



American Finance Association

Bad News Travels Slowly: Size, Analyst Coverage, and the Profitability of Momentum Strategies

Author(s): Harrison Hong, Terence Lim, Jeremy C. Stein

Source: *The Journal of Finance*, Vol. 55, No. 1 (Feb., 2000), pp. 265-295

Published by: Blackwell Publishing for the American Finance Association

Stable URL: <http://www.jstor.org/stable/222556>

Accessed: 24/11/2009 12:22

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>

Bad News Travels Slowly: Size, Analyst Coverage, and the Profitability of Momentum Strategies

HARRISON HONG, TERENCE LIM, and JEREMY C. STEIN*

ABSTRACT

Various theories have been proposed to explain momentum in stock returns. We test the gradual-information-diffusion model of Hong and Stein (1999) and establish three key results. First, once one moves past the very smallest stocks, the profitability of momentum strategies declines sharply with firm size. Second, holding size fixed, momentum strategies work better among stocks with low analyst coverage. Finally, the effect of analyst coverage is greater for stocks that are past losers than for past winners. These findings are consistent with the hypothesis that firm-specific information, especially negative information, diffuses only gradually across the investing public.

SEVERAL RECENT PAPERS HAVE DOCUMENTED that, at medium-term horizons ranging from three to 12 months, stock returns exhibit momentum—that is, past winners continue to perform well, and past losers continue to perform poorly. For example, Jegadeesh and Titman (1993), using a U.S. sample of NYSE/AMEX stocks over the period from 1965 to 1989, find that a strategy that buys past six-month winners (stocks in the top performance decile) and shorts past six-month losers (stocks in the bottom performance decile) earns approximately one percent per month over the subsequent six months. Not only is this an economically interesting magnitude, but the result also appears to be robust: Rouwenhorst (1998) obtains very similar numbers in a sample of 12 European countries over the period from 1980 to 1995.¹

* Hong is from the Stanford Business School, Lim is from Goldman Sachs, and Stein is from the MIT Sloan School of Management and the National Bureau of Economic Research. This research is supported by the National Science Foundation and the Finance Research Center at MIT. We are grateful to Joseph Chen for research assistance and to Ken French, Paul Pfleiderer, Geert Rouwenhorst, David Scharfstein, Ken Singleton, René Stulz, three anonymous referees, and seminar participants at MIT, Yale, UCLA, Berkeley, Stanford, Illinois, the Norwegian School of Management, and the Stockholm School of Economics for helpful comments and suggestions. Data on analyst coverage were provided by I/B/E/S Inc. under a program to encourage academic research. Thanks also to Lisa Meulbroek for sharing the data on options listings.

¹Rouwenhorst (1997) finds that momentum strategies also earn significant profits on average in a sample of 20 emerging markets. See Haugen and Baker (1996) for confirmatory evidence from the United States and several European countries.

While the existence of momentum in stock returns does not seem to be too controversial, it is much less clear what might be driving it. Some (e.g., Conrad and Kaul (1998)) have suggested a risk-based interpretation of momentum. This is certainly a logical possibility, although there is little evidence that cuts clearly in favor of a risk story. In this vein, Fama and French (1996) note that momentum effects are not subsumed by their three-factor model.

Turning to “behavioral” (i.e., non-risk-based) explanations, there are a number of theories that can give rise to positive medium-term return autocorrelations. In some of these, prices initially *overreact* to news about fundamentals, then continue to overreact further for a period of time. The positive-feedback-trader model of DeLong et al. (1990) fits in this camp, as does the overconfidence model of Daniel, Hirshleifer, and Subrahmanyam (1998). In other models, momentum is a symptom of *underreaction*—prices adjust too slowly to news.

The set of underreaction theories can be further subdivided according to the exact mechanism that is at work. In Barberis, Shleifer, and Vishny (1998), there is a representative investor who suffers from a conservatism bias, and who does not update his beliefs sufficiently when he observes new *public* information. In Hong and Stein (1999) the emphasis is instead on heterogeneities across investors, who observe different pieces of *private* information at different points in time. Hong and Stein make two key assumptions: (1) firm-specific information diffuses gradually across the investing public; and (2) investors cannot perform the rational-expectations trick of extracting information from prices. Taken together, these two assumptions generate underreaction and positive return autocorrelations.

Our goal in this paper is to test the Hong–Stein version of the underreaction hypothesis. In other words, we look for evidence that momentum reflects the gradual diffusion of firm-specific information.² To do so, we begin by sorting stocks into different classes, for which information is *a priori* more or less likely to spread gradually. The central prediction is then that stocks with slower information diffusion should exhibit more pronounced momentum.³

One natural sorting variable—which forms the basis for our first set of tests—is firm size. It seems plausible that information about small firms gets out more slowly; this would happen if, for example, investors face fixed costs of information acquisition, and hence choose in the aggregate to devote more effort to learning about those stocks in which they can take large positions.

² A recent paper that can be thought of in a similar spirit is Chan, Jegadeesh, and Lakonishok (1996). They show that momentum strategies are profitable even after controlling for post-earnings-announcement drift (Bernard and Thomas (1989, 1990), Bernard (1992)). This suggests that momentum at least in part reflects the adjustment of stock prices to the sort of information that (unlike earnings news) is not made publicly available to all investors simultaneously.

³ To obtain this prediction, we are assuming that smart-money arbitrage does not completely eliminate differences in momentum across stocks. This property holds in a wide range of settings. For example, if there is a pool of arbitrageurs that operate across all stocks, it suffices to assume that they are risk-averse and hence prefer to hold diversified portfolios.

Unfortunately, even if firm size is in fact a useful measure of the rate of information diffusion, it is likely to capture other things as well, potentially confounding our inferences. For example, Merton (1987) and Grossman and Miller (1988) argue that market making or arbitrage capacity may be less in small-capitalization stocks. On the one hand, if there are supply shocks, this could lead to a greater tendency toward reversals (i.e., negatively correlated returns) in small stocks, which would obscure the gradual-information-flow effect we are interested in. On the other hand, one might argue that *whatever* behavioral phenomenon is driving positive serial correlation in returns, less arbitrage means that it will have a *bigger* impact in small stocks, leading us to overstate the importance of gradual information flow as the specific mechanism at work. The bottom line is that although it is certainly interesting to see how momentum profits vary with firm size, this probably does not by itself constitute a clean test of our central hypothesis.

As an alternative proxy for the rate of information flow, we consider analyst coverage. The idea here is that stocks with lower analyst coverage should, all else equal, be ones where firm-specific information moves more slowly across the investing public. Thus our second set of tests boils down to checking whether momentum strategies work better in low-analyst-coverage stocks. The one important caveat is that analyst coverage is very strongly correlated with firm size (Bhushan (1989)). So in this second set of tests, we control for the influence of size on analyst coverage by sorting stocks into groups according to their *residual analyst coverage*, where the residual comes from a regression of coverage on firm size.⁴

To preview, we obtain the predicted results for both firm size and residual analyst coverage. First, with respect to size, once one moves past the very smallest capitalization stocks (where thin market making capacity does indeed appear to be an issue) the profitability of momentum strategies declines sharply with market capitalization. Second, holding size fixed, momentum strategies work particularly well among stocks that have low analyst coverage. Moreover, size and coverage interact in a plausible fashion: The marginal importance of analyst coverage is greatest among small stocks. Beyond being statistically significant, these effects are also of an economically interesting magnitude. For example, across our entire sample, momentum profits are roughly 60 percent greater among the one-third of the stocks with the lowest residual coverage, as compared to the one-third with the highest residual coverage.

In addition to these basic findings, we uncover another interesting regularity. There is a strong asymmetry, in that the effect of analyst coverage is much more pronounced for stocks that are past losers than for stocks that

⁴ Our use of residual analyst coverage as a forecaster of stock returns links us to work by Brennan, Jegadeesh, and Swaminathan (1993). They are interested in understanding a higher frequency phenomenon—the fact that at daily and weekly horizons, small stocks seem to lag large stocks (Lo and MacKinlay (1990)). They show that holding size fixed, low-coverage stocks also tend to lag high-coverage stocks, which they interpret as evidence that analysts are important in helping stocks adjust to *common* information. Note that this is quite different from our story, which focuses on the role of analysts in propagating *firm-specific* information.

are past winners. In other words, low-coverage stocks seem to react more sluggishly to bad news than to good news. This makes intuitive sense in the context of a theory based on the flow of firm-specific information. Think of a firm that has no analyst coverage but is sitting on good news. To the extent that its managers prefer higher to lower stock prices, they will push the news out the door themselves, via increased disclosures, etc. On the other hand, if the same firm is sitting on bad news, its managers will have much less incentive to bring investors up to date quickly. Thus the marginal contribution of outside analysts in getting the news out is likely to be greater when the news is bad.

Although all of our evidence is consistent with the sort of gradual-information-flow model in Hong and Stein (1999), it is also possible to put forward an alternative explanation of the data. In particular, it may be that analyst coverage is a proxy for differences in transactions costs that are somehow not well captured by firm size. To take a concrete example, consider two stocks A and B of equal size, where A is harder to sell short than B, and also attracts fewer analysts. Since short-sales constraints can impede the adjustment of prices to negative information, (Diamond and Verrecchia (1987)) this could explain why the low-coverage stock A reacts more slowly—especially to bad news—than the high-coverage stock B.

In an effort to confront this alternative hypothesis, we experiment with two further proxies for transactions costs: share turnover and a dummy variable for the existence of listed options on a given stock. The latter variable might be expected to be particularly useful in picking up cross-sectional differences in ease of shorting, since investors who are not adept at directly shorting a stock can use put options as a substitute. As it turns out, our results are robust to both of these controls. Nevertheless, although these checks are helpful, we recognize that we do not have a perfect measure of transactions costs at the individual stock level, and so cannot definitively rule out all variations of the alternative hypothesis. This is an inevitable shortcoming of our approach.

The remainder of the paper is organized as follows. In Section I we describe our data and analyze in detail the cross-sectional determinants of analyst coverage. Section II contains our main results on momentum strategies sorted by firm size and residual coverage. In Section III we present complementary results based on an alternative, much more parametrically structured, regression approach. Section IV concludes.

I. Cross-Sectional Determinants of Analyst Coverage

Our data come from three primary sources. The stock return and turnover data are from the CRSP Monthly Stocks Combined File, which includes NYSE, AMEX, and Nasdaq stocks. Throughout, we exclude ADRs, REITs, closed-end funds, and primes and scores—that is, stocks that do not have a CRSP share type code of 10 or 11. The data on analyst coverage are from the I/B/E/S Historical Summary File, and are available on a monthly basis beginning in 1976. For each stock on CRSP, we set the coverage in any given

month equal to the number of I/B/E/S analysts who provide fiscal year 1 earnings estimates that month. If no I/B/E/S value is available (i.e., the CRSP cusip is not matched in the I/B/E/S database), we set the coverage to zero. Finally, the options-listing data come from the Options Clearing Corporation, and cover options listed on the CBOE, NYSE, AMEX, Philadelphia, Pacific, and Midwest exchanges.

