



American Finance Association

Reputation and Performance Among Security Analysts

Author(s): Scott E. Stickel

Source: *The Journal of Finance*, Vol. 47, No. 5 (Dec., 1992), pp. 1811-1836

Published by: Blackwell Publishing for the American Finance Association

Stable URL: <http://www.jstor.org/stable/2328997>

Accessed: 25/11/2009 12:00

Your use of the JSTOR archive indicates your acceptance of JSTOR's Terms and Conditions of Use, available at <http://www.jstor.org/page/info/about/policies/terms.jsp>. JSTOR's Terms and Conditions of Use provides, in part, that unless you have obtained prior permission, you may not download an entire issue of a journal or multiple copies of articles, and you may use content in the JSTOR archive only for your personal, non-commercial use.

Please contact the publisher regarding any further use of this work. Publisher contact information may be obtained at <http://www.jstor.org/action/showPublisher?publisherCode=black>.

Each copy of any part of a JSTOR transmission must contain the same copyright notice that appears on the screen or printed page of such transmission.

JSTOR is a not-for-profit service that helps scholars, researchers, and students discover, use, and build upon a wide range of content in a trusted digital archive. We use information technology and tools to increase productivity and facilitate new forms of scholarship. For more information about JSTOR, please contact support@jstor.org.



Blackwell Publishing and American Finance Association are collaborating with JSTOR to digitize, preserve and extend access to *The Journal of Finance*.

<http://www.jstor.org>

Reputation and Performance Among Security Analysts

SCOTT E. STICKEL*

ABSTRACT

Members of the *Institutional Investor* All-American Research Team supply more accurate earnings forecasts than other analysts when forecasts are matched by the corporation followed and by the date of brokerage house issuance. This contemporaneous advantage is complemented by a timing advantage; All-Americans supply forecasts more often than other analysts. Stocks returns immediately following large upward forecast revisions suggest that All-Americans impact prices more than other analysts. However, there is virtually no difference in returns following large downward revisions. Nevertheless, the collective results suggest a positive relation between reputation and performance, and, assuming that All-Americans are better paid, pay and performance.

THIS PAPER EXAMINES THE performance of security analysts on the *Institutional Investor* (*II*) All-American Research Team relative to the performance of other analysts. Each year, for its October issue, *II* asks about 2,000 money managers to evaluate analysts on the basis of four criteria: stock picking, earnings forecasts, written reports, and overall service. The association between the results of this poll and analyst performance is important for increasing our understanding of the relation between reputation and performance and, to the extent that All-Americans are better paid, the relation between pay and performance.

Position on the All-American Research Team can be viewed as a proxy for relative reputation and compensation. A recent *Wall Street Journal* article reported that, at most brokerage houses, position on the All-American Research Team is one of the three most important criteria for determining analyst pay.¹ Conversations with research directors at major brokerage

* School of Business Administration, La Salle University. I wish to express my gratitude to the following: Larry Brown, Vic Defeo, Bob Holthausen, Don Lewin, Dave Mayers, Pat O'Brien, Jim Ohlson, K. Ramesh, René Stulz, an anonymous referee, and workshop participants at Columbia, Duke, Illinois at Chicago, Northwestern, Rutgers, Salisbury State, SUNY at Buffalo, Villanova, and Wharton for helpful comments; the KPMG Peat Marwick Foundation and Deloitte & Touche for financial support; Mark Millender for research assistance; and Zacks Investment Research, Inc. for supplying the analyst forecasts.

¹The article reported that, at *most* firms, the important factors affecting pay are an evaluation of the analyst by the brokerage sales force, standing in the *II* poll, and job offers from competitors. A smaller set of firms expand the set of factors to include investment banking business generated, trading volume in recommended stocks, and the success of buy and sell recommendations. Accuracy of earnings forecasts is rarely an *explicit* factor, but is subsumed in the other factors. As one analyst put it, "If your estimates aren't accurate, nobody's going to buy your stocks." See the "Heard on the Street" column from October 29, 1991.

houses confirmed that All-Americans are generally higher paid.² Another confirmation of the importance of the *II* poll is that analysts attempt to influence the poll by visiting money managers about the time they vote, which suggests that analyst utility, if not reputation and compensation, is affected by position on the All-American Research Team.³

While All-Americans are generally higher paid, there is some dispute over whether or not All-Americans perform any better than other analysts. An *II* editor said that the magazine has been criticized by some analysts who say the poll is simply a “beauty contest” and that All-Americans perform no better than other analysts. By examining the relative performance of All-Americans, I provide direct evidence on the relation between reputation and performance among analysts. To the extent that All-Americans are better paid, I also provide evidence on the relation between pay and performance.⁴

I measure analyst performance in three dimensions: (1) forecast accuracy, (2) frequency of forecast issuance, and (3) impact of forecast revisions on security prices. Summarizing my results, there is a positive relation between reputation and performance. All-Americans supply more accurate and more frequent earnings forecasts than other analysts.⁵ Matching forecasts by the date of brokerage house issuance, All-Americans are more accurate by an average of 2.8 cents per share for shares that average \$37 in price and \$2.81 in earnings (EPS), respectively. This contemporaneous advantage is complemented by a timing advantage. All-Americans supply forecasts more often, revising an average of every 86 calendar days, whereas NonAll-Americans revise an average of every 93 calendar days.

There is strong evidence that election to the Team follows a period of more frequent forecast issuance, relative to other NonAll-Americans, but little evidence that removal from the Team follows a period of less-frequent forecast issuance, relative to other All-Americans. I find evidence that removal from the Team follows a period of relatively poor forecast accuracy, relative to other All-Americans, but little evidence that election to the Team follows

²The research directors said that although All-Americans are generally higher paid, analyst compensation is not simply based on the four criteria used in the *II* poll. This is corroborated by the *Wall Street Journal* article and by an All-American, who (in conversation) said that a specific portion of her pay is based on her ability to generate investment banking business through finding, endorsing, sponsoring, and supporting relations between companies (e.g., stock offerings and mergers and acquisitions). Thus, by polling only institutions, *II* does not capture some important determinants of analyst reputation and compensation at some brokerage houses.

³Barbara Bent, an *II* editor, told me of this practice. This activity is confirmed by a research director quoted in the *Wall Street Journal* article as saying, “Most of the guys know that they’ll be visiting for *II* in the spring. I’m a lonely guy in March and April shortly before the balloting.”

⁴Murphy (1985) and Coughlan and Schmidt (1985) find evidence of a positive relation between pay and performance among corporate executives.

⁵Brown and Chen (1991) find that All-Americans are more accurate using a different experimental design that compares a consensus of All-American forecasts with a consensus of all sell-side analyst forecasts. On the other hand, O’Brien (1990) finds no evidence of systematic differences in individual analyst forecast accuracy.

a period of greater relative forecast accuracy, relative to other NonAll-Americans.

In terms of relative impact on security prices, abnormal returns over the first 11 days after large upward forecast revisions indicate that All-Americans have a 0.21% greater average impact on security prices than NonAll-Americans. This difference occurs in a test that controls for the magnitude of the revision and firm size. Nevertheless, the price impact evidence is mixed to the extent that the mean marginal effect of All-Americans vis-à-vis NonAll-Americans varies with the length of the cumulation period and, using an 11-day cumulation period, varies from first-team members, an incremental 0.65%, to second-team members, an incremental 0.09%, to third-team members, an incremental -0.03% , to runners-up, an incremental 0.27%. Moreover, there is virtually no difference between the impact of large downward revisions made by All-Americans and NonAll-Americans. As expected, the impact of revisions on prices is generally positively related to the change in forecast and is generally greater for smaller firms.

Section I states the forecast accuracy hypotheses, details the test method, and reports the results. Frequency of forecast issuance tests are reported in Section II. Tests of security price impact are in Section III. Section IV concludes this study.

I. Forecast Accuracy

A. Hypotheses and Experiment Design for Relative Forecast Accuracy

Earnings forecast accuracy is one measurable performance characteristic that establishes analyst reputation. My three forecast accuracy hypotheses are listed in order of the life cycle of an analyst's All-American status, from admission to removal.

Hypothesis 1: *NonAll-Americans who are about to become Team members forecast earnings more accurately than other NonAll-Americans.*

Hypothesis 2: *All-Americans are more accurate forecasters than NonAll-Americans.*

Hypothesis 3: *All-Americans who are about to lose Team membership forecast earnings less accurately than other All-Americans.*

Forecast errors are defined as actual annual EPS minus forecasted EPS

$$E_{a,i,y(h)} = A_{i,y} - F_{a,i,y(h)} \quad (1)$$

where

$E_{a,i,y(h)}$ = forecast error by analyst a , forecasting EPS for firm i for fiscal year y , where the forecast horizon is h days prior to the fiscal year-end

$A_{i,y}$ = actual EPS for firm i for fiscal year y

$F_{a,i,y(h)}$ = forecast by analyst a , forecasting EPS for firm i for fiscal year y , where the forecast horizon is h days prior to the fiscal year-end.

