

Cross-Sectional and Time-Series Determinants of Momentum Returns

Narasimhan Jegadeesh

University of Illinois

Sheridan Titman

University of Texas and the NBER

Portfolio strategies that buy stocks with high returns over the previous 3–12 months and sell stocks with low returns over this same time period perform well over the following 12 months. A recent article by Conrad and Kaul (1998) presents striking evidence suggesting that the momentum profits are attributable to cross-sectional differences in expected returns rather than to any time-series dependence in returns. This article shows that Conrad and Kaul reach this conclusion because they do not take into account the small sample biases in their tests and bootstrap experiments. Our unbiased empirical tests indicate that cross-sectional differences in expected returns explain very little, if any, of the momentum profits.

Portfolio strategies that buy stocks that performed well over the previous 3–12 months and sell stocks that performed poorly over this same time period have historically earned profits of about 1% per month over the following 12 months. The original results documented by Jegadeesh and Titman (1993) have subsequently been extended in several studies. For example, Rouwenhorst (1998) finds similar momentum profits in the European markets, Moskowitz and Grinblatt (1999) find momentum profits across industry-sorted portfolios, and Grundy and Martin (2001) document that momentum strategies have been consistently profitable in the United States since the 1920s.¹

The momentum literature has attracted considerable attention because the consistent profitability of the strategy poses a strong challenge to the efficient markets hypothesis. Indeed recent articles by Barberis, Shleifer, and Vishny (1998), Daniel, Hirshleifer, and Subrahmanyam (1998), and Hong and Stein (1999) develop models that appeal to behavioral biases to capture this phenomenon. In these articles, cognitive biases lead investors to either

This article has benefited from the excellent research assistance of Fei Zou and helpful comments from Gopal Basak, Eugene Fama, Ravi Jagannathan, Gautam Kaul, and participants of finance workshops at Indiana University, University of Chicago, University of Illinois, University of Texas, and Vanderbilt University. Address correspondence to Prof. Sheridan Titman, Department of Finance, College of Business Administration, University of Texas at Austin, Austin, TX 78712-1179.

¹ See also recent empirical by Chan, Jegadeesh, and Lakonishok (1996, 1999), Daniel and Titman (1999), Haugen (1999), and Hong, Lim, and Stein (1999) for additional discussion of the momentum effect.

underreact to information or follow positive feedback strategies that lead to a delayed overreaction to information. Early evidence in Jegadeesh and Titman and more recent evidence in Jegadeesh and Titman (2001) and Lee and Swaminathan (2000) support the implications of these behavioral models.

Others have argued, however, that the returns associated with momentum strategies are attributable to risk that may not have been detected with traditional measures such as the capital asset pricing model (CAPM) or the Fama and French (1993) three-factor model.² This explanation merits serious consideration, particularly in the case of momentum strategies, since the winner and loser portfolios are classified based on past returns. As Jegadeesh and Titman (1993) point out, to the extent that high past returns may be partly due to high expected returns, the winner portfolios could potentially contain high-risk stocks that would continue to earn higher expected returns in the future.

To examine this possibility, Jegadeesh and Titman calculates momentum profits within subsamples with lower dispersion in expected returns (e.g., size-based and beta-based subsamples). They find that momentum profits are not necessarily smaller within samples with lower dispersion in expected returns. Based on this evidence, Jegadeesh and Titman concludes that the dispersion in expected returns is not the source of momentum profits.

However, the idea that cross-sectional variation in expected returns can generate momentum has attracted renewed attention in the theoretical as well as the empirical literature. In particular, Chordia and Shivakumar (2000) address this issue empirically and Berk, Green, and Naik (1999) develop a theoretical model where the cross-sectional dispersion in risk and expected returns generate momentum profits. In addition, Conrad and Kaul (1998) examine this possibility in great detail and provide empirical results and simulations that lead them to conclude that most, and perhaps all of the observed momentum profits are explained by cross-sectional differences in expected returns rather than any "time-series patterns in stock returns." In contrast to Berk, Green, and Naik (1999) and Chordia