Table I provides an overview of the extent of analyst coverage for both our full sample (Panel A) as well as for five size-based subsamples (Panel B). The first striking thing that emerges from the table is how many firms show up as having zero analysts. This is especially true in the first few years of the sample period, 1976 to 1978. For example, in 1976, 77.3 percent of all firms appear as having zero analysts. There is a marked deepening of coverage around 1980, with the fraction of uncovered firms dropping to 58.2 percent. After that, things change much more smoothly, with the fraction of uncovered firms declining gradually to 36.9 percent in 1996.

While the numbers no doubt largely reflect the reality that many firms are simply not covered by analysts, we worry that they may also be somewhat contaminated by measurement error. It is possible that the I/B/E/S data set is missing information on some firms' analysts. Alternatively, it is possible that I/B/E/S has the data, but has assigned a different cusip number to a firm than CRSP. In this case, we would mistakenly code the CRSP firm as having no analysts. In principle, such measurement error should make our tests err on the side of conservatism—it should be harder to discern significant differences across stocks that we classify as low coverage versus high coverage. Because of this concern, and because the number of zeros is so much higher in the first few years, all the tests we present below use a sample period that runs from 1980 to 1996.⁵ However, it should be noted that none of our results are materially altered if we begin in 1976 instead.

A second key fact that comes out of Table I is that for the smallest firms, there is simply *no variation in coverage*. Consider those firms that are smaller than the 20th percentile NYSE/AMEX firm. As can be seen, almost all of them have zero analysts—82 percent are not covered in 1988, which is roughly the midpoint of the sample period we use. Consequently, we simply cannot use this part of the population to test any hypotheses having to do with analyst coverage. Hence, all our coverage-related tests begin with a subsample that excludes those firms that are below the 20th percentile NYSE/AMEX breakpoint in any given month.⁶ Note that there is much more variation in analyst coverage in the next size class, which runs from the 20th to the 40th percentile of NYSE/AMEX—in 1988, only 41.7 percent of the firms in this class are not covered, and a substantial fraction have as many as three or four analysts.

⁵ For reasons that we explain later, we typically measure analyst coverage six months before we actually begin to implement our momentum strategies. Since our sample period for measuring returns begins in 1980, we use analyst data as far back as 1979.

⁶ The cutoff point is around \$30 million in market capitalization as of the midpoint of the sample period, and rises to almost \$60 million by 1996.

Table I
Descriptive Statistics for Analyst Coverage

Descriptive statistics for analyst coverage for NYSE, AMEX, and Nasdaq stocks, excluding ADRs, REITs, closed-end funds, and primes and scores during the period 1976 to 1996. Panel A reports for the even years between 1976 and 1996 the number of firms in the sample, their mean and median size, the number of analysts at various coverage percentiles, and the percentage of firms that had no coverage. Panel B reports for 1988 by firm size the same statistics as in Panel A.

Panel A: All Stocks, 1976–1996													
Year	No. of Firms	Mean Size (millions)	Median Size (millions)	No. of Analysts at Coverage Percentiles							Percentage of firms uncovered		
				10	20	30	40	50	60	70		80	90
76	4402	183.6	18.7	0	0	0	0	0	0	0	1	4	77.3%
78	4472	176.4	22.7	0	0	0	0	0	0	0	2	5	71.5%
80	4329	248.9	34.6	0	0	0	0	0	1	2	4	9	58.2%
82	4754	249.3	30.3	0	0	0	0	0	1	2	5	11	59.3%
84	5049	332.3	44.4	0	0	0	0	0	1	3	6	12	50.8%
86	5364	387.4	42.5	0	0	0	0	0	1	3	6	14	50.5%
88	5932	402.2	32.6	0	0	0	0	0	1	3	5	12	50.1%
90	5567	500.7	34.5	0	0	0	0	1	2	3	7	13	45.4%
92	5438	672.8	49.8	0	0	0	0	1	2	3	6	13	46.7%
94	5890	802.9	81.1	0	0	0	0	1	3	4	7	13	40.0%
96	6460	978.1	90.8	0	0	0	1	2	3	4	7	12	36.9%

Panel B: Breakdown of Analyst Coverage by Firm Size for 1988													
NYSE/AMEX Breakpoints	No. of Firms	Mean Size (millions)	Median Size (millions)	No. of Analysts at Coverage Percentiles							Percentage of firms uncovered		
				10	20	30	40	50	60	70		80	90
Below the 20th percentile	2597	9.6	8.3	0	0	0	0	0	0	0	0	1	82.0%
Between the 20th & 40th percentiles	1363	45.1	42.5	0	0	0	0	1	1	2	3	4	41.7%
Between the 40th & 60th percentiles	937	147.1	133.3	0	0	1	2	3	4	5	7	9	21.5%
Between the 60th & 80th percentiles	607	554.0	495.8	1	4	6	7	8	10	12	14	17	7.7%
Above the 80th percentile	431	4235.7	2390.7	8	13	16	19	21	23	26	28	30	5.6%

In Table II, we examine the cross-sectional determinants of analyst coverage. When we actually implement our trading strategies in the next section, we run a separate regression every month to create our measure of residual coverage. Because the regressions look so similar month to month, we only present one set in Table II for illustrative purposes, corresponding to December 1988, which is around the midpoint of our sample period. Again, note that in each case, the regression is run only on those stocks that are larger than the 20th percentile NYSE/AMEX breakpoint in the given month.

The first point to note is that unlike some previous researchers who have run similar regressions (e.g., Bhushan (1989) and Brennan and Hughes (1991)) we use as our left-hand side variable $\log(1 + \text{Analysts})$, rather than the raw number of analysts. We do this because we ultimately want to use the residuals from our analyst-coverage regressions to explain momentum, and it seems plausible that one extra analyst should matter much more in this regard if a firm has few analysts than if it has many.

In Model 1, we use OLS, and the only two right-hand side variables are $\log(\text{Size})$, where Size is current market capitalization, and a Nasdaq dummy variable.⁷ The size variable is clearly enormously important, generating an R^2 of 0.61. In Model 2, we add 15 industry dummies to the regression.⁸ This has a small effect, raising the R^2 to 0.63.

In Models 3 and 4, we try adding the firm's book-to-market ratio. We do this because book-to-market is known to forecast returns (Fama and French (1992), Lakonishok, Shleifer, and Vishny (1994)) and we want to make sure that any return-predicting power we get out of analyst coverage is not simply capturing a book-to-market effect. As it turns out, the coefficient on book-to-market is positive and significant, but it adds nothing at all to the R^2 . Thus it is unlikely that any of the results we report below are driven by anything to do with book-to-market.⁹ In Models 5 and 6, we undertake a similar experiment with beta.¹⁰ The coefficient on beta is positive and strongly significant, and in this case, the R^2 increases marginally, going from 0.61 to 0.63 when we exclude industry dummies.

⁷ The Nasdaq dummy is the only variable whose behavior changes much over the sample period. In earlier years, it is strongly negative, which is why we include it in our baseline model. However, by the late 1980s, it is typically positive, though not always significantly so.

⁸ The dummies correspond to the following grouping of two-digit SIC codes: (1) SIC 01–09; (2) SIC 10–14; (3) SIC 15–19; (4) SIC 20–21; (5) SIC 22–23; (6) SIC 24–27; (7) SIC 28–32; (8) SIC 33–34; (9) SIC 35–39; (10) SIC 40–48; (11) SIC 49; (12) SIC 50–52; (13) SIC 53–59; (14) SIC 60–69; and (15) SIC 70–79.

⁹ Even if high-coverage stocks do have higher mean returns because they have a higher loading on book-to-market, this cannot explain our central result, namely that high-coverage stocks exhibit less momentum.

¹⁰ Throughout, we calculate beta with the Scholes–Williams (1977) method, using daily returns and the value-weighted CRSP index in the prior calendar year. We require that 50 percent of single-day trade-only returns (computed using closing prices, not bid/ask averages) be available. This is the same approach used by CRSP in its NYSE/AMEX Excess Returns File.

Table II
Determinants of Analyst Coverage, 12/1988

Dependent variable is $\log(1 + \text{Analyst coverage})$. Log Size is the log of a firm's year-end market value. NASD is a Nasdaq dummy. Book/Mkt is the ratio of a firm's year-end book-to-market value. Beta is a firm's market beta. P is a firm's share price. Var is the variance of a firm's return using the last 200 observations from year-end. R_k is the rate of return of a firm lagged k years for $k = 0, 1, 2, 3, 4$. T-O is a firm's turnover defined as the prior six months' trading volume divided by shares outstanding. NASD * T-O is the Nasdaq dummy times firm turnover. OPT is a dummy for whether a firm has options trading on CBOE, NYSE, AMEX, Philadelphia, or Pacific stock exchanges. IND is a set of CRSP industry dummies. There are 2,012 observations. t -statistics are in parentheses.

Model No.	Log Size	NASD	Book/Mkt	Beta	1/P	Var	R ₀	R ₁	R ₂	R ₃	R ₄	T-O	NASD * T-O	OPT	IND	R ²
1	0.54 (52.67)	0.03 (0.99)													No	0.61
2	0.56 (52.90)	0.04 (1.21)													Yes	0.63
3	0.55 (53.03)	0.05 (1.50)	0.12 (3.15)												No	0.61
4	0.57 (52.22)	0.07 (2.00)	0.17 (4.30)												Yes	0.63
5	0.50 (48.41)	0.07 (2.28)		0.38 (11.54)											No	0.64
6	0.51 (46.11)	0.09 (2.62)		0.40 (10.94)											Yes	0.65
7	0.57 (49.87)	0.09 (2.59)			-0.52 (-3.12)	-1.27 (-3.23)	-0.50 (-9.46)	-0.28 (-6.06)	-0.28 (-6.00)	-0.04 (-0.85)	-0.16 (-3.46)				Yes	0.65
8	0.52 (51.46)	-0.02 (-0.54)										3.82 (8.18)	-0.53 (-0.93)		No	0.64
9	0.50 (38.83)	-0.02 (-0.48)										3.52 (7.32)	-0.37 (-0.64)	0.12 (2.48)	No	0.64

In Model 7, we add to the industry-dummy specification of Model 2 a number of variables that are considered in Brennan and Hughes (1991): $1/P$, where P is the price of a share; the variance of daily returns; and five years' worth of annual lagged returns. Although many of the coefficients are individually significant, the overall impression is that these extra variables are not very important in explaining the variation in coverage—jointly they raise the R^2 from 0.63 to 0.65.¹¹

In Model 8, we take the baseline specification of Model 1 and add a turnover measure, defined as the number of shares traded over the prior six months divided by total shares outstanding. (Because turnover numbers may not have the same interpretation in a dealer market, we allow the coefficient on turnover to be different for Nasdaq firms.) Turnover is significantly positively correlated with coverage on all exchanges, and it raises the R^2 somewhat, from 0.61 to 0.64. However, with this regression, one needs to be especially careful in attaching any causal interpretation. On the one hand, it is possible that turnover causes coverage: Analysts may be more inclined to follow naturally high-turnover stocks if this makes it easier to generate brokerage commissions for their employers (Hayes (1996)). On the other hand, Brennan and Subrahmanyam (1995) find evidence of causality running in the other direction: More analysts reduce the adverse-selection costs of trading, and thereby attract a greater volume of trade. As we argue in Section II.D below, depending on which story one believes, it may or may not make sense to control for turnover in generating our measure of residual analyst coverage.