The hypothesis that greater relative accuracy leads to All-American status, Hypothesis 1, is tested by comparing the past (fiscal year $y - 3$, $y - 2$, and $y - 1$) accuracy of past (calendar year $y - 3$, $y - 2$, and $y - 1$) NonAll-Americans who become first-time All-Americans in calendar year y with that of past NonAll-Americans who remain NonAll-Americans in calendar year y .⁶ Specifically, a forecast made h days prior to the end of fiscal year $y - v$ ($v = 1, 2, 3$) by a first-time All-American in calendar year y is matched with a forecast made h days prior to the same fiscal year-end by a NonAll-American in calendar years $y - v$ through y .⁷ Because analysts become more accurate as the announcement date approaches (see Crichfield, Dyckman, and Lakonishok (1978)), the accuracy of All-Americans and NonAll-Americans is compared by matching forecasts of *identical* horizons.⁸ This design allows a determination of whether or not All-Americans possess a contemporaneous advantage. A separate issue, investigated below, is whether or not All-Americans offer a timing advantage by issuing forecasts more often than NonAll-Americans.

Unsigned (absolute) forecast errors (measures of accuracy) are compared as

$$|E_{a(Ay),i,y-v(h)}| - |E_{a(Ny),i,y-v(h)}| \quad (2)$$

where $a(A_y)$ represents a first-time All-American analyst in calendar year y and $a(N_y)$ represents a NonAll-American analyst. Results are reported for unscaled absolute forecast errors and for absolute forecast errors that are scaled by price per share at fiscal year-end.⁹

⁶*II* is reluctant to reveal when the voting occurs because some analysts attempt to enhance their prospects by contacting money managers about the time they vote. An *II* editor did say that voting occurs within April–June. Chambers and Penman (1984) find that about 97% of annual EPS announcements are made within 12 weeks of fiscal year-end (annual 10Ks must be filed with the SEC within 90 days). Thus, in virtually all cases, forecast accuracy for fiscal year $y - 1$ is known before voting occurs for the calendar year y Team.

⁷Over 1979–88, *II* increased the number of industries and investment specialties covered in the poll from 48 to 65. Thus, a consistently superior analyst over fiscal years $y - 3$ through $y - 1$ may appear on the calendar year y Team because of a new category (e.g., biotechnology) rather than an *improvement* in relative forecasting ability over years $y - 3$ through $y - 1$. Analysts admitted to the Team because of a new category are eliminated from tests of Hypothesis 1.

⁸Zacks dates forecasts by the date of issuance by the brokerage house employing the analyst. This dating convention contrasts with that used by I/B/E/S, which dates forecasts by the date of receipt by I/B/E/S. Requiring the same calendar date of issuance for matched forecasts results in discarding approximately 88% of all forecasts, but more importantly reduces the potential of NonAll-Americans (All-Americans) mimicking All-American (NonAll-American) forecast revisions.

⁹Conclusions are generally not changed if absolute forecast errors are scaled by EPS. Scaling factors less than \$0.25, including any negative numbers, are arbitrarily set equal to \$0.25 to mitigate “small denominators” and scaled forecast errors are truncated at -200% and $+200\%$ to avoid giving undue weight to outliers.

Accuracy differences are not identically distributed because analyst forecasts become more accurate as the earnings announcement date approaches. To mitigate this problem, the sample is segregated into 60 subsamples on the basis of the calendar month in which the forecast date falls and mean results are calculated by subsample as

$$\text{MDE}_m = \sum^{\text{all } N_m} [|E_{a(Ay), i, y-v(h)}| - |E_{a(Ny), i, y-v(h)}|] / N_m \quad (3)$$

where MDE_m is the mean difference in absolute forecast errors and N_m is the number of matched errors within monthly period m . Subsample means are averaged as

$$\text{MDE} = \sum_{m=1}^{60} \text{MDE}_m / 60. \quad (4)$$

The significance of this overall mean is determined by

$$t = \text{MDE} / [\sigma_{\text{MDE}_m} / \sqrt{60}] \quad (5)$$

where σ_{MDE_m} is the estimated standard deviation of the monthly MDE_m . This test statistic, which is assumed to be distributed Student t with 59 degrees of freedom, uses monthly *mean* differences, which, from the Central Limit Theorem, are normally distributed. This experimental design is analogous to that in Jaffe (1974) and Mandelker (1974).¹⁰ Results are also reported for a nonparametric binomial test that uses monthly mean differences in paired forecast errors. Hypothesis 1 predicts a negative MDE, indicating greater relative accuracy precedes election to the Team.

The hypothesis that All-Americans are more accurate forecasters than other analysts, Hypothesis 2, is tested by comparing the current (fiscal year y) accuracy of current (calendar year y) All-Americans with that of current (calendar year y) NonAll-Americans. Specifically, a forecast made h days prior to the year-end of fiscal year y by an All-American in calendar year y is matched with a forecast of identical horizon made by a NonAll-American in calendar year y . Using equation (2), v is set equal to zero to test Hypothesis 2. Hypothesis 2 predicts a negative MDE.

The hypothesis that poor relative accuracy precedes removal from the Team, Hypothesis 3, is tested by comparing the past (fiscal year $y - 3$, $y - 2$, and $y - 1$) accuracy of past (calendar year $y - 3$, $y - 2$, and $y - 1$) All-Americans who become NonAll-Americans in calendar year y with the past accuracy of past All-Americans who remain All-Americans in calendar year y . A forecast made h days prior to the fiscal year-end of year $y - v$ ($v = 1, 2, 3$) by a first-time NonAll-American in calendar year y (an All-American in years $y - v$ through $y - 1$) is matched with a forecast made h days prior to the same fiscal year-end by an All-American in calendar years $y - v$ through y .

¹⁰ Forecast errors are cross-sectionally dependent due to economy-wide and industry-wide factors. However, there is no reason to suspect that *differences* in errors are cross-sectionally dependent. Paired analysts are limited to one inclusion per firm per fiscal year to avoid cross-sectional dependence.

Equation (2) is used, but now year y is defined as the year of removal from the Team rather than the year of admission. Hypothesis 3 predicts a negative MDE, indicating a deterioration in the relative accuracy of All-Americans who become NonAll-Americans in year y . Unfortunately, Hypothesis 3 is confounded to the extent that All-Americans may leave the Team for reasons unrelated to poor performance.¹¹

B. Empirical Results for Forecast Accuracy

Individual analyst forecasts of annual EPS are supplied by Zacks Investment Research (Zacks) for firms with fiscal year-ends within 1981–85.¹² The names of members of the All-American Research Team from 1979–88 are collected from back issues of *II*. A cursory review of *II* suggests that most All-Americans remain All-Americans for more than one year. In addition to covering industry categories, *II* has All-Americans for the categories of portfolio strategy, quantitative analysis, economics, and market timing. Analysts in these investment specialty categories generally do not make forecasts for specific firms.

For tests of the forecast accuracy Hypotheses 1 through 3, forecasts that are dated within 1981–85 and within the current fiscal year of the firm are used. Actual EPS are obtained from the 1988 COMPUSTAT Annual Industrial File, Full Coverage File, and Research File. Table I reports the number of All-Americans, the number of matched forecasts, and the number of firms for the forecast accuracy and forecast frequency tests. Testing using larger sample sizes (i.e., Hypotheses 2 and 5) are more powerful (a greater chance of rejecting the null when it is false), *ceteris paribus*.

Table II reports results for tests of Hypothesis 1, that greater relative accuracy leads to All-American status. Forecast accuracy is reported for fiscal years $y - 3$, $y - 2$, and $y - 1$ for NonAll-Americans who become first-time All-Americans in calendar year y along with that of past NonAll-Americans who remain NonAll-Americans in calendar year y . Although differences in forecast accuracy are generally negative as predicted, the evidence is not statistically significant at conventional levels.

Table III reports the results of the test of Hypothesis 2, that All-American analysts are more accurate forecasters than other analysts. Forecast bias (signed forecast errors) and accuracy are reported for fiscal year y for All-Americans and NonAll-Americans of calendar year y . The negative

¹¹ Within 1979–88, I found 15 cases in which *II* noted that an All-American left the Team for a reason unrelated to poor performance (e.g., became a research director at a brokerage house, became a money manager, or died). These analysts are excluded from the test of Hypothesis 3. However, this does not ensure the exclusion of *all* analysts who have left the Team for reasons unrelated to poor performance.

