck, 2005). Once it is accepted as a valid scientific enterprise, there are different ways to do it. Essentially, there are polls, models, and markets. Polls examine vote intentions (or expectations), as expressed in current

public opinion surveys (Lewis-Beck & Tien, 1999). Models explore statistical patterns in aggregate election data. Markets involve the purchase of political stocks in the different candidates, with forecasts based on the most frequently purchased stock (Forsythe, Nelson, Neumann, & Wright, 1992). It is also possible to combine information from the different approaches, and each method is aware of the existence of others. For example, market players may use information from polls and models in making their purchases. While the market approach is gaining in popularity, polls and models remain the dominant approaches. Often, these latter two are pitted against each other, in competition (Lewis-Beck, 2001). It is possible to show that, in the long-run, the two methods yield about the same error (Campbell, 2004; Lewis-Beck, 2005).

Here, we address the topic of forecasting using statistical models, and how they might be improved. Our real-world example is the forecasting of US presidential

* Corresponding author. Tel.: +1 319 335 2358; fax: +1 319 335 3400.

E-mail addresses: michael-lewis-beck@uiowa.edu (M.S. Lewis-Beck), ctien@hunter.cuny.edu (C. Tien).

¹ Tel.: +1 212 772 5494; fax: +1 212 650 3669.

Table 1

The independent variables of some US Presidential Election forecasting models

Forecaster	Independent variables
Abramowitz (1988)	Popularity, GNP, In-party Terms.
Brody and Sigelman (1983)	Popularity.
Campbell and Wink (1990)	Presidential trial-heat poll, GNP.
Fair (1978)	GNP, incumbency, time trend.
Hibbs (1982)	Personal income.
Holbrook (1996)	Popularity, pocketbook, In-party Terms.
Lewis-Beck and Rice (1984)	Popularity, GNP.
Lockerbie (2000)	Personal income, economic future, years.
Norpoth (1996)	Past votes, GNP, inflation, primary.
Wlezien and Erikson (1996)	Leading indicators, popularity.

Source: Lewis-Beck and Tien (2002), with some updates. Some of these forecasters have published more than one paper on presidential forecasting, in which case the earliest one is cited. To completely appreciate these specifications, the work itself must of course be read.

elections. First, we ask whether a model should be changed. To illustrate problems and opportunities, we examine the forecasting history of different models, in particular our own, which has tried to foresee presidential selection since 1984. We apply what we learn to the question of whether our Jobs model, which offered an accurate ex ante point estimate for 2004, should be changed for 2008. We conclude that there is room for judicious, theory-driven adjustment, but also raise a caution about inadvertent curve-fitting. Some evidence is offered that simple core models, based on strong theory, may perform almost as well as more stretched models.

2. Should a model be changed?

Suppose that Jane Smith has the following forecasting model, $V = a + bX + cD + e$, which predicts ex ante the precise percentage point vote share (V) for the winning presidential candidate in 2004. Since she is exactly right, she has a relatively low incentive to change the model before forecasting the upcoming 2008 election. However, say she is off by +4.0 percentage points. She may conclude that the model did not work very well. Her incentive for changing the model now appears to be relatively higher. In other words, in the face of a large prediction error, there is a temptation to go back to the drawing board. Nevertheless, some analysts suggest that models should not change even then. Campbell (2004, p. 735), for example, contends that “Model stability (the constancy of model specification from one election to the next) must be a goal of election forecasting...”. Lewis-Beck and Rice (1992), however, contradict this view, arguing that with each election trial we identify new sources of error, thus providing us with valuable information to be incorpo-

rated into model revisions. We will return to this controversy at more length later. At this point, suffice it to say that we have traditionally followed the second strategy—that of revision based on our errors.²

3. The evolution of US presidential election forecasting models

A typical US presidential election forecasting model expresses vote choice as a function of key prior macroeconomic and macro-political conditions, estimated over a time-series of contests. For example,

$$\text{Vote} = f(\text{Economic Growth, Presidential Popularity}), \quad (1)$$

where Vote=percentage vote for the incumbent party, Economic Growth=percentage change in GNP, and Presidential Popularity=percentage approval of the President in a national opinion poll, measured over the last 15 presidential elections back to 1948. While this core political economy model may be typical, it is not the only available model. In Table 1, a summary of various leading US presidential election models is provided.

One can see that the specification tends to change from model to model. For example, while most of the models include an economic measure, that of Brody and Sigelman (1983) does not. And, while most include a popularity measure, those of Hibbs (1982), Norpoth (1996), and Fair

² Note that both the order of the errors and the size of the errors matter. That is, if the model has a four percentage point error for the 1956 election but gets the 2004 election right, there is a low incentive to change the forecast for the 2008 election.

Download English Version:

<https://daneshyari.com/en/article/998449>

Download Persian Version:

<https://daneshyari.com/article/998449>

[Daneshyari.com](https://daneshyari.com)