Continuing in a similar vein, Model 9 adds to the turnover measure of Model 8 another proxy for transactions costs, a dummy variable that takes on the value one if the stock in question has listed options. (About 25 percent of our sample firms have listed options in 1988, with the fraction rising to 49 percent by 1996.) As can be seen, the options-listing dummy has the expected positive sign and is statistically significant. However, unlike turnover, it adds virtually nothing to the explanatory power of the regression—the R^2 remains at 0.64, just as in Model 8.

Overall, the results in Table II make it clear that although a number of other variables are significantly related to analyst coverage, firm size is by far the dominant factor. Thus, in addition to worrying about the influence of these other variables, it is also important to think about potential nonlinearities in the relationship between $\log(1 + \text{Analysts})$ and $\log(\text{Size})$. In this spirit, we proceed as follows. We start in Section II.B by using the simple size-based regression in Model 1 as our baseline method for generating re-

¹¹ Interestingly, our results call into question the conclusions of Brennan and Hughes (1991), who obtain significant positive coefficients on $1/P$. In our regressions, we tend to get the opposite sign. We conjecture that this arises because we are using $\log(1 + \text{Analysts})$ on the left-hand side, rather than the raw number of analysts. Because $1/P$ is correlated with firm size, and because firm size is of such dominant importance, any differences in how one models the analyst-size relationship is likely to have a strong influence on the $1/P$ coefficient.

sidual analyst coverage. Next, in Section II.C we rerun all of our tests separately for each of the size classes (except the very smallest) in Table I. In this case, we run a separate cross-sectional analyst regression each month for firms in the 20th–40th NYSE/AMEX percentiles, for firms in the 40th–60th percentiles, and so on. Among other things, this approach allows the relationship between $\log(1 + \text{Analysts})$ and $\log(\text{Size})$ to take on a piecewise linear form, hopefully correcting any deficiencies that arise from imposing an overly simple linear structure on the entire sample.

Moreover, in Section II.D we also report on sensitivity checks that take into account the potential for analyst coverage to be correlated with some of the other variables considered in Table II. For example, we experiment with alternative definitions of residual coverage based on Model 2, which includes the industry dummies, and Models 8 and 9, which include turnover and the options-listing dummy. Furthermore, we redo our tests in terms of beta-adjusted returns in case the pronounced relationship between beta and analyst coverage is affecting the results.

II. Momentum Strategies, Cut Different Ways

A. Cuts on Raw Size

We begin our analysis of momentum strategies in Table III. In this table, unlike in the tables that come later, we look at the entire universe of stocks without dropping those below the 20th NYSE/AMEX percentile. In so doing, we closely follow the methodology of Jegadeesh and Titman (1993) in many respects. In particular, we focus on their preferred six-month/six-month strategy, we couch everything in terms of raw returns, and we equal-weight these returns. But there are three noteworthy differences. First, our sample period from 1980 to 1996 is more recent. Second, we do not exclude Nasdaq stocks. And third, our measure of momentum differs from theirs. They sort stocks into 10 deciles according to past performance, and then measure the return differential of the most extreme deciles—which they denote by $P10 - P1$. In contrast, we place less emphasis on the tails of the performance distribution. We sort our sample into only three parts based on past performance: P1, which includes the worst-performing 30 percent; P2 which includes the middle 40 percent; and P3, which includes the best-performing 30 percent. Our basic measure of momentum is then $P3 - P1$. This is similar to the measure used by Moskowitz (1997) and Rouwenhorst (1997).

We use this alternative, broader-based measure of momentum in order to generate better signal-to-noise properties for our tests. Unlike Jegadeesh and Titman (1993), we are not so much interested in establishing the existence of momentum per se, but in *comparing* momentum effects across subsamples of stocks. In some cases, we look at as many as 12 subsamples, when we sort by size and residual analyst coverage simultaneously. (See Table V below.) If we also were to use 10 performance deciles, we would end

Table III
Momentum Strategies, 1/1980–12/1996, Using Raw Returns and Sorting by Size

This table includes all stocks. The relative momentum portfolios are formed based on six-month lagged raw returns and held for six months. The stocks are ranked in ascending order on the basis of six-month lagged returns. Portfolio P1 is an equally weighted portfolio of stocks in the worst-performing 30 percent, portfolio P2 includes the middle 40 percent, and portfolio P3 includes the best-performing 30 percent. This table reports the average monthly returns of these portfolios and portfolios formed using size-based subsamples of stocks. Using NYSE/AMEX decile breakpoints, the smallest firms are in size class 1, the next in 2, and the largest are in 10. Mean (median) size is in millions. *t*-statistics are in parentheses.

Past	Size Class (NYSE/AMEX Decile Breakpoints)										
	All Stocks	1	2	3	4	5	6	7	8	9	10
P1	0.01043 (2.44)	0.02106 (4.44)	0.00653 (1.37)	0.00231 (0.52)	0.00194 (0.43)	0.00469 (1.05)	0.00573 (1.32)	0.00606 (1.43)	0.01010 (2.51)	0.00922 (2.25)	0.01258 (3.37)
P2	0.01378 (4.48)	0.01662 (4.97)	0.01290 (3.84)	0.01280 (3.88)	0.01244 (3.75)	0.01395 (4.18)	0.01374 (4.14)	0.01375 (4.27)	0.01393 (4.40)	0.01401 (4.43)	0.01355 (4.50)
P3	0.01570 (4.35)	0.01733 (4.40)	0.01507 (3.89)	0.01664 (4.35)	0.01570 (4.05)	0.01655 (4.26)	0.01608 (4.26)	0.01491 (4.13)	0.01436 (4.04)	0.01363 (3.96)	0.01278 (3.84)
P3 – P1	0.00527 (2.61)	-0.00374 (-1.77)	0.00854 (3.60)	0.01433 (6.66)	0.01376 (6.10)	0.01187 (5.32)	0.01035 (4.80)	0.00885 (3.72)	0.00425 (1.90)	0.00441 (1.73)	0.00021 (0.08)
P2-P1 P3-P1		—	0.746	0.732	0.763	0.780	0.774	0.869	0.901	1.086	—
Mean size		7	21	44	79	138	242	437	806	1658	7290
Median size		7	21	43	78	136	237	430	786	1612	4504
Mean analyst		0.1	0.5	1.1	2.0	3.2	5.0	7.3	10.6	15.3	21.4
Median analyst		0.0	0.0	0.7	1.3	2.5	4.4	6.9	10.5	15.7	22.4

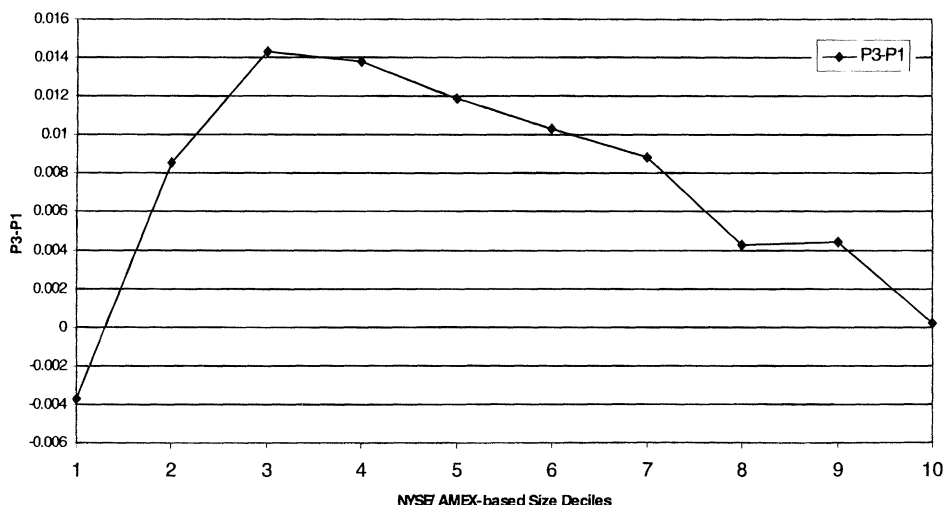


Figure 1. Momentum profits and firm size. Momentum profits ($P3 - P1$) plotted against NYSE/AMEX-based size deciles, 1 (smallest) to 10 (largest).

up chopping the universe of stocks into 120 portfolios, and we would reach a point where some of the individual portfolios are quite undiversified, thereby creating larger standard errors in our test statistics.¹²

The first column in Table III confirms that there is significant momentum in the full sample: The baseline strategy that buys top-30 percent ($P3$) winners and shorts bottom-30 percent ($P1$) losers generates 0.53 percent per month (t -statistic = 2.61).¹³ The next columns break the momentum effect down by size (measured six months before the start of the ranking period). We use an independent sort to generate 10 subsamples, with the break-points determined by NYSE/AMEX deciles. Figure 1 illustrates the results, plotting the relationship between size and the magnitude of the momentum effect. As can be seen, there is a pronounced, inverted U-shape. In the very smallest stocks (which are tiny, with a mean market capitalization of \$7 million) momentum is actually *negative*. By the second size decile, momen-

¹² In fact, we have redone all our key tests, using the Jegadeesh and Titman (1993) $P10 - P1$ momentum measure in place of our $P3 - P1$ measure. As might be expected, the point estimates of interest—that is, the differences in momentum between low- and high-coverage firms—are typically larger in absolute value. However, the standard errors are also larger, so in many cases the t -statistics turn out to be smaller. This confirms the notion that our $P3 - P1$ measure has better signal-to-noise properties for the particular type of tests we focus on.