¹² References to EPS are to primary EPS before extraordinary items and discontinued operations. However, Zacks varies from the APBO #30 definition to coincide with, in their view, Wall Street's perception of earnings, which is similar to "permanent earnings." This increases measured mean absolute forecast error, but affects All-Americans and NonAll-Americans equally. EPS figures are adjusted for any changes in shares outstanding.

Table I
Sample Characteristics for the Forecast Accuracy and Forecast Frequency Hypotheses

Hypothesis 1: NonAll-Americans who are about to become Team members forecast earnings more accurately than other NonAll-Americans. Year y is the year of admission to the Team. Hypothesis 2: All-Americans are more accurate forecasters than NonAll-Americans. Year y represents any year an analyst is on the Team. Hypothesis 3: All-Americans who are about to lose Team membership forecast earnings less accurately than other All-Americans. Year y is the year of removal from the Team. Hypothesis 4: NonAll-Americans who are about to become Team members forecast earnings more frequently than other NonAll-Americans. Year y is the year of admission to the Team. Hypothesis 5: All-Americans are more frequent forecasters than NonAll-Americans. Year y represents any year an analyst is on the Team. Hypothesis 6: All-Americans who are about to lose Team membership forecast earnings less frequently than other All-Americans. Year y is the year of removal from the Team.

	Forecast Year				
	1981	1982	1983	1984	1985
Hypotheses 1 and 4: Using forecasts for fiscal year $y - 3$					
Number of year y first-time All-Americans	18	23	16	12	15
Number of matched forecasts	70	125	43	63	64
Number of firms	45	92	33	48	37
Hypotheses 1 and 4: Using forecasts for fiscal year $y - 2$					
Number of year y first-time All-Americans	12	24	31	24	17
Number of matched forecasts	69	150	146	102	88
Number of firms	44	103	99	58	63
Hypotheses 1 and 4: Using forecasts for fiscal year $y - 1$					
Number of year y first-time All-Americans	15	23	43	53	36
Number of matched forecasts	76	152	277	315	180
Number of firms	47	99	149	182	93
Hypotheses 2 and 5: Using forecasts for fiscal year y					
Number of year y All-Americans	195	229	252	283	314
Number of matched forecasts	1009	1517	1393	1859	1866
Number of firms	345	428	436	515	538
Hypotheses 3 and 6: Using forecasts for fiscal year $y - 3$					
Number of year y first-time NonAll-Americans	24	31	33	39	45
Number of matched forecasts	37	40	55	76	101
Number of firms	30	28	47	49	64
Hypotheses 3 and 6: Using forecasts for fiscal year $y - 2$					
Number of year y first-time NonAll-Americans	26	35	19	77	60
Number of matched forecasts	38	66	28	147	95
Number of firms	32	48	21	102	71
Hypotheses 3 and 6: Using forecasts for fiscal year $y - 1$					
Number of year y first-time NonAll-Americans	48	39	37	35	102
Number of matched forecasts	83	63	70	77	248
Number of firms	51	48	50	52	134

Table II

Hypothesis 1: Does Greater Relative Accuracy Precede All-American Status?

Absolute forecast errors (abs[actual EPS – forecasted EPS]) for fiscal year $y - 3$, $y - 2$, and $y - 1$ made by calendar year y first-time All-Americans and calendar year $y - 3$, $y - 2$, $y - 1$, and y NonAll-Americans. Forecasts are dated within the firm's current fiscal year and within calendar 1981–85 and fiscal year-ends are within calendar 1981–85. Forecasts are matched by date of brokerage house issuance.

Panel A: Forecast Accuracy									
Fiscal Year	Scaling Factor	Forecaster	Distribution of Absolute Forecast Errors						
			N	Mean	Min	1st Qrt	Median	3rd Qrt	Max
y - 3	Unscaled	Year y first-time All-Americans	365	0.9445	0.0000	0.0900	0.3600	1.038	13.24
		NonAll-Americans in years y - 3 through y	365	0.9590	0.0000	0.0925	0.3600	0.9700	15.49
y - 3	Price	Year y first-time All-Americans	365	0.0345	0.0000	0.0026	0.0111	0.0373	1.351
		NonAll-Americans in years y - 3 through y	365	0.0349	0.0000	0.0027	0.0107	0.0371	1.581
y - 2	Unscaled	Year y first-time All-Americans	555	0.9365	0.0000	0.0900	0.3100	0.8600	29.56
		NonAll-Americans in years y - 3 through y	555	0.9486	0.0000	0.1000	0.3000	0.8700	28.30
y - 2	Price	Year y first-time All-Americans	555	0.0319	0.0000	0.0028	0.0091	0.0249	1.060
		NonAll-Americans in years y - 3 through y	555	0.0322	0.0000	0.0033	0.0094	0.0244	1.070
y - 1	Unscaled	Year y first-time All-Americans	1000	0.9677	0.0000	0.1000	0.3000	0.8500	44.36
		NonAll-Americans in years y - 3 through y	1000	0.9678	0.0000	0.0900	0.2800	0.8400	43.52
y - 1	Price	Year y first-time All-Americans	1000	0.0309	0.0000	0.0030	0.0092	0.0268	0.8613
		NonAll-Americans in years y - 3 through y	1000	0.0312	0.0000	0.0027	0.0093	0.0273	0.8450

Panel B: Matched Pairs Differences in Forecast Accuracy														
Fiscal Year	Scaling Factor	Distribution of Differences in Matched Pairs Absolute Forecast Errors						Monthly Mean Differences						
		N	Mean	Min	1st Qrt	Median	3rd Qrt	Max	% < 0	% > 0	Mean ^a	% < 0 ^a	% > 0	
y - 3	Unscaled	365	-0.0145	-3.070	-0.1000	0.0000	0.1400	2.000	44.4	12.6	43.0	-0.0015	59.6	40.4
y - 3	Price	365	-0.0004	-0.2296	-0.0035	0.0000	0.0038	0.1094	same as for unscaled			0.0002	50.9	49.1
y - 2	Unscaled	555	-0.0121	-7.590	-0.1200	0.0000	0.1000	3.000	46.8	10.3	42.9	-0.0027	55.0	45.0
y - 2	Price	555	-0.0003	-0.3814	-0.0035	0.0000	0.0030	0.1252	same as for unscaled			-0.0001	51.7	48.3
y - 1	Unscaled	1000	-0.0001	-3.105	-0.1200	0.0000	0.1000	4.250	44.2	12.0	43.8	0.0077	55.0	45.0
y - 1	Price	1000	-0.0004	-0.1744	-0.0039	0.0000	0.0035	0.2038	same as for unscaled			-0.0001	58.3	41.7

Table III

Hypothesis 2: Are All-Americans More Accurate Forecasters?

Forecast errors (actual EPS – forecasted EPS) and absolute errors (abs[actual EPS – forecasted EPS]) for fiscal year *y* made by calendar year *y* All-Americans and NonAll-Americans. Forecasts are dated within the firm's current fiscal year and within calendar 1981–85 and fiscal year-ends are within calendar 1981–85. Forecasts are matched by date of brokerage house issuance.

Panel A: Forecast Bias														
Fiscal Year	Scaling Factor	Forecaster	N	Distribution of Forecast Errors										
				Mean	Min	1st Qrt	Median	3rd Qrt	Max					
y	Unscaled	All-Americans	7644	− 0.7278	− 44.88	− 0.7600	− 0.1400	0.0500	14.46					
		NonAll-Americans	7644	− 0.7421	− 47.20	− 0.7748	− 0.1500	0.0600	11.57					
y	Price	All-Americans	7644	− 0.0292	− 2.000	− 0.0269	− 0.0044	0.0014	2.000					
		NonAll-Americans	7644	− 0.0298	− 2.000	− 0.0280	− 0.0048	0.0015	2.000					
Panel B: Forecast Accuracy														
Fiscal Year	Scaling Factor	Forecaster	N	Distribution of Absolute Forecast Errors										
				Mean	Min	1st Qrt	Median	3rd Qrt	Max					
y	Unscaled	All-Americans	7644	0.9536	0.0000	0.1000	0.3000	0.9200	44.88					
		NonAll-Americans	7644	0.9818	0.0000	0.1000	0.3200	0.9400	47.20					
y	Price	All-Americans	7644	0.0364	0.0000	0.0028	0.0095	0.0311	2.000					
		NonAll-Americans	7644	0.0376	0.0000	0.0031	0.0099	0.0319	2.000					
Panel C: Matched Pairs Differences in Forecast Accuracy														
Fiscal Year	Scaling Factor	Distribution of Differences in Matched Pairs Absolute Forecast Errors						Monthly Mean Differences						
		N	Mean	Min	1st Qrt	Median	3rd Qrt	Max	% < 0	% > 0				
y	Unscaled	7644	− 0.0282	− 10.17	− 0.1200	0.0000	0.1000	7.740	46.9	12.3	40.8	− 0.0272***	71.7***	28.3
y	Price	7644	− 0.0013	− 0.8895	− 0.0040	0.0000	0.0028	0.4651	same as for unscaled	− 0.0012***	70.0***	30.0		

^aThree asterisks indicate significance at the 0.01 level.

forecast errors suggest that analysts overestimated EPS over this period.¹³ There is strong evidence that All-Americans are more accurate than NonAll-Americans.