¹³ This is lower than the Jegadeesh–Titman (1993) figure of 0.95 percent per month. The difference arises for two distinct reasons noted above. First, our strategy invests in stocks with less-extreme past performance. And second, it turns out that including the smaller Nasdaq firms substantially damps the results since, as can be seen from Table III, the momentum measure is actually negative for the very smallest firms. The different sample period is *not* responsible for the difference in results because when we use an NYSE/AMEX sample and a $P10 - P1$ momentum measure over our sample period we obtain numbers almost identical to Jegadeesh and Titman.

tum profits are significantly positive, and they reach a peak in the third size decile, where market capitalization averages about \$45 million and where the profits are a striking 1.43 percent per month (t -statistic = 6.66), which is almost three times the value for the sample as a whole. After the third size decile, momentum profits decline monotonically to the point where they are essentially zero in the largest stocks.¹⁴

The nonmonotonic effect of raw size can be easily understood in the context of the informal theory sketched in the Introduction: Smaller firms may have slower information diffusion, which would lead to greater momentum, but they probably also have more limited investor participation (i.e., thinner market making capacity) which can lead to more pronounced supply-shock-induced reversals.¹⁵ Under this interpretation, the sharp decline in momentum profits that occurs between the third and the tenth size classes is testament to the economic importance of gradual information diffusion in mid-cap stocks.

Another interesting pattern that emerges in Table III is that the bulk of the momentum effect seems to come from losers, as opposed to winners. Consider for example, the column corresponding to the third size class, where, as noted above, the $P3 - P1$ winners-minus-losers measure is 1.43 percent per month. Of that, 1.05 percent, or about three-quarters of the total, comes from the difference between average performers and losers—that is, from $P2 - P1$. As can be seen from the table, this tendency holds with remarkable consistency in every one of the size classes (i.e., deciles two through eight) where there are positive momentum profits to begin with.¹⁶ It suggests that to the extent that stock prices do underreact, they are more prone to underreact to bad news than to good news. We return to this theme in greater detail below.

B. Cuts on Residual Analyst Coverage

Next we turn to the cuts based on residual analyst coverage. Here, and in everything that follows, we exclude all stocks that are below the 20th percentile NYSE/AMEX breakpoint. Again, this is because the vast majority of these small stocks simply never have any analyst coverage, so there is no

¹⁴ Jegadeesh and Titman (1993) also find that momentum profits follow a hump shape with respect to size (see their Table III, p. 78). But they document only small differences across subsamples. This is because they only use three size classes, and exclude Nasdaq firms; much of the variation in size is thus either blurred or omitted. It should also be noted that the hump shape is robust to a number of variations—for example, skipping a month between the ranking period and the holding period, or eliminating January returns. The latter reduces overall momentum, but does not alter the nonmonotonic relationship between momentum and size.

¹⁵ Alternatively, it may be that many of the tiniest stocks trade at very low dollar prices, so we are picking up some discreteness-induced negative correlation. Since we do not pay any further attention to this class of stocks in what follows, we do not pursue this possibility.

¹⁶ In Jegadeesh and Titman's (1993) full sample, the asymmetry between winners and losers is not so big. This discrepancy appears to come from the behavior of the very smallest loser stocks, which, as Table III shows, actually exhibit strong reversals. When one excludes these tiny stocks, as we do, the winner-loser asymmetry becomes much more pronounced.

variation to work with. Within this truncated universe, we create three subsamples based on residual analyst coverage, with the residuals coming from month-by-month cross-sectional regressions of $\log(1 + \text{Analysts})$ on $\log(\text{Size})$ and a Nasdaq dummy, just as in Model 1 of Table II.

In implementing this technique, we choose to measure residual coverage *six months before* we start our preformation ranking period.¹⁷ We use slightly “stale” data on analyst coverage in order to address a possible endogeneity concern. McNichols and O’Brien (1996) find that analysts are more likely to begin covering firms when they are optimistic about their near-term prospects. When one combines this finding with Womack’s (1996) evidence that there is stock price drift for up to six months in response to analyst recommendations, it raises the possibility that recent innovations in analyst coverage may be informative about future returns. Although we have no reason to expect that this form of endogeneity would bias any of our key tests one way or another, we adopt the stale data approach as a simple precaution. Intuitively, any patterns that we now find are driven by the permanent component of coverage, and not by recent (and possibly return-predicting) innovations in coverage. These caveats notwithstanding, our results seem very insensitive to exactly when we measure analyst coverage. We have experimented with measuring it zero, 12, and 18 months prior to our ranking period, and in each case we obtain very similar results.

Table IV presents the results of this approach. Before getting to the returns for the three subsamples, it is important to check that they have the desired characteristics with respect to size and coverage. Ideally, the subsamples will contain stocks of the same size, yet will display a healthy spread in coverage. As can be seen from the table, the variation in coverage is certainly there. The low-coverage subsample, which we denote Sub1, has median coverage of 0.1 (mean of 1.5), and the high-coverage subsample Sub3 has median coverage of 7.6 (mean of 9.7). We do a little less well in terms of size matching. Sub1 has a somewhat larger mean size than Sub3 (\$962 million versus \$455 million) and at the same time a smaller median size (\$103 million versus \$180 million). Evidently, due to nonlinearities in the analyst-size relationship, the simple linear regression technique is giving us residuals that do not have exactly the same size distribution across the three subsamples.¹⁸ We attempt to remedy this deficiency shortly, in Table V. For the moment, it suffices to say that the imperfect size matching in Table IV does not color any of the conclusions.

¹⁷ Concretely, our first month’s worth of observations has the following timing: (1) we measure residual coverage based on a regression using data as of January 1979; (2) in an independent sort, we rank stocks on their performance in the six months from June 30, 1979 to December 31, 1979 and assign them to either P1, P2, or P3; and (3) we then calculate the realized returns for the coverage/past-performance portfolios over the next six months, which run until June 30, 1980.

¹⁸ What seems to be going on is this: After a point, the number of analysts simply maxes out, and no longer increases with size. Thus with a linear model, the very largest firms—the Intels and GMs of the world—tend to show up as having abnormally low coverage relative to their size, thereby landing in Sub1. This pushes the mean size in Sub1 up relative to that in Sub3.

Turning to the returns numbers, two patterns emerge that hold up throughout our subsequent analysis. First, as predicted by the theory, there is more momentum in stocks with low residual coverage. The P3 – P1 momentum measure is 1.13 percent per month in the low-residual-coverage subsample Sub1, and only 0.72 percent per month in the high-residual-coverage subsample Sub3.¹⁹ The difference of 0.42 percent between Sub1 and Sub3 in this regard is highly statistically significant, with a *t*-statistic of 3.50. Moreover, the economic magnitude is clearly important—momentum profits are roughly 60 percent higher in Sub1 than in Sub3.

The second key finding is that the effect of residual coverage on the P3 – P1 momentum measure is entirely driven by what happens in the loser stocks in P1.²⁰ P1/Sub1 stocks underperform P1/Sub3 stocks by 0.70 percent per month. This difference is also highly significant, with a *t*-statistic of 5.16. In other words, one attractive strategy, which we call the “loser-analyst-spread trade,” or “LAST” strategy, is simply to buy the stocks in P1/Sub3 and short those in P1/Sub1, without ever dealing with any of the winner stocks in P3. This strategy is not only size-neutral, but also (unlike the Jegadeesh–Titman strategy) momentum-neutral. So to the extent that anybody ever makes an argument that momentum returns are proxying for a risk factor, our LAST strategy earns 0.70 percent per month *with no loading on that risk factor*.

Taken together, these two patterns suggest that analyst coverage is especially important in propagating bad news. This ties together nicely with our earlier finding that the bulk of momentum profits seem to come from loser stocks. And as we noted in the Introduction, it also makes intuitive economic sense. When firms are sitting on good news, managers probably have every incentive to push this news out to investors as fast as possible, which makes analysts less important. In contrast, when there is bad news, managers are likely to be less forthcoming, so outside analysts have a more crucial role to play.²¹

¹⁹ For the full sample in Table IV, the P3 – P1 value is 0.94 percent per month. This is higher than in Table III because we have now dropped the smallest firms, which as seen above, have negative momentum.

²⁰ Indeed, the numbers in P3 go slightly the “wrong way”—the continuing performance of low-coverage winners is a bit *worse* than that of high-coverage winners. Although this difference between P3/Sub1 and P3/Sub3 is statistically significant in Table IV, it, much more so than our other results, appears to be fragile. For example, it totally disappears when we work with beta-adjusted returns in Table VI below. To the extent that there is a premium for beta in our sample period, this should not be surprising since, as we saw in Table II, low coverage is associated with lower values of beta. In fact, the median beta in Sub1 is 0.75, versus 0.95 in Sub3.

²¹ A large literature finds that analysts tend to be too optimistic about firms’ prospects (see Lim (1998) and Easterwood and Nutt (1998) for recent examples). Note that there need be no contradiction between this work and our claim that analysts are important for propagating bad news. Smart investors will deflate analysts’ overhyped reports, so a “hold” recommendation as opposed to a “buy” can be a powerful bad signal, even if analysts rarely say “sell.”

Table IV
Momentum Strategies, 1/1980–12/1996, Using Raw Returns
and Sorting by Model 1 Residuals

This table includes only stocks above the NYSE/AMEX 20th percentile. The relative momentum portfolios are formed based on six-month lagged raw returns and held for six months. The stocks are ranked in ascending order on the basis of six-month lagged returns. Portfolio P1 is an equally weighted portfolio of stocks in the worst-performing 30 percent, portfolio P2 includes the middle 40 percent, and portfolio P3 includes the best-performing 30 percent. This table reports the average monthly returns of these portfolios and portfolios formed using an independent sort on Model 1 analyst coverage residuals of log size and a Nasdaq dummy. The least-covered firms are in Sub1, the medium covered firms in Sub2, the most covered firms in Sub3. Mean (median) size is in millions. *t*-statistics are in parentheses.

Past	Residual Coverage Class				
	All Stocks	Low: Sub1	Medium: Sub2	High: Sub3	Sub1 – Sub3
P1	0.00622 (1.54)	0.00271 (0.66)	0.00669 (1.70)	0.00974 (2.31)	–0.00703 (–5.16)
P2	0.01367 (4.40)	0.01257 (4.20)	0.01397 (4.58)	0.01439 (4.29)	–0.00182 (–2.11)
P3	0.01562 (4.35)	0.01402 (3.95)	0.01583 (4.52)	0.01690 (4.45)	–0.00288 (–2.80)
P3 – P1	0.00940 (4.89)	0.01131 (5.46)	0.00915 (4.64)	0.00716 (3.74)	0.00415 (3.50)
Mean size		962	986	455	
Median size		103	200	180	
Mean analyst		1.5	6.7	9.7	
Median analyst		0.1	3.5	7.6	

C. Two-Way Cuts on Size and Residual Coverage

In Table V, we disaggregate the analysis of Table IV by size. The methodology is exactly the same except that we look at four separate subsamples. The first includes all stocks between the 20th and 40th NYSE/AMEX percentiles, the second includes those between the 40th and 60th percentiles, and so forth. We have two motivations for doing this disaggregation. First, as a matter of economics, it seems reasonable to conjecture that the marginal importance of coverage is greater in the smaller stocks, which have fewer analysts on average, and are probably less well researched in other ways. Second, as a matter of methodology, this approach should give us better size matches across residual coverage classes, since we now run the analyst coverage regressions separately for each size-based subsample. Compared to our earlier approach, this is like allowing the analyst-size relationship to be piecewise linear.

As can be seen from the table, the size matching is now almost flawless, except in the largest class of stocks. Consider first the results for the smallest size class, that corresponds to the 20th to 40th percentile range. The mean size is \$63 million in Sub1, versus \$64 million in Sub3. (The medians are

Table V

Momentum Strategies, 1/1980–12/1996, Using Raw Returns and Sorting by Size and Model 1 Residuals

This table includes only stocks above the NYSE/AMEX 20th percentile. The relative momentum portfolios are formed based on six-month lagged raw returns and held for six months. The stocks are ranked in ascending order on the basis of six-month lagged returns. Portfolio P1 is an equally weighted portfolio of stocks in the worst-performing 30 percent, portfolio P2 includes the middle 40 percent, and portfolio P3 includes the best-performing 30 percent. This table reports the average monthly returns to portfolios formed by sorts on size and Model 1 analyst coverage residuals of log size and a Nasdaq dummy. Size is sorted using NYSE/AMEX breakpoints. The least covered firms are in Sub1, the medium covered firms in Sub2, the most covered firms in Sub3. Mean (median) size is in millions. *t*-statistics are in parentheses.

Residual Coverage Class	Size Class:			
	1: 20th–40th Percentile	2: 40th–60th Percentile	3: 60th–80th Percentile	4: 80th–100th Percentile
Low: Sub1	P3 – P1 = 0.01511 (6.46)	P3 – P1 = 0.01057 (4.49)	P3 – P1 = 0.00605 (3.11)	P3 – P1 = 0.00092 (0.49)
Mean size	63	199	653	5056
Median size	59	183	592	2363
Median coverage	0.0	0.6	3.7	11.1
Medium: Sub2	P3 – P1 = 0.01389 (5.48)	P3 – P1 = 0.00975 (4.95)	P3 – P1 = 0.00316 (1.62)	P3 – P1 = 0.00009 (0.05)
Mean size	61	207	678	5163
Median size	56	193	629	2853
Median coverage	0.9	3.6	9.0	18.8
High: Sub3	P3 – P1 = 0.01147 (5.10)	P3 – P1 = 0.00730 (3.60)	P3 – P1 = 0.00424 (2.02)	P3 – P1 = 0.00070 (0.33)
Mean size	64	202	663	3650
Median size	61	188	615	2511
Median coverage	3.1	7.6	14.7	24.9
Sub1 – Sub3	P3 – P1 = 0.00364 (2.13)	P3 – P1 = 0.00327 (1.95)	P3 – P1 = 0.00180 (1.18)	P3 – P1 = 0.00023 (.14)

\$59 and \$61 million respectively.) Yet we still have a good spread in coverage, with a median of 0.0 analysts in Sub1 and 3.1 analysts in Sub3. And the basic results from Table IV carry over. The $P3 - P1$ momentum measure is 1.51 percent per month in Sub1, and 1.15 percent per month in Sub3. The difference of 0.36 percent is statistically significant (t -statistic of 2.13) even though the standard errors are naturally quite a bit higher with the smaller sample.

As we move to progressively larger size classes, two things happen. First, the overall momentum effect shrinks, just as in Table III. Second, the *differential* in momentum between Sub1 and Sub3 shrinks also, consistent with the hypothesis that the marginal importance of analysts should decline with size. In the next size class, covering the 40th–60th percentile range, where stocks average approximately \$200 million in market capitalization, the Sub3 – Sub1 momentum differential is not much smaller, at 0.33 percent (t -statistic = 1.95). But by the time we get to the 60th–80th percentile range, where mean size is close to \$700 million, the differential is down to 0.18 percent (t -statistic = 1.18). And it is essentially zero for the largest size class.

Overall, the size disaggregation effort in Table V lends further credence to our interpretation of the evidence. It makes it clear that the earlier numbers in Table IV are not an artifact of imperfect size matching in the full sample. And it is comforting to know that analyst coverage has more of an impact on momentum in precisely those parts of the size distribution where one *a priori* suspects that gradual information diffusion is likely to be important and where momentum effects are most pronounced to begin with.

Table V also helps put into perspective the extent to which firm size and residual coverage might each be capturing something related to the phenomenon of gradual information flow. On the one hand, it is natural to focus most of the attention on residual coverage as a proxy for this phenomenon—it makes for a cleaner test of our hypothesis because it is less likely than size to be bringing in other confounding factors. But in gauging the quantitative significance of the results, it is important to recognize that, if we hold size fixed, we cannot hope to capture the full magnitude of any gradual-information-flow effect.

To be specific, return to the results for the smallest set of firms in Table V—those in the 20th–40th percentile range. Among these firms, those with the fewest analysts have momentum of 1.51 percent per month; those with the most analysts have momentum of 1.15 percent per month. Although the difference of 0.36 percent is substantial, it is still just a fraction of the total momentum effect. One reading of this might be that gradual information diffusion can only “explain” a fraction of the overall momentum in stock returns. However, such an inference is at best superficial. Recall that even the most heavily covered stocks in this class have only three or four analysts, and only average \$60 million in market capitalization. Thus they might naturally be expected to have slower information diffusion than, say, a \$10 billion company with 25 analysts. The bottom line is that residual analyst coverage, viewed in isolation, is unlikely to provide a full picture of the importance of gradual information flow. This is where the cuts on raw size in Tables III and V add potentially useful evidence.

Table VI
Momentum Strategies, 1/1980–12/1996, Using Beta-Adjusted
Returns and Sorting by Model 1 Residuals

This table includes only stocks above the NYSE/AMEX 20th percentile. The relative momentum portfolios are formed based on six-month lagged beta-adjusted returns and held for six months. The stocks are ranked in ascending order on the basis of six-month lagged returns. Portfolio P1 is an equally weighted portfolio of stocks in the worst-performing 30 percent, portfolio P2 includes the middle 40 percent, and portfolio P3 includes the best-performing 30 percent. This table reports the average monthly beta-adjusted returns of these portfolios and portfolios formed using an independent sort on Model 1 analyst coverage residuals of log size and a Nasdaq dummy. The least covered firms are in Sub1, the medium covered firms in Sub2, the most covered firms in Sub3. Mean (median) size is in millions. *t*-statistics are in parentheses.

Past	All Stocks	Residual Coverage Class			
		Low: Sub1	Medium: Sub2	High: Sub3	Sub1 – Sub3
P1	–0.00753 (–3.29)	–0.01007 (–3.97)	–0.00712 (–3.30)	–0.00511 (–2.13)	–0.00497 (–3.64)
P2	0.00280 (2.44)	0.00313 (2.48)	0.00299 (2.92)	0.00231 (1.73)	0.00081 (1.06)
P3	0.00444 (3.17)	0.00423 (2.76)	0.00454 (3.50)	0.00430 (2.74)	–0.00006 (–0.06)
P3 – P1	0.01197 (5.99)	0.01431 (6.79)	0.01167 (5.76)	0.00940 (4.62)	0.00491 (4.04)
Mean size		1070	998	464	
Median size		106	221	186	
Mean analyst		1.8	7.1	9.9	
Median analyst		0.2	4.0	7.9	

D. Sensitivities

We now discuss several variations on the baseline analysis of Table IV. First, in Table VI, we depart from Jegadeesh and Titman's (1993) focus on raw returns. Given that our economic story is all about firm-specific information, it seems sensible to focus on returns adjusted for any marketwide factors. This is also a useful precaution since, as seen in Table II, analyst coverage is correlated with beta. In Table VI all the returns, both in the preformation and postformation periods, are market-model adjusted, using individual stock betas. As it turns out, the use of this beta adjustment does not significantly alter our central results. The P3 – P1 momentum measure for the entire sample actually rises somewhat, to 1.20 percent per month (from 0.94 percent in Table IV), and the difference between the low-coverage Sub1 and the high-coverage Sub3 also goes up a bit, to 0.49 percent, with a *t*-statistic of 4.04 (from 0.42 percent in Table IV). Finally, the LAST strategy, which is long P1/Sub3 and short P1/Sub1, continues to do well, though not quite as well as before, generating an average beta-adjusted return of 0.50 percent per month (*t*-statistic = 3.64).

In a second sensitivity check, we go back to using raw returns, but generate the coverage residuals from Model 2 of Table II, which includes the 15 industry dummies. To save space, we do not report the results in a table here (see the NBER working paper version for full details) as they are not much changed. The difference in $P3 - P1$ momentum between Sub1 and Sub3 falls slightly, to 0.33 percent per month, but is still strongly significant, with a t -statistic of 3.06. As for our LAST strategy which operates only in P1, it now generates a monthly return of 0.60 percent (t -statistic = 5.03). Thus it appears that one can design a profitable LAST strategy that is not only size-neutral and momentum-neutral but beta-neutral, as well as neutral to any industry factors. This makes it all the more improbable that one can explain the substantial returns to this strategy based on any kind of risk story.²²

However, a final caveat on this point is that we have not checked whether the profits to the LAST strategy continue to be large after controlling for book-to-market effects. One might think that this correction would be relevant in light of the evidence in Table II that analyst coverage is positively correlated with book-to-market. As it turns out, though, the differences in book-to-market across Sub1 and Sub3 are too small to matter much. Using our Model 1 residuals, the median value of book-to-market is 0.57 in Sub1 and 0.69 in Sub3 (the means are 0.67 and 0.78 respectively). Based on the evidence in Fama and French (1992), this book-to-market spread corresponds to a return differential of roughly 0.10 percent per month, only a small fraction of the profits to our LAST strategy.²³

In another untabulated check (again see the NBER version for details), we do everything else the same as in Table IV except that we skip a month between the six-month ranking period and the six-month investment holding period. Jegadeesh and Titman (1993) suggest this approach as a way to check that neither bid-ask bounce nor any other high-frequency phenomenon is coloring any of the results. As it turns out, nothing changes—the numbers come out almost identical to those in Table IV.