Table IV reports the results of tests of Hypothesis 3, that lower relative accuracy precedes removal from the Team. Forecast accuracy is reported for past All-Americans who become NonAll-Americans in calendar year y and past All-Americans who remain All-Americans in calendar year y . Differences in forecast accuracy are generally negative as predicted, and, in some cases, statistically significant at conventional levels.

II. Frequency of Forecast Issuance

A. Hypotheses and Experiment Design for Relative Frequency of Forecasts

If All-Americans are more accurate forecasters, holding the forecast release date constant, then, since analyst forecasts become more accurate as the earnings announcement date approaches, *ceteris paribus*, more frequent forecasts are advantageous. More frequent forecasts are more advantageous because they can incorporate the latest earnings-relevant information (e.g., recent interim earnings announcements and industry reports). Forecast frequency also proxies for the third and fourth criteria polled: written reports and overall service. My three frequency of forecasts hypotheses are analogous to my three forecast accuracy hypotheses, and are listed in order of the life cycle of an analyst's All-American status, from admission to removal.

Hypothesis 4: *NonAll-Americans who are about to become Team members issue earnings forecasts more frequently than other NonAll-Americans.*

Hypothesis 5: *All-Americans revise their forecasts more frequently than Non-All-Americans.*

Hypothesis 6: *All-Americans who are about to lose Team membership issue earnings forecasts less frequently than other All-Americans.*

These hypotheses are tested by comparing the number of days elapsed between forecasts

$$\text{DYS}_{a(Ay), i, y-v(h)} - \text{DYS}_{a(Ny), i, y-v(h)}, \quad (6)$$

where $\text{DYS}_{a, i, y-v(h)}$ is the number of calendar days prior to the issuance of $F_{a, i, y-v(h)}$ that another forecast was issued by the same analyst for the same firm for the same fiscal year. The sample is again segregated into 60

¹³A similar upward bias over the same period has been documented for Value Line Investment Survey forecasts by Abarbanell (1989) and in my own unpublished work using I/B/E/S mean consensus forecasts. The median forecast errors are also negative, but are approximately one-third the magnitude of the mean forecast errors.

Table IV

Hypothesis 3: Does Lower Relative Accuracy Precede Removal From the Team?

Absolute forecast errors (abslactual EPS – forecasted EPS) for fiscal year $y - 3$, $y - 2$, and $y - 1$ made by calendar year $y - 3$, $y - 2$, and $y - 1$ All-Americans who are calendar year y NonAll-Americans and calendar year $y - 3$, $y - 2$, $y - 1$, and y All-Americans. Forecasts are dated within the firm's current fiscal year and within calendar 1981–85 and fiscal year-ends are within calendar 1981–85. Forecasts are matched by date of brokerage house issuance.

Panel A: Forecast Accuracy												
Fiscal Year	Scaling Factor	Forecaster	Distribution of Absolute Forecast Errors						N	Mean	Min	Max
			1st Qrt	Median	3rd Qrt	Max						
y - 3	Unscaled	All-Americans in years y - 3 through y	0.0900	0.2450	0.8450	30.45	0.8732	0.0000	309	0.8732	0.0000	30.45
		Year y first-time NonAll-Americans	0.0910	0.2500	0.8400	30.79	0.9169	0.0000	309	0.9169	0.0000	30.79
y - 3	Price	All-Americans in years y - 3 through y	0.0026	0.0087	0.0260	0.5399	0.0265	0.0000	309	0.0265	0.0000	0.5399
		Year y first-time NonAll-Americans	0.0029	0.0089	0.0261	0.5459	0.0279	0.0000	309	0.0279	0.0000	0.5459
y - 2	Unscaled	All-Americans in years y - 3 through y	0.1300	0.3600	0.9325	42.78	0.8779	0.0000	374	0.8779	0.0000	42.78
		Year y first-time NonAll-Americans	0.1100	0.3500	1.080	44.88	0.9042	0.0000	374	0.9042	0.0000	44.88
y - 2	Price	All-Americans in years y - 3 through y	0.0036	0.0108	0.0333	1.183	0.0332	0.0000	374	0.0332	0.0000	1.183
		Year y first-time NonAll-Americans	0.0033	0.0109	0.0318	0.8754	0.0327	0.0000	374	0.0327	0.0000	0.8754
y - 1	Unscaled	All-Americans in years y - 3 through y	0.1300	0.4000	1.1500	13.85	1.008	0.0000	541	1.008	0.0000	13.85
		Year y first-time NonAll-Americans	0.1200	0.4100	1.1550	14.88	1.027	0.0000	541	1.027	0.0000	14.88
y - 1	Price	All-Americans in years y - 3 through y	0.0037	0.0135	0.0379	0.9814	0.0370	0.0000	541	0.0370	0.0000	0.9814
		Year y first-time NonAll-Americans	0.0034	0.0132	0.0388	0.9563	0.0375	0.0000	541	0.0375	0.0000	0.9563

Panel B: Matched Pairs Differences in Forecast Accuracy													
Fiscal Year	Scaling Factor	Distribution of Differences in Matched Pairs Absolute Forecast Errors						Monthly Mean Differences					
		1st Qrt	Median	3rd Qrt	Max	% < 0	% = 0	% > 0	Mean ^a	% < 0 ^a	% > 0 ^a		
y - 3	Unscaled	-0.0436	-6.000	-0.1200	0.0000	0.1000	2.900	44.7	13.3	42.1	-0.0225	52.6	47.4
y - 3	Price	-0.0014	-0.1875	-0.0034	0.0000	0.0024	0.1695	same as for unscaled			-0.0011	61.4*	38.6
y - 2	Unscaled	-0.0263	-2.097	-0.1500	0.0000	0.1000	2.000	45.5	11.8	42.7	-0.0348	64.3**	35.7
y - 2	Price	0.0004	-0.1080	-0.0040	0.0000	0.0032	0.3077	same as for unscaled			0.0014	67.9***	32.1
y - 1	Unscaled	-0.0196	-3.560	-0.1200	0.0000	0.1000	3.990	42.3	13.9	43.8	-0.0431*	58.3	41.7
y - 1	Price	-0.0005	-0.1250	-0.0039	0.0000	0.0034	0.1327	same as for unscaled			-0.0008	58.3	41.7

^a One, two, and three asterisks indicate significance at the 0.10, 0.05, and 0.01 levels, respectively.

subsamples on the basis of the calendar month in which the forecast date falls and mean results are calculated by subsample as

$$\text{MDDYS}_m = \sum^{\text{all } N_m} [|\text{DYS}_{a(Ay), i, y-v(h)}| - |\text{DYS}_{a(Ny), i, y-v(h)}|] / N_m \quad (7)$$

where MDDYS_m is the mean difference in number of days between forecasts for subsample m . Subsample means are averaged and a test of statistical significance is performed in a manner analogous to that used for differences in absolute forecast errors in equations (4) and (5). Hypotheses 4, 5, and 6 all predict a negative MDDYS.

B. Empirical Results for Forecast Frequency

For tests of Hypotheses 4, 5, and 6, the matched forecasts from the tests of the accuracy Hypotheses 1, 2, and 3, are used. Table V reports results for tests of Hypothesis 4, that more frequent forecasts lead to Team membership. Forecast frequency is reported for fiscal years $y - 3$, $y - 2$, and $y - 1$ for NonAll-Americans who become first-time All-Americans in calendar year y along with that of past NonAll-Americans who remain NonAll-Americans in calendar year y . There is strong evidence that future All-Americans revise forecasts more frequently in the two years prior to admission to the Team.