In a similar spirit, we again redo Table IV, this time looking only at investment returns in non-January months. We do so because Jegadeesh and Titman (1993) find that small firms exhibit large negative momentum in January, and we worry that this might somehow be influencing the results. Once more, nothing much changes. Although overall momentum is noticeably higher outside of January, the Sub1 versus Sub3 differential is only marginally affected, rising from its Table IV value of 0.42 percent per month to 0.46 percent per month (t -statistic = 3.75).

²² Moskowitz (1997) argues that momentum effects are in part explained by industry factors. Whether or not this is correct *on average*, it appears that our results about *cross-sectional differences* in the power of momentum strategies are not driven by industry factors.

²³ See their Table IV (pp. 442–443), which covers the period from 1963 to 1990. Our Sub1 and Sub3 median values of book-to-market correspond roughly to the fourth and fifth deciles of their book-to-market distribution, respectively. On average, for each decile one moves between the second and the ninth, there is a 0.10 percent per month return increment.

Table VII
Momentum Strategies, 1/1984–12/1996, Using Raw Returns
and Sorting by Model 8 Residuals

This table includes only stocks above the NYSE/AMEX 20th percentile. The relative momentum portfolios are formed based on six-month lagged raw returns and held for six months. The stocks are ranked in ascending order on the basis of six-month lagged returns. Portfolio P1 is an equally weighted portfolio of stocks in the worst-performing 30 percent, portfolio P2 includes the middle 40 percent, and portfolio P3 includes the best-performing 30 percent. This table reports the average monthly returns of these portfolios and portfolios formed using an independent sort on Model 8 analyst coverage residuals of log size, a Nasdaq dummy, firm turnover, and Nasdaq dummy times firm turnover. The least covered firms are in Sub1, the medium covered firms in Sub2, the most covered firms in Sub3. Mean (median) size is in millions. *t*-statistics are in parentheses.

Past	All Stocks	Residual Coverage Class			
		Low: Sub1	Medium: Sub2	High: Sub3	Sub1 – Sub3
P1	0.00498 (1.11)	0.00190 (0.42)	0.00553 (1.26)	0.00747 (1.56)	–0.00557 (–3.58)
P2	0.01209 (3.44)	0.01126 (3.44)	0.01273 (3.67)	0.01229 (3.16)	–0.00103 (–1.00)
P3	0.01351 (3.38)	0.01210 (3.20)	0.01377 (3.54)	0.01458 (3.31)	–0.00248 (–2.11)
P3 – P1	0.00853 (4.22)	0.01020 (4.67)	0.00824 (3.92)	0.00711 (3.46)	0.00309 (2.23)
Mean size		1412	1078	442	
Median size		124	282	180	
Mean analyst		2.9	8.3	10.0	
Median analyst		0.4	4.9	7.7	

In Table VII, we again use raw returns, and this time generate the coverage residuals from Model 8 of Table II, which includes the turnover variables. But before turning to the numbers, we should point out that it is far from clear that it makes economic sense to control for turnover in this way. As noted above, it may well be that the positive correlation of coverage and turnover reflects causality running from the former to the latter: High-coverage stocks have lower adverse-selection costs of trading, and hence attract more trading volume (Brennan and Subrahmanyam (1995)). To the extent that this story is true, we *should not* use Model 8 to generate our residuals, for we would just be reducing the exogenous variation in coverage by regressing it on a noisy proxy for itself, thereby weakening the power of our tests.

However, there are other stories, according to which it is more sensible to use Model 8. To take a simple example, one might argue that our basic measure of firm size is misleading, because for some stocks the “float” (i.e., those shares that trade on a regular basis in the public market) is much

smaller than the market capitalization. And it is possible that both analyst coverage, as well as transactions costs of arbitrage, are driven primarily by float, rather than by market capitalization. In this setting, a turnover control—presumably a good proxy for float—would be warranted.

Overall, this discussion suggests that by using a turnover control, as in Table VII, we are erring on the side of being too conservative: The control may or may not make economic sense, and it potentially wastes some statistical power. We also end up sacrificing further power because of two data limitations: (1) we can only run the turnover-adjusted tests for the shorter sample period from 1984 to 1996, due to a lack of earlier turnover data on Nasdaq; and (2) we also lose roughly 12 percent of the firms—typically among the smaller ones—from our Table IV sample because of the requirement that turnover numbers be available for six months prior to the measurement of analyst coverage. With all these flags in mind, the results in Table VII are surprisingly strong. The difference in $P3 - P1$ momentum between Sub1 and Sub3 falls slightly relative to Table IV, to 0.31 percent per month, but even with the shorter sample it is still significant, with a t -statistic of 2.23. The return to the LAST strategy is now 0.56 percent per month, with a t -statistic of 3.58. The bottom line is that our results appear to be robust, even to this (possibly ill-conceived) control for the correlation between turnover and analyst coverage.

A similarly motivated check is to rerun our tests using the residuals from Model 9, which, in addition to the turnover measure, incorporates the options-listing dummy. One would not expect this to make much difference, since we have already seen in Table II that this dummy has virtually no incremental explanatory power for analyst coverage. And, indeed, the results from this untabulated set of tests are almost identical to those displayed in Table VII. Thus we can safely conclude that our inferences are not colored by any cross-sectional differences in transactions costs or short-sales constraints *that we can reasonably measure*.

Finally, in Table VIII, we break our sample into three subperiods: 1980 to 1984, 1985 to 1990, and 1991 to 1996. We then exactly repeat our baseline analysis from Table IV for each subperiod. Our results hold up well to this time disaggregation. The $P3 - P1$ momentum measure is meaningfully larger for the low-coverage Sub1 in each of the three subperiods: the difference between Sub1 and Sub3 bounces around from 0.65 percent to 0.31 percent. Even more impressive, the LAST strategy earns positive and statistically significant returns in each of the three subperiods.

In fact, the only surprise in Table VIII is that there appears to be little momentum on average in the last subperiod, which runs from 1991 to 1996. The overall point estimate for $P3 - P1$ over this period is only 0.33 percent, compared to values of 1.14 percent and 1.38 percent for the first two subperiods respectively. It is hard to say whether this just reflects noise in a short sample, or the fact that more arbitrageurs have caught on to momentum effects and are beginning to drive them out of existence. In any case, what is noteworthy from our perspective is that though the *average* degree

Table VIII
Momentum Strategies for Subperiods, 1/1980–12/1996,
Using Raw Returns and Sorting by Model 1 Residuals

This table includes only stocks above the NYSE/AMEX 20th percentile. The relative momentum portfolios are formed based on six-month lagged raw returns and held for six months. The stocks are ranked in ascending order on the basis of six-month lagged returns. Portfolio P1 is an equally weighted portfolio of stocks in the worst-performing 30 percent and portfolio P3 includes the best-performing 30 percent. This table reports the average monthly returns of these portfolios and portfolios formed using an independent sort on Model 1 analyst coverage residuals of log size and a Nasdaq dummy. The least covered firms are in Sub1, the medium covered firms in Sub2, the most covered firms in Sub3. Mean (median) size is in millions. *t*-statistics are in parentheses.

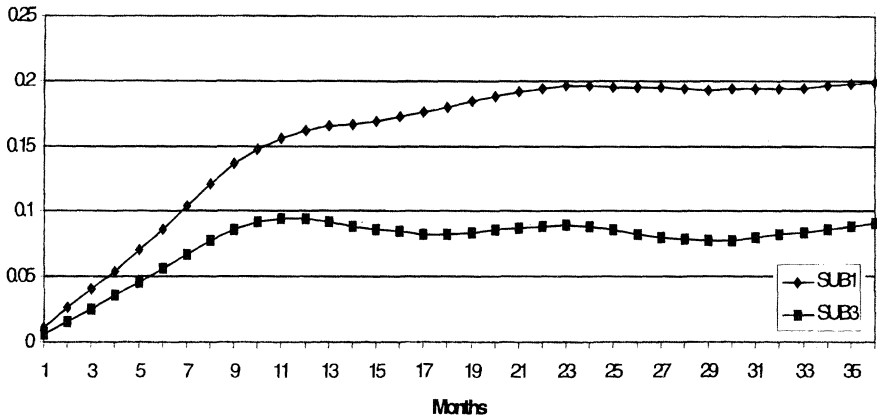
Subperiods	Past	All Stocks	Residual Coverage Class			
			Low: Sub1	Medium: Sub2	High: Sub3	Sub1 – Sub3
1/1980–12/1984	P1	0.00713 (0.94)	0.00282 (0.35)	0.00806 (1.09)	0.01215 (1.63)	–0.00933 (–3.48)
	P3	0.01852 (2.62)	0.01706 (2.30)	0.01850 (2.69)	0.01991 (2.79)	–0.00286 (–1.31)
	P3 – P1	0.01139 (3.27)	0.01424 (3.61)	0.01044 (2.88)	0.00777 (2.52)	0.00647 (2.90)
1/1985–12/1990	P1	–0.00302 (–0.41)	–0.00617 (–0.85)	–0.00205 (–0.28)	–0.00081 (–0.10)	–0.00536 (–2.21)
	P3	0.01079 (1.54)	0.00920 (1.38)	0.01164 (1.70)	0.01145 (1.50)	–0.00225 (–1.29)
	P3 – P1	0.01381 (4.65)	0.01538 (4.86)	0.01369 (4.42)	0.01227 (4.05)	0.00311 (1.62)
1/1991–12/1996	P1	0.01472 (2.49)	0.01151 (1.91)	0.01428 (2.49)	0.01828 (2.95)	–0.00677 (–3.34)
	P3	0.01805 (4.06)	0.01632 (3.79)	0.01781 (4.05)	0.01983 (4.16)	–0.00351 (–2.36)
	P3 – P1	0.00333 (0.97)	0.00481 (1.33)	0.00353 (1.02)	0.00155 (0.43)	0.00326 (1.60)

of momentum may be declining over time, there is not yet any evidence that the *cross-sectional differences* in momentum that we are emphasizing have begun to disappear.