Table VI reports results that strongly support Hypothesis 5, that All-Americans supply forecasts more frequently than NonAll-Americans. All-Americans revise earnings an average of every 86 calendar days, whereas NonAll-Americans revise an average of every 93 calendar days.¹⁴

A regression of absolute forecast errors on the forecast horizon produces an intercept of \$0.50 and a slope of \$0.0025 using All-American forecasts and an intercept of \$0.51 and a slope of \$0.0026 using NonAll-American forecasts. Thus, an approximation of the marginal improvement in forecast accuracy attributable to a 7-day mean difference in forecast frequency is 1.75 cents per share (7 days times \$0.0025 per day).¹⁵

Table VII reports results for Hypothesis 6, that less frequent forecasts precede removal from the Team. In the year preceding removal from the Team, All-Americans who are in their final year revise earnings an average of every 96 calendar days, whereas All-Americans who remain on the Team revise an average of every 84 days. However, both a parametric test and a nonparametric test are unable to reject the null of no difference at conventional levels.

¹⁴On average, All-Americans forecast earnings for 15.3 firms per year and NonAll-Americans follow 10.5 firms. The median numbers of firms followed are 14 and 8 for All-Americans and NonAll-Americans, respectively.

¹⁵This measure of marginal improvement assumes that investors who use an All-American forecast do not also use NonAll-American forecasts issued within 7 days after the All-American forecast. Such NonAll-American forecasts could incorporate the All-American forecast, negating the timing advantage. This assumption is consistent with the results in Stickel (1991) (see footnotes 15, 18, and 21).

Table V
Hypothesis 4: Does Higher Relative Frequency of Forecast Issuance Precede Admission to the Team?

Number of calendar days elapsed since the date of the prior forecast for forecast revisions dated within a firm's current fiscal year. Forecasts are matched by date of brokerage house issuance. This is the same sample used to test Hypothesis 1.

Panel A: Number of Calendar Days Elapsed between Forecasts													
Fiscal Year	Forecaster	Distribution of Number of Days Elapsed between Forecasts											
		N	Mean	Min	1st Qrt	Median	3rd Qrt	Max					
y - 3	Year y first-time All-Americans	365	83.72	1	30	61	102	553					
	NonAll-Americans in years y - 3 through y	365	93.11	1	33	68	120	563					
y - 2	Year y first-time All-Americans	555	80.02	2	30	60	98	518					
	NonAll-Americans in years y - 3 through y	555	93.47	1	31	63	120	547					
y - 1	Year y first-time All-Americans	1000	84.89	1	30	61	108	660					
	NonAll-Americans in years y - 3 through y	1000	93.78	1	32	65	122	538					
Panel B: Matched Pairs Differences in Number of Calendar Days Elapsed between Forecasts													
Fiscal Year	Distribution of Differences in Matched Pairs Number of Days Elapsed between Forecasts								Monthly Mean Differences				
	N	Mean	Min	1st Qrt	Median	3rd Qrt	Max	% < 0	% = 0	% > 0	Mean ^a	% < 0 ^a	% > 0 ^a
y - 3	365	-9.38	-502	-62	-2	36	472	51.2	1.6	47.1	-5.84	56.1	43.9
y - 2	555	-13.44	-527	-56	-5	32	352	52.8	1.3	45.9	-17.45***	60.0	40.0
y - 1	1000	-8.89	-482	-60	-7	36	642	55.5	1.5	43.0	-7.91*	61.7	38.3

^aOne and three asterisks indicate significance at the 0.10 and 0.01 levels, respectively.

Table VI
Hypothesis 5: Do All-Americans Supply Forecasts More Frequently Than NonAll-Americans?

Number of calendar days elapsed since the date of the prior forecast for forecast revisions dated within a firm's current fiscal year. Forecasts are matched by date of brokerage house issuance. This is the same sample used to test Hypothesis 2.

Panel A: Number of Calendar Days Elapsed between Forecasts										
Fiscal Year	Forecaster	N	Distribution of Number of Days Elapsed							
			Mean	Min	1st Qrt	Median	3rd Qrt	Max		
y	All-Americans	7644	85.83	1	31	61	107	730		
	NonAll-Americans	7644	93.25	1	32	63	120	671		
Panel B: Matched Pairs Differences in Number of Calendar Days Elapsed between Forecasts										
Fiscal Year	Distribution of Differences in Matched Pairs Number of Days Elapsed between Forecasts							Monthly Mean Differences		
	N	Mean	Min	1st Qrt	Median	3rd Qrt	Max	% < 0	% = 0	% > 0
y	7644	-7.42	-585	-56	-2	38	674	51.1	1.9	47.0
								-7.73***	73.3***	26.7

^aThree asterisks indicate significance at the 0.01 level.

Table VII
Hypothesis 6: Does Lower Relative Frequency of Forecast Issuance Precede Removal From the Team?

Number of calendar days elapsed since the date of the prior forecast for forecast revisions dated within a firm's current fiscal year. Forecasts are matched by date of brokerage house issuance. This is the same sample used to test Hypothesis 3.

Panel A: Number of Calendar Days Elapsed between Forecasts											
Fiscal Year	Forecaster	Distribution of Number of Days Elapsed between Forecasts								N	Max
		1st Qrt	Median	3rd Qrt	Mean	Min	1st Qrt	Median	3rd Qrt		
y - 3	All-Americans in years y - 3 through y				84.69	2	30	60	107	309	631
	Year y first-time NonAll-Americans				80.53	1	30	61	98	309	538
y - 2	All-Americans in years y - 3 through y				87.69	2	34	62	103	374	577
	Year y first-time NonAll-Americans				91.44	1	33	63	122	374	546
y - 1	All-Americans in years y - 3 through y				83.90	1	31	59	94	541	549
	Year y first-time NonAll-Americans				96.08	1	34	70	125	541	609

Panel B: Matched Pairs Differences in Number of Calendar Days Elapsed between Forecasts											
Fiscal Year	N	Distribution of Differences in Matched Pairs Number of Days Elapsed between Forecasts						Monthly Mean Differences			
		Mean	1st Qrt	Median	3rd Qrt	Max	% < 0	% = 0	% > 0	Mean ^a	% < 0 ^a % > 0
y - 3	309	4.16	-378	-45	0	600	49.8	0.6	49.5	7.66	42.1 57.9
y - 2	374	-3.75	-516	-54	1	408	48.7	1.1	50.3	-2.74	48.2 51.8
y - 1	541	-12.19	-442	-60	-7	475	54.5	2.2	43.3	-9.00	56.7 43.3

^aThe monthly mean differences are *not* significantly different from zero at conventional levels.

III. Impact of Forecast Revisions on Security Prices

A. Hypothesis and Experiment Design for Relative Impact on Security Prices

The impact of forecast revisions on stock prices is another measurable performance characteristic that establishes analyst reputation. My security price impact hypothesis is

Hypothesis 7: Forecast revisions made by All-Americans have a greater impact on security prices than revisions made by NonAll-Americans.

Because Imhoff and Lobo (1984) and Stickel (1991) find mean abnormal returns following forecast revisions are greatest for large revisions that are scaled by the standard deviation of outstanding forecasts, this study measures information content by the scaled change in forecast

$$\text{SUF}_{a,i,t} = (F_{a,i,t} - F_{a,i,t-d}) / \sigma_{i,t-1}. \quad (8)$$

$\text{SUF}_{a,i,t}$ is defined as the standardized unexpected forecast made by analyst a for firm i on day t , $F_{a,i,t}$ is analyst a 's EPS forecast for company i on day t , a day on which the forecast is revised, $F_{a,i,t-d}$ is analyst a 's prior EPS forecast, a forecast which is dated d days prior to day t , and $\sigma_{i,t-1}$ is the cross-sectional standard deviation of all outstanding forecasts for firm i on day $t - 1$.^{16,17}

Expected returns are estimated from the market model (e.g., Fama (1976)). Event day 0 is defined as the date of the revision. Returns from event (business) days +251 to +350 are used to obtain ordinary least squares estimates of model parameters with a minimum of 30 returns required for parameter estimation.¹⁸ The period preceding the revision is rejected as the estimation period because analysts are likely to respond to the information in past returns when making a revision (see Copeland and Mayers (1982)). The year immediately following the revision is rejected because of evidence that prices are slow to assimilate revision information (e.g., Givoly and Lakonishok (1979), Brown, Foster, and Noreen (1985), and Stickel (1991)).

The deviation of actual return from expectations is

$$\text{AR}_{it} = R_{it} - (\hat{\alpha}_i + \hat{\beta}_i R_{mt}) \quad (9)$$

¹⁶Again, a divisor less than \$0.25 is arbitrarily set equal to \$0.25 to mitigate small denominators and scaled forecast revisions are truncated at -200% and +200%. Scaling forecast revisions by the absolute value of the prior forecast or by price per share changes no conclusions.

¹⁷Stickel (1991) finds that market reaction to forecast revisions is more closely associated with the deviation of an analyst's new forecast from the analyst's old forecast than the deviation of the new forecast from an updated version of the analyst's old forecast.

¹⁸About 2% of the sample did not meet the 30 returns minimum. For these revisions, market-adjusted returns ($R_{it} - R_{mt}$) are used. The conclusions do not change if these observations are excluded from the analysis.

where AR_{it} is defined as the abnormal return for security i on event day t and R_{mt} is the equally weighted market return.¹⁹ Isolating the impact of forecast revisions is difficult because forecasts are sequentially disseminated to market participants by brokerage house sales personnel rather than immediately disseminated by major news services, such as the Broad Tape. Stickel (1991) finds a six-month drift in security prices following large forecast revisions.²⁰ Lengthening the cumulation period for a test of differential market reaction involves a tradeoff between the increase in power derived from a more complete reflection of the mean price effect, and the decrease in power caused by a larger variance in cumulative abnormal performance. For the tests of significance reported below, the primary measure of the impact of forecast revisions on security prices is $CAR_{i,(0,+10)}$, cumulative abnormal returns from day 0 through day 10.²¹ Because this choice is somewhat arbitrary, the sensitivity of the results to the length of the cumulation period is also examined.

The hypothesis that forecast revisions made by All-Americans have a greater impact on security prices, Hypothesis 7, is tested by cross-sectional regressions of $CAR_{i,(0,+10)}$ on the magnitude of the revision (see Beaver, Clarke, and Wright (1979)), firm size (see Atiase (1985)), and dummy variables indicating All-American status. The following three regression equations are used,

$$CAR_{i,(0,+10)} = \beta_0 + \beta_1 SUF_{a,i,0} + \beta_2 LSIZE_{i,-1} + \beta_3 DA_a + \varepsilon \quad (10)$$

¹⁹The sensitivity of the results to the measure of abnormal return is examined by calculating abnormal returns as $AR_{it} = (R_{it} - R_{mt}) - MMAR$, where $MMAR$ is the mean market-adjusted return over days +251 to +350 (see Penman (1987)). The sensitivity of the results to the choice of benchmark period is examined by estimating market model parameters over event days -400 to -301. The results using these alternative abnormal performance measures did not change any conclusions and are not reported.

²⁰Briefly, the major findings in Stickel (1991) are (1) there is a six-month post-announcement drift following revisions, 25% of which occurs within the first 11 days after the revision date, (2) there is an effect of revisions on prices that is independent of earnings announcements, (3) price reaction is more closely associated with a traditional measure of information content (equation (8) in this study) than an updated measure that incorporates other analyst forecasts released since an analyst's prior forecast into the expectation of the analyst's next forecast, (4) a trading strategy that incorporates the results in (3) produces a difference in abnormal returns between securities predicted to perform the best and the worst of approximately 13% over the six-month period following forecast revisions, and (5) abnormal performance is insignificantly different from zero in the second six-month period following forecast revisions, which is indicative of an unbiased performance measure.

²¹Stickel (1989) finds that 2.63% of all analysts revise annual earnings forecasts on interim earnings announcement dates. This percentage is significantly greater than a nonevent period benchmark percentage and indicates that some analysts quickly respond to earnings announcements, confounding measurement of market reaction to revisions. Nevertheless, this percentage is small and tests of the *difference* in market reaction between All-Americans and NonAll-Americans are unbiased. The results of tests using cumulative abnormal returns over event days +1 to +10 change no conclusions.

$$\begin{aligned} \text{CAR}_{i,(0,+10)} = & \beta_0 + \beta_1 \text{SUF}_{a,i,0} + \beta_2 \text{LSIZE}_{i,-1} \\ & + \beta_3 D1_a + \beta_4 D2_a + \beta_5 D3_a + \beta_6 D4_a + \varepsilon \end{aligned} \quad (11)$$

$$\begin{aligned} \text{CAR}_{i,(0,+10)} = & \beta_0 + \beta_1 \text{DA}_a^* \text{SUF}_{a,i,0} + \beta_2 \text{DN}_a^* \text{SUF}_{a,i,0} \\ & + \beta_3 \text{LSIZE}_{i,-1} + \varepsilon \end{aligned} \quad (12)$$

where $\text{LSIZE}_{i,-1}$ is the natural logarithm of the market value of common stock the day prior to the revision, DA_a is a binary variable set equal to 1 if the analyst is an All-American (first-, second-, third-team, or runner-up) and 0 otherwise, $\text{DN}_a = 1$ if the analyst is a NonAll-American and 0 otherwise, $D1_a = 1$ if the analyst is a first-team All-American and 0 otherwise, $D2_a = 1$ if the analyst is second-team All-American and 0 otherwise, $D3_a = 1$ if the analyst is third-team All-American and 0 otherwise, and $D4_a = 1$ if the analyst is a runner-up and 0 otherwise. Forecasts made by calendar year y All-Americans between October 1 of year y and September 30 of year $y + 1$ are defined as All-American forecasts in the price impact tests.

Equations (10) and (11) assume the magnitude of price reaction to forecast revisions is not a function of the magnitude of the forecast revision. Violations of this assumption are mitigated by restricting the regressions to revisions of a given (relatively constant) magnitude (i.e., the top 10% SUF or the bottom 10% SUF). An advantage of equations (10) and (11) over (12) is that the estimated coefficient for the dummy variable represents the marginal abnormal return associated with large All-American forecast revisions. Equation (12) allows the magnitude of the price reaction to vary with the magnitude of the forecast revision, but the estimated coefficients for SUF have no apparent economic interpretation.

To mitigate potential problems from calendar clustering and heteroscedasticity, the sample is segregated in 60 subsamples on the basis of the calendar month in which day 0 falls and tests are performed on the data by subsample. The significance of the mean coefficients from these 60 regressions is determined by dividing the mean by its standard error, which is the estimated standard deviation of the 60 coefficients divided by the square root of 60 (see Fama and MacBeth (1973)). To test the hypothesis that All-Americans have a greater impact on security prices using equation (12), the difference between β_1 and β_2 ($\beta_1 - \beta_2$) is computed for each subsample and a t -test is performed using the 60 differences.

B. Empirical Results for Impact on Security Prices

For tests of Hypothesis 7, the security price impact hypothesis, returns for firms on the New York Stock Exchange, American Stock Exchange, or NASDAQ are provided by the Center for Research in Security Prices. Again, forecasts that are dated within 1981–85 and within the current fiscal year of the firm are used. Other requirements for sample inclusion are that the forecast revision date is within 200 trading days of the date of the analyst's prior forecast, and that at least two analysts have an outstanding forecast on

the date of the revision. Table VIII reports some sample characteristics for the price impact tests.

Panels A and B of Table IX report the results of cross-sectional regression tests of equations (10) and (11) using the top 10% SUF and the bottom 10% SUF.²² There is evidence that suggests that large upward revisions by All-Americans affect security prices more than large upward revisions by NonAll-Americans.²³ Using the top 10% SUF and the (0, +10) cumulation period, All-Americans have a 0.21% greater average impact on security prices, which is significant at approximately the 0.05 level.²⁴ Distinguishing among All-Americans, first-team All-Americans have a 0.65% greater average impact, which is significant at less than the 0.01 level. However, the magnitude and statistical significance of the differential effect declines using the (0, +20) cumulation period, and, for downward revisions, there is no significant differential effect. As expected, the impact of revisions on prices is generally greater for smaller firms.

Panel C of Table IX reports regression results for tests of equation (12). Using the top 10% SUF, the difference between the coefficients for All-American revisions and NonAll-American revisions is significant at less than the 0.05 level. Using the top 50% SUF, the difference is significant at less than the 0.10 level. For the bottom 10% and bottom 50% SUF, there is no significant differential effect.

Table X documents the economic significance of abnormal returns following forecast revisions without the linearity assumption of regression. Mean cumulative abnormal returns are segregated by the magnitude of the revision, the length of the cumulation period, firm size, and by whether or not the analyst is an All-American. The initial impact of forecast revisions on security prices is greatest for smaller firms experiencing large positive revisions.

²² Each $SUF_{a,i,t}$ is compared to the SUF cutoffs of the most recently ended monthly period. This avoids using distributional information not available at the date of the revision (see Holthausen (1983)).