E. Cumulative Returns in Event Time

We have focused throughout on the six-month/six-month strategy, because it has become a standard benchmark for evaluating momentum strategies. But of course this is somewhat arbitrary. To provide more information, Figure 2 plots cumulative returns in event time. In so doing, we use the methodology of Table VI—we assign stocks to performance categories based on

PANEL A



PANEL B

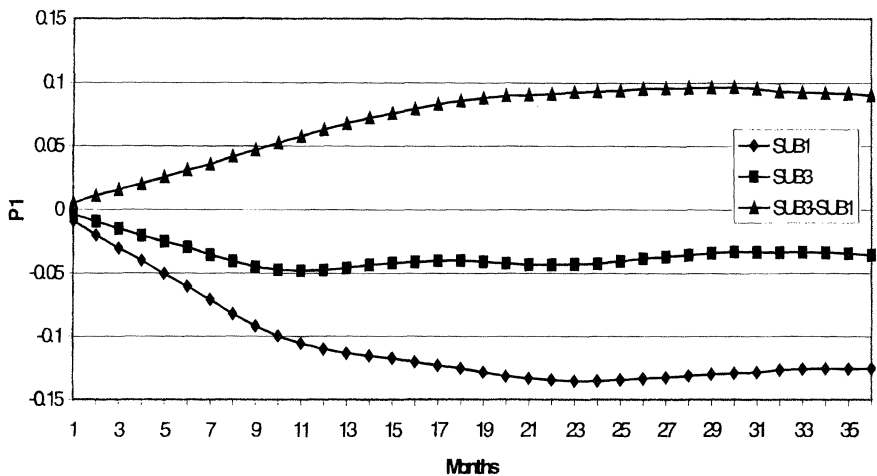


Figure 2. Cumulative beta-adjusted returns in event time. We assign stocks to performance categories based on six-months-prior beta-adjusted returns, and do an independent sort based on the analyst coverage residuals from Model 1. In Panel A we show momentum profits for low and high coverage stocks. We track the cumulative beta-adjusted momentum portfolio returns ($P3 - P1$) on a month-by-month basis, out to 36 months for low coverage (Sub1) and high coverage (Sub3) stocks. Panel B shows profits to the LAST strategy. Here we plot the cumulative beta-adjusted returns ($P1$) for the low coverage (Sub1) past losers, the high coverage (Sub3) past losers, and the LAST portfolio that is short the former and long the latter (Sub3 - Sub1).

six months' prior beta-adjusted returns, and do an independent sort based on the analyst-coverage residuals from Model 1. We then track cumulative beta-adjusted returns on a month-by-month basis, out to 36 months.

In Panel A, we plot the cumulative returns to the $P3 - P1$ momentum strategy separately for the low-coverage subsample Sub1 and the high-coverage subsample Sub3. There appear to be two distinct things going on.

First, up to about the 10-month mark, we see roughly a linear extrapolation of our earlier results: Momentum strategies continue to earn incremental monthly profits in both Sub3 and Sub1, but the effect is stronger in Sub1 so that the cumulative differential keeps on widening. After this point, something else quite interesting happens. The cumulative performance of the high-coverage subsample Sub3 flattens out; in other words, there is no more momentum left after 10 months for the high-coverage stocks. But the low-coverage subsample Sub1 continues to display some momentum out to about the two-year mark. Consequently, the cumulative differential between Sub1 and Sub3 keeps on growing until this point. Twenty-four months after portfolio formation, the total $P3 - P1$ profit for Sub1 is 19.63 percent, versus 8.90 percent for Sub3, a difference of 10.73 percent.

This dynamic pattern is, of course, completely consistent with the theory of gradual information diffusion that we have emphasized. In the context of this theory one would interpret Figure 2, Panel A, as follows: High-coverage Sub3 firms underreact by roughly nine percent to the information contained in lagged six-month returns, and it takes them a little less than a year to fully catch up. In contrast, low-coverage Sub1 firms underreact by more, on the order of 20 percent. Their adjustment to long-run equilibrium not only involves more movement in the first year, but also requires a longer period of time to fully play itself out.

In Panel B of Figure 2, we explore the dynamics of our LAST strategy. Focusing only on the past-loser stocks in P1, we plot the cumulative returns for P1/Sub1, P1/Sub3, and the LAST portfolio that is short the former and long the latter. The time profile that emerges is almost identical to that in Panel A, and is consistent with our earlier conclusion that virtually all of the Sub1 versus Sub3 action is coming from the losers in P1. In particular, the high-coverage P1/Sub3 stocks continue to perform poorly for about 10 months, and then flatten out. The low-coverage P1/Sub1 stocks not only perform worse over the first 10 months, but continue to do poorly until about two years out. Consequently, the LAST strategy keeps on earning incremental profits up to the two-year mark, with the cumulated profit amounting to 9.32 percent.

III. An Alternative, More Tightly Structured Regression Approach

In this section, we take a somewhat different approach to measuring the same basic phenomenon. In the most general terms, our central hypothesis is that stocks that are small and that have low residual analyst coverage should display more positively autocorrelated returns at medium horizons. A simple (perhaps naive) way to test this would be to estimate a serial correlation coefficient for each stock, and then regress this serial correlation coefficient on measures of the stock's analyst coverage and size.

This is what we attempt to do now. More precisely, at the beginning of each year t , we collect all stocks that have a market capitalization greater than the 20th percentile NYSE/AMEX breakpoint, and for which we have

complete return data through year $t + 5$. We then estimate for each stock i the serial correlation of its six-month excess returns (relative to T-bills), using 49 overlapping observations over the five-year period from t to $t + 5$, and call this variable RHO_{it} .²⁴ Next, we perform a cross-sectional regression, running RHO_{it} against $\log(1 + \text{Analysts}_{it})$ and $\log(\text{Size}_{it})$, as well as a Nasdaq dummy variable. All the right-hand-side variables are measured at the start of year t , so one can think of this regression as an attempt to forecast stock i 's serial correlation over the next five years.

We should note one caveat associated with this method. For any stock i , our measure of serial correlation RHO_{it} is affected not only by the correlation of its firm-specific information, but also by its loading on any common factors. To see this, suppose the returns on stock i , r_{it} , are given by a one-factor model (suppressing constants):

$$r_{it} = b_i m_t + e_{it}, \quad (1)$$

where m_t is the common factor, b_i is the loading on this factor, and e_{it} represents firm-specific information. Even if we assume for simplicity that the common factor is serially uncorrelated, ($\text{cov}(m_t, m_{t-1}) = 0$) a regression of r_{it} on r_{it-1} produces the following theoretical coefficient ρ_i^* :

$$\rho_i^* = \text{cov}(e_{it}, e_{it-1}) / (b_i^2 \text{var}(m_t) + \text{var}(e_{it})). \quad (2)$$

This suggests that, all else equal, our constructed left-hand-side variable RHO_{it} is lower for stocks with higher factor loadings—that is, higher betas. This is potentially a matter of concern because, as we have seen in Table II, there is a positive cross-sectional correlation between beta and analyst coverage. Thus one might mistakenly conclude that high coverage is reducing RHO_{it} by reducing the serial correlation of firm-specific information, when in fact it is proxying for a beta effect. In order to address this issue, we rerun the regressions that we present below, adding firm betas to the right-hand side. As it turns out, none of our results is materially altered.

Before turning to these results, it is useful to discuss how this general approach compares to what we have done above. The main difference is that it imposes more parametric structure, some of which may be unwarranted. For example, the regression approach we are now proposing does not allow for asymmetries across winners and losers; yet we have seen that such asymmetries are pronounced in the data. Additionally, the regression approach only makes sense if residual analyst coverage is a firm-level attribute that is “quasi-fixed”—that is, it does not vary much over five-year periods of

²⁴ It is well known that in a small sample one obtains downward-biased measures of serial correlation. Kendall (1954) shows that the bias is given by $-(1+3\rho)/T$, where ρ is the true value and T is the number of independent observations. This does not affect the conclusions from our cross-sectional regressions, however. We could easily rescale all our estimates of RHO_{it} to de-bias them, and none of our regression t -statistics would change.

time. If there is significant high-frequency variation in residual coverage, this is again something that the less-structured method of the previous section is better equipped to handle.

The offsetting advantage is that if the parametric structure we impose with the regression is not too inappropriate, our statistical power along certain dimensions should be enhanced. In particular, if we are interested in doing the analysis over very short intervals of time—(e.g., to check the stability of our estimates) the regression approach may be especially useful.

Table IX summarizes the results. In Panel A, we present the coefficients on the coverage and size variables from cross-sectional regressions run each year over the 14 years from 1979 to 1992.²⁵ We also aggregate the annual information in two different ways. First, we calculate Fama–MacBeth (1973) time-series averages of the coefficients. Second, we run a giant pooled regression with year dummies. Not surprisingly, this latter approach tends to produce point estimates almost identical to the Fama–MacBeth method, but higher *t*-statistics.

All the evidence in Panel A points to a consistent negative effect of analyst coverage on a stock's serial correlation. Of the yearly coefficients, 13 out of 14 are negative, the majority significantly so. The Fama–MacBeth and pooled estimates are strongly significant. The point estimates for size are also negative, but statistically insignificant.

In Panel B, we modify the specification by adding an interaction term, given by $\log(1 + \text{Analysts}) * \log(\text{Size})$. This is motivated by our evidence in Table V that the importance of analyst coverage is decreasing in firm size. The cross-sectional regressions substantiate this finding. The coverage and size terms increase in magnitude relative to Panel A (the size term is now statistically significant) and the interaction term is positive, as expected, implying that the negative influence of coverage on serial correlation becomes weaker for larger firms.

It is interesting to compare the economic magnitudes implied by Table IX to those in our earlier tables. Think of two equal-sized firms, one with the Sub1 median coverage of 0.1 (from Table IV), the other with the Sub3 median coverage of 7.6. According to the Fama–MacBeth coverage-term estimate of -0.0125 in Panel A of Table IX, the Sub1 firm should have a serial correlation coefficient that is 0.026 higher than that of the Sub3 firm ($0.0125 \times (\log(8.6) - \log(1.1)) = 0.026$). When one combines this with the observation that the past return differential between P1 and P3 stocks is approximately 60 percent, this implies that a P3 – P1 momentum strategy should be expected to return 1.56 percent more over six months for the Sub1 firm, ($0.026 \times 60 \text{ percent} = 1.56 \text{ percent}$), or about 0.25 percent per month extra. This is very much in the same ballpark as—albeit a bit smaller than—the Sub1/Sub3 differential of 0.42 percent per month reported in Table IV.

²⁵ We have to stop in 1992 because we need to go five years forward from that point to calculate RHO_{it} .

Table IX
Cross-Sectional Momentum Regressions, 1979–1992

This table includes only stocks above the NYSE/AMEX 20th percentile. The dependent variable is RHO: regression coefficient of six-month returns (net risk-free) on lagged six-month returns. Panel A: Independent variables are $\log(1 + \text{Analyst coverage})$, $\log \text{size}$, and a Nasdaq dummy. Panel B: Independent variables are $\log(1 + \text{Analyst coverage})$, $\log \text{size}$, interaction of $\log(1 + \text{Analyst coverage})$ and $\log \text{size}$ and a Nasdaq dummy. Note: t -statistics are adjusted for serial correlation.