²³ Because NonAll-Americans revise forecasts less frequently than All-Americans, there is some concern that using the prior forecast as the expectation of the forecast revision (see equation (8)) results in a less well-specified measure of the information content of NonAll-American revisions and a bias in favor of finding greater information content for All-American revisions. However, Stickel (1991) finds that the prior forecast is a better measure of market expectations of the forecast revision than an updated measure that incorporates other analyst forecasts released *since the date* of an analyst's prior forecast into the expectation of the analyst's next forecast. Moreover, including a variable for the number of days elapsed since the date of the prior forecast in regression equation (10), a variable that is predicted to have a negative coefficient for positive revisions if the bias hypothesis is true, results in an insignificant coefficient for that variable (a t -statistic of -0.75 using equation (10) and the top 10% SUF) and no changes in any conclusions (the mean coefficient on DA remains 0.21 with a t -statistic of 1.93).

²⁴ The estimated coefficients from one regression with all revisions are unbiased, but the t -statistics are suspect. Using the top 10% SUF in a single regression, the estimated coefficients (t -statistics) for β_0 , β_1 , β_2 , and β_3 in equation (10) are 4.43 (9.40), -0.04 (-0.42), -0.28 (-8.42), and 0.19 (1.93).

Table VIII
Sample Characteristics for Hypothesis Tests Using Abnormal Returns
Hypothesis 7: Forecast revisions made by All-Americans have a greater impact on security prices than NonAll-American revisions.

	Calendar Year of All-American Status, Brokerage House Affiliation, and Forecast Issuance				
	1981	1982	1983	1984	1985
Panel A: All-Americans					
All-Americans on the Zacks database following firms on the CRSP databases:					
First team	39	45	54	56	53
Second team	33	46	43	49	51
Third team	70	71	72	97	87
Runners-up	122	110	118	125	167
Total number of All-Americans in tests	264	272	287	327	358
Total number of All-Americans per II	324	306	302	383	395
Percentage of All-Americans in tests	81%	89%	95%	85%	91%
Panel B: Brokerage Houses					
Total number of brokerage houses with All-Americans in tests	28	29	30	30	30
Total number of brokerage houses with All-Americans per II	35	34	32	35	38
Panel C: Revisions					
All-Americans	11,268	15,500	14,595	14,699	15,519
NonAll-Americans	20,491	29,796	28,281	32,327	28,578
Total number of revisions	31,759	45,296	42,876	47,026	44,097
Total number of firms	1,449	1,657	1,832	2,027	1,937

IV. Conclusions

All-Americans forecast EPS more accurately and more often than other analysts. Matching forecasts by date of brokerage house issuance, I find mean absolute forecast errors of 95 cents and 98 cents for All-Americans and NonAll-Americans, respectively, for shares that average \$37 in price and \$2.81 in EPS. This contemporaneous advantage is complemented by a timing advantage. All-Americans supply forecasts more often. The mean numbers of calendar days since the prior forecast for revisions dated within a firm's current fiscal year are 86 and 93 for All-Americans and NonAll-Americans, respectively.

There is statistically significant evidence that NonAll-Americans who are about to become All-Americans issue forecasts more frequently than other NonAll-Americans, but insignificant evidence that All-Americans who are about to become NonAll-Americans issue forecasts less frequently than other All-Americans. On the other hand, there is statistically significant evidence that All-Americans who are about to become NonAll-Americans issue less accurate forecasts than other All-Americans, but insignificant evidence that NonAll-Americans who are about to become All-Americans issue more accurate forecasts than other NonAll-Americans.

In terms of security price impact, abnormal returns over the first 11 days after large upward forecast revisions indicate that All-Americans have a 0.21% greater average impact on security prices than NonAll-Americans. Nevertheless, the price impact evidence is mixed to the extent that the economic magnitude and statistical significance of this differential diminishes with longer cumulation periods and, using an 11-day cumulation period, varies substantially from first-team members, an incremental 0.65%, to second-team members, an incremental 0.09%, to third-team members, an incremental -0.03% , to runners-up, an incremental 0.27%. Moreover, for downward revisions, there is virtually no difference between the impact of All-Americans and NonAll-Americans. As expected, the impact of revisions on prices is generally positively related to the change in forecast and is generally greater for smaller firms.

Although the positive relations between reputation and performance are generally statistically significant, there is reason to question the economic significance of the differential performance. However, analyst reputation is based on other performance measures (e.g., buy/hold/sell recommendations and the ability to generate investment banking business) besides those examined here. The association between the performance measures examined here and the other performance measures (e.g., the association between forecast revisions and buy/hold/sell recommendations) is unknown. A study that includes those other performance measures may find the differences in analyst performance to be of greater economic consequence.

There is a positive relation between reputation and performance among security analysts. To the extent that All-Americans are better paid, there is also a relation between pay and performance. Stickel (1990) finds that

Table IX
Hypothesis 7: Do All-American Revisions Affect Security Prices More Than NonAll-American Revisions?

Mean coefficients from 60 regressions of percentage cumulative abnormal returns following forecast revisions on the magnitude of the revision, firm size, and dummy variables indicating All-American status. Revisions are dated within the firm's current fiscal year and within calendar 1981-85 and fiscal year-ends are within 1981-85. $CAR_{i,(a,b)}$ is the percentage cumulative abnormal return to security i from event day a to b . SUF is the standardized unexpected forecast, which is defined as the change in forecast divided by the cross-sectional standard deviation of outstanding forecasts. $LSIZE$ is the natural log of $(1 + \text{Market Value of Common Stock})$. $DA = 1$ if 1st, 2nd, 3rd Team All-American or runner-up; $= 0$ otherwise. $DN = 1$ if NonAll-American; $= 0$ otherwise. $D1 = 1$ if 1st Team All-American; $= 0$ otherwise. $D2 = 1$ if 2nd Team All-American; $= 0$ otherwise. $D3 = 1$ if 3rd Team All-American; $= 0$ otherwise. $D4 = 1$ if runner-up; $= 0$ otherwise. t -Statistics are in parentheses.^a For each regression, the adjusted R^2 is less than 0.01.

Panel A: $CAR_{i,(a,b)} = \beta_0 + \beta_1SUF_{a,i,0} + \beta_2LSIZE_{i,-1} + \beta_3DA_a + \varepsilon$					
SUF Cutoff	Dependent Variable	β_0	β_1	β_2	β_3
Top 10%	$CAR_{i,(0,+5)}$	2.21 (2.76)	-0.06 (-0.61)	-0.14 (-2.30)	0.14 (1.88)
Top 10%	$CAR_{i,(0,+10)}$	3.32 (2.71)	-0.02 (-0.15)	-0.20 (-2.16)	0.21 (1.94)
Top 10%	$CAR_{i,(0,+20)}$	5.31 (2.59)	0.03 (0.16)	-0.33 (-2.03)	0.07 (0.38)
Bottom 10%	$CAR_{i,(0,+5)}$	1.20 (0.96)	1.21 (2.27)	0.05 (0.97)	0.09 (0.86)
Bottom 10%	$CAR_{i,(0,+10)}$	0.45 (0.26)	1.13 (1.54)	0.08 (1.00)	0.02 (0.17)
Bottom 10%	$CAR_{i,(0,+20)}$	-0.60 (-0.26)	1.32 (1.46)	0.16 (1.18)	-0.13 (-0.69)

Table IX—Continued

Panel B: $CAR_{i,(a,b)} = \beta_0 + \beta_1SUF_{a,t,0} + \beta_2LSIZE_{i,-1} + \beta_3D1_a + \beta_4D2_a + \beta_5D3_a + \beta_6D4_a + \varepsilon$									
SUF Cutoff	Dependent Variable	β_0	β_1	β_2	β_3	β_4	β_5	β_6	
Top 10%	$CAR_{i,(0,+10)}$	3.29 (2.71)	-0.02 (-0.14)	-0.20 (-2.15)	0.65 (2.74)	0.09 (0.44)	-0.03 (-0.17)	0.27 (1.84)	
Bottom 10%	$CAR_{i,(0,+10)}$	0.47 (0.27)	1.08 (1.46)	0.07 (0.90)	-0.37 (-1.42)	0.06 (0.20)	0.16 (0.78)	-0.02 (-0.10)	
Panel C: $CAR_{i,(a,b)} = \beta_0 + \beta_1DA_a^*SUF_{a,t,0} + \beta_2DN_a^*SUF_{a,t,0} + \beta_3LSIZE_{i,-1} + \varepsilon$									
SUF Cutoff	Dependent Variable	β_0	β_1	β_2	β_3	$\beta_1 - \beta_2$			
Top 10%	$CAR_{i,(0,+10)}$	3.40 (2.78)	0.11 (0.75)	-0.09 (-0.70)	-0.20 (-2.17)	0.19 (2.28)			
Top 50%	$CAR_{i,(0,+10)}$	1.97 (2.02)	0.63 (5.69)	0.48 (6.06)	-0.13 (-1.74)	0.14 (1.70)			
Bottom 10%	$CAR_{i,(0,+10)}$	0.46 (0.26)	1.13 (1.54)	1.13 (1.54)	0.08 (0.99)	-0.01 (-0.10)			
Bottom 50%	$CAR_{i,(0,+10)}$	-1.03 (-1.22)	0.14 (1.52)	0.17 (1.96)	0.05 (0.78)	-0.04 (-0.64)			

^a *t*-Statistics of 1.65, 1.96, and 2.58 indicate that the coefficient is significantly different from zero at the 0.10, 0.05, and 0.01 levels, respectively.

Table X
Percentage Mean Cumulative Abnormal Returns Following Analyst Forecast Revisions

Revisions are segregated by the magnitude of the revision, firm size, and All-American status. Revisions are dated within the firm's current fiscal year and within calendar 1981-85 and fiscal year-ends are within 1981-85. $MCAR_{(a,b)}$ is the mean percentage cumulative abnormal return from event day a to b . SIZE is the market value of common stock. SUF is the standardized unexpected forecast.

SUF Percentiles	SIZE Quartile	1st Quartile (Smallest Firms)		4th Quartile (Largest Firms)	
		All-Americans	NonAll-Americans	All-Americans	NonAll-Americans
SUF $\leq 10\%$ (most negative)	$MCAR_{(0,+5)}$	-0.53%	-0.48%	-0.29%	-0.34%
	$MCAR_{(0,+10)}$	-0.78%	-0.64%	-0.41%	-0.43%
	$MCAR_{(0,+20)}$	-1.22%	-1.21%	-0.72%	-0.59%
	$MCAR_{(0,+60)}$	-3.62%	-3.27%	-0.32%	-1.14%
	$MCAR_{(0,+120)}$	-5.67%	-5.32%	0.65%	-0.68%
10% < SUF $\leq 25\%$	$MCAR_{(0,+5)}$	0.08%	-0.42%	-0.19%	-0.22%
	$MCAR_{(0,+10)}$	-0.10%	-0.72%	-0.23%	-0.41%
	$MCAR_{(0,+20)}$	-0.54%	-1.08%	-0.37%	-0.63%
	$MCAR_{(0,+60)}$	-2.52%	-2.48%	-0.93%	-1.17%
	$MCAR_{(0,+120)}$	-3.81%	-3.81%	-0.87%	-1.41%
25% < SUF $\leq 50\%$	$MCAR_{(0,+5)}$	-0.17%	-0.22%	-0.12%	-0.33%
	$MCAR_{(0,+10)}$	-0.31%	-0.39%	-0.29%	-0.50%
	$MCAR_{(0,+20)}$	-0.75%	-0.81%	-0.50%	-0.70%
	$MCAR_{(0,+60)}$	-3.03%	-2.40%	-0.81%	-1.16%
	$MCAR_{(0,+120)}$	-4.53%	-3.32%	-0.93%	-1.46%

Table X—Continued

SIZE Quartile	1st Quartile (Smallest Firms)		4th Quartile (Largest Firms)		
	All-Americans	NonAll-Americans	All-Americans	NonAll-Americans	
50% < SUF ≤ 75%	MCAR _(0, + 5)	0.02%	0.03%	− 0.09%	− 0.11%
	MCAR _(0, + 10)	0.08%	− 0.02%	− 0.21%	− 0.24%
	MCAR _(0, + 20)	0.13%	− 0.06%	− 0.42%	− 0.41%
	MCAR _(0, + 60)	− 0.57%	− 0.58%	− 1.28%	− 0.84%
	MCAR _(0, + 120)	− 1.21%	− 0.67%	− 1.45%	− 1.10%
75% < SUF ≤ 90%	MCAR _(0, + 5)	0.85%	0.76%	0.09%	0.15%
	MCAR _(0, + 10)	1.15%	1.06%	− 0.05%	0.21%
	MCAR _(0, + 20)	1.59%	1.68%	0.11%	0.41%
	MCAR _(0, + 60)	3.27%	2.80%	− 0.31%	0.29%
	MCAR _(0, + 120)	3.94%	3.31%	− 0.53%	0.45%
90% < SUF (most positive)	MCAR _(0, + 5)	0.81%	0.79%	0.19%	0.14%
	MCAR _(0, + 10)	1.33%	1.14%	0.34%	0.20%
	MCAR _(0, + 20)	1.83%	1.86%	0.39%	0.50%
	MCAR _(0, + 60)	3.12%	2.67%	− 0.07%	0.61%
	MCAR _(0, + 120)	3.98%	3.46%	0.76%	0.49%

All-American forecasts are less likely to “follow the crowd” and are less predictable, which is consistent with All-Americans being leaders in their profession. The results here provide some justification for All-Americans being the leaders.

REFERENCES

- Abarbanell, Jeffery, 1989, Information acquisition by financial analysts in response to security price changes, Ph.D. dissertation, University of Pennsylvania.
- Atiase, Rowland K., 1985, Predisclosure information, firm capitalization, and security price behavior around earnings announcements, *Journal of Accounting Research* 23, 21–36.
- Beaver, William, Roger Clarke, and William Wright, 1979, The association between unsystematic security returns and the magnitude of the earnings forecast error, *Journal of Accounting Research* 17, 316–340.
- Brown, Lawrence D. and David M. Chen, 1991, How good is the All-America Research Team in forecast earnings?, *Journal of Business Forecasting* 9, 14–18.
- Brown, Philip, George Foster, and Eric Noreen, 1985, *Security Analyst Multi-Year Earnings Forecasts and the Capital Market* (American Accounting Association, Sarasota, Florida).
- Brown, Stephen J. and Jerold B. Warner, 1985, Using daily stock returns: The case of event studies, *Journal of Financial Economics* 14, 3–31.
- Chambers, Anne E. and Stephen H. Penman, 1984, Timeliness of reporting and the stock price reaction to earnings announcements, *Journal of Accounting Research* 22, 21–47.
- Copeland, Thomas E. and David Mayers, 1982, The Value Line enigma (1965–1978): A case study of performance evaluation issues, *Journal of Financial Economics* 10, 289–321.
- Coughlan, Anne T. and Ronald M. Schmidt, 1985, Executive compensation, management turnover, and firm performance: An empirical investigation, *Journal of Accounting and Economics* 7, 43–66.
- Crichfield, Timothy, Thomas Dyckman, and Josef Lakonishok, 1978, An evaluation of security analysts' forecasts, *The Accounting Review* 53, 651–668.
- Fama, Eugene F., 1976, *Foundations of Finance* (Basic Books, New York).
- and James D. MacBeth, 1973, Risk, return, and equilibrium: Empirical tests, *Journal of Political Economy* 81, 607–636.
- Givoly, Dan and Josef Lakonishok, 1979, The information content of financial analysts' forecasts of earnings, *Journal of Accounting and Economics* 1, 165–185.
- Holthausen, Robert, 1983, Abnormal returns following quarterly earnings announcements, *Proceedings of the Seminar on the Analysis of Security Prices* (University of Chicago), 37–59.
- Imhoff, Eugene A. and Gerald J. Lobo, 1984, Information content of analysts' composite forecast revisions, *Journal of Accounting Research* 22, 541–554.
- Jaffe, Jeffrey, 1974, The effect of regulation changes on insider trading, *Bell Journal of Economics and Management Science* 5, 92–121.
- Mandelker, Gershon, 1974, Risk and return: The case of merging firms, *Journal of Financial Economics* 1, 303–336.
- Murphy, Kevin J., 1985, Corporate performance and managerial remuneration: An empirical analysis, *Journal of Accounting and Economics* 7, 3–42.
- O'Brien, Patricia C., 1990, Forecast accuracy of individual analysts in nine industries, *Journal of Accounting Research* 28, 286–304.
- Penman, Stephen H., 1987, The distribution of earnings news over time and seasonalities in aggregate stock returns, *Journal of Financial Economics* 18, 199–228.
- Stickel, Scott E., 1989, The timing of and incentives for annual earnings forecasts near interim earnings announcements, *Journal of Accounting and Economics* 11, 275–292.
- , 1990, Predicting individual analyst earnings forecasts, *Journal of Accounting Research* 28, 409–417.
- , 1991, Common stock returns surrounding earnings forecast revisions: More puzzling evidence, *The Accounting Review* 66, 402–416.