Panel A						
Year	Coverage	t -statistics	Size	t -statistics		
79	-0.0015	-0.1800	-0.0097	-1.7530		
80	-0.0014	-0.2040	-0.0188	-3.8600		
81	-0.0039	-0.6090	-0.0061	-1.2800		
82	0.0040	0.5500	-0.0259	-4.5520		
83	-0.0136	-1.9020	0.0050	0.8990		
84	-0.0280	-3.9300	0.0168	2.9200		
85	-0.0166	-2.1330	0.0146	2.4060		
86	-0.0357	-5.6310	0.0240	4.7650		
87	-0.0111	-1.8160	0.0040	0.8480		
88	-0.0163	-2.5820	-0.0108	-2.2560		
89	-0.0141	-2.2900	-0.0071	-1.5550		
90	-0.0208	-3.2060	-0.0004	-0.0860		
91	-0.0126	-1.7100	0.0059	1.1680		
92	-0.0031	-0.4720	0.0019	0.4070		
Fama–MacBeth	-0.0125	-3.8023	-0.0005	-0.1265		
Pooled with year dummies	-0.0127	-5.0898	-0.0004	-0.2832		

Panel B						
Year	Coverage	t -statistics	Size	t -statistics	Interaction:	
					Coverage * Size	t -statistics
79	0.0306	0.7260	-0.0060	-0.8320	-0.0027	-0.775
80	0.1000	2.5930	-0.0064	-0.9480	-0.0084	-2.674
81	0.0055	0.1430	-0.0049	-0.6980	-0.0008	-0.248
82	-0.0382	-0.8500	-0.0321	-3.7080	0.0035	0.951
83	-0.0053	-0.1270	0.0061	0.7730	-0.0007	-0.205
84	-0.1441	-3.3310	-0.0001	-0.0060	0.0095	2.721
85	-0.1618	-3.3860	-0.0092	-0.9330	0.0118	3.079
86	-0.0457	-1.1950	0.0224	2.8280	0.0008	0.265
87	-0.0664	-1.8020	-0.0051	-0.6720	0.0044	1.521
88	-0.1622	-4.2580	-0.0359	-4.4700	0.0118	3.884
89	-0.0837	-2.4370	-0.0189	-2.5800	0.0057	2.059
90	-0.1372	-3.6960	-0.0212	-2.6360	0.0094	3.184
91	-0.0898	-2.2350	-0.0084	-0.9450	0.0063	1.954
92	-0.0836	-2.3560	-0.0118	-1.5720	0.0065	2.308
Fama–MacBeth	-0.0630	-1.8920	-0.0094	-2.3701	0.0041	1.5423
Pooled with year dummies	-0.0648	-5.0533	-0.0087	-3.7494	0.0043	4.4487

A similar calculation based on the interactive specification in Panel B can be used to back out the implied momentum differentials for firms in varying size classes. For example, consider the smallest class of firms (those between the 20th and 40th NYSE/AMEX percentiles) in the first column of Table V, which have a mean market capitalization of about \$60 million. Comparing a Sub1 firm in this class with median coverage of 0.0 to a Sub3 firm with median coverage of 3.1, the Fama–MacBeth coefficients in Panel B imply that a momentum strategy returns 3.91 percent more over six months for the Sub1 firm, or approximately 0.60 percent per month extra. This is again roughly in line with—although in this case somewhat larger than the analogous number of 0.36 percent reported in Table V.

Overall then, Table IX provides further comfort as to the robustness of our central results. Even with a very different measurement approach, we get not only the same qualitative outcome—higher six-month return autocorrelations among lower coverage stocks—but remarkably comparable economic magnitudes.

IV. Conclusions

Recently, a number of researchers (e.g., Barberis et al. (1998), Daniel et al. (1998), and Hong and Stein (1999)) have begun to develop behavioral models that aim to unify a range of previously documented “anomalies” in asset returns. In a critique of this work, Fama (1998) argues that one should not be too impressed if these models simply rationalize those existing patterns that they were specifically designed to capture. Rather, the acid test should be the “out-of-sample” one: The ability to generate *new* hypotheses that are ultimately borne out in future empirical work: “The over-riding question should always be: Does the new model produce coherent rejectable predictions . . .”.

We agree wholeheartedly with this sentiment, and this paper represents an attempt to take one step in the indicated direction.²⁶ The gradual-information-diffusion model of Hong and Stein (1999) was built for the express purpose of delivering both medium-term momentum and long-term reversals in stock returns; in the spirit of Fama (1998), then, it should be evaluated more on the basis of other, previously untested auxiliary predictions. Here we have focused on one relatively simple and clear-cut such hypothesis, namely: If momentum comes from gradual information flow, then there should be more momentum in those stocks for which information gets out more slowly.

Rather than restating all our findings, at this point it suffices to say that they are strongly consistent with the above hypothesis. This is not to claim that alternative interpretations of some or all of the evidence cannot be put

²⁶ A recent paper with a similar motivation is Klibanoff, Lamont, and Wizman (1998). They test the behavioral hypothesis that investors react more strongly to news that is “salient”—in this case, news about countries that appears on the front page of *The New York Times*.

forth. If concrete new alternatives are in fact offered, it will be necessary to do more refined testing to sort things out. But in any case, we hope that this effort has demonstrated at least one point: Nonclassical models of asset pricing can do more than just provide ex post rationalizations of existing anomalies; they can—and should—be subject to the same standards of out-of-sample empirical testing as more traditional theories.

REFERENCES

- Barberis, Nicholas, Andrei Shleifer, and Robert Vishny, 1998, A model of investor sentiment, *Journal of Financial Economics* 49, 307–343.
- Bernard, Victor L., 1992, Stock price reactions to earnings announcements; in R. Thaler, ed.: *Advances in Behavioral Finance* (Russell Sage Foundation, New York).
- Bernard, Victor L., and J. Thomas, 1989, Post-earnings announcement drift: Delayed price response or risk premium?, *Journal of Accounting Research*, 27, 1–48.
- Bernard, Victor L., and J. Thomas, 1990, Evidence that stock prices do not fully reflect the implications of current earnings for future earnings, *Journal of Accounting and Economics* 13, 305–340.
- Bhushan, Ravi, 1989, Firm characteristics and analyst following, *Journal of Accounting and Economics* 11, 255–274.
- Brennan, Michael J., and Patricia Hughes, 1991, Stock prices and the supply of information, *Journal of Finance* 46, 1655–1691.
- Brennan, Michael J., Narasimhan Jegadeesh, and Bhaskaran Swaminathan, 1993, Investment analysis and the adjustment of stock prices to common information, *Review of Financial Studies* 6, 799–824.
- Brennan, Michael J., and Avanidhar Subrahmanyam, 1995, Investment analysis and price formation in securities markets, *Journal of Financial Economics* 38, 361–381.
- Chan, Louis K. C., Narasimhan Jegadeesh, and Josef Lakonishok, 1996, Momentum strategies, *Journal of Finance* 51, 1681–1713.
- Conrad, Jennifer, and Gautam Kaul, 1997, An anatomy of trading strategies, *Review of Financial Studies* 11, 489–519.
- Daniel, Kent D., David Hirshleifer, and Avanidhar Subrahmanyam, 1998, Investor psychology and security market under- and overreactions, *Journal of Finance* 53, 1839–1885.
- DeLong, J. Bradford, Andrei Shleifer, Lawrence H. Summers, and Robert Waldmann, 1990, Positive feedback investment strategies and destabilizing rational speculation, *Journal of Finance* 45, 379–395.
- Diamond, Douglas, and Robert Verrecchia, 1987, Constraints on short-selling and asset price adjustment to private information, *Journal of Financial Economics* 18, 277–311.
- Easterwood, John C., and Stacey R. Nutt, 1999, Inefficiency in analysts' earnings forecasts: Systematic misreaction or systematic optimism?, *Journal of Finance* 54, 1777–1797.
- Fama, Eugene F., 1998, Market efficiency, long-term returns, and behavioral finance, *Journal of Financial Economics* 49, 283–306.
- Fama, Eugene F., and Kenneth R. French, 1992, The cross-section of expected stock returns, *Journal of Finance* 47, 427–465.
- Fama, Eugene F., and Kenneth R. French, 1996, Multifactor explanations of asset pricing anomalies, *Journal of Finance* 51, 55–84.
- Fama, Eugene F., and James D. MacBeth, 1973, Risk, return and equilibrium: Empirical tests, *Journal of Political Economy* 81, 607–636.
- Grossman, Sanford J., and Merton H. Miller, 1988, Liquidity and market structure, *Journal of Finance* 43, 617–633.
- Haugen, Robert A., and Nardin L. Baker, 1996, Commonality in the determinants of expected stock returns, *Journal of Financial Economics* 41, 401–439.

- Hayes, Rachel M., 1996, The impact of trading commission incentives on stock coverage and earnings forecast decisions by security analysts, Stanford GSB thesis.
- Hong, Harrison, Terence Lim, and Jeremy C. Stein, 1998, Bad news travels slowly: Size, analyst coverage and the profitability of momentum strategies, NBER working paper 6553.
- Hong, Harrison, and Jeremy C. Stein, 1999, A unified theory of underreaction, momentum trading and overreaction in asset markets, *Journal of Finance* 54, 2143–2184.
- Jegadeesh, Narasimhan, and Sheridan Titman, 1993, Returns to buying winners and selling losers: Implications for stock market efficiency, *Journal of Finance* 48, 65–91.
- Kendall, M. C., 1954, Note on bias in the estimation of autocorrelation, *Biometrika* 41, 403–404.
- Klibanoff, Peter, Owen Lamont, and Thierry A. Wizman, 1998, Investor reaction to salient news in closed-end country funds, *Journal of Finance* 53, 673–699.
- Lakonishok, Josef, Andrei Shleifer, and Robert Vishny, 1994, Contrarian investment, extrapolation and risk, *Journal of Finance* 49, 1541–1578.
- Lim, Terence, 1998, Rationality and analysts' forecast bias, Working paper, Amos Tuck School.
- Lo, Andrew, and A. Craig MacKinlay, 1990, When are contrarian profits due to stock market overreaction, *Review of Financial Studies* 3, 175–206.
- McNichols, Maureen, and Patricia C. O'Brien, 1996, Self-selection and analyst coverage, Working paper, London Business School.
- Merton, Robert C., 1987, A simple model of capital market equilibrium with incomplete information, *Journal of Finance* 42, 483–510.
- Moskowitz, Tobias J., 1997, Industry factors as an explanation for momentum in stock returns, Working paper, UCLA Anderson School.
- Rouwenhorst, K. Geert, 1997, Local return factors and turnover in emerging stock markets, Working paper, Yale University.
- Rouwenhorst, K. Geert, 1998, International momentum strategies, *Journal of Finance* 53, 267–284.
- Scholes, Myron, and Joseph Williams, 1977, Estimating betas from nonsynchronous data, *Journal of Financial Economics* 5, 309–328.
- Womack, Kent L., 1996, Do brokerage analysts' recommendations have investment value? *Journal of Finance* 51, 137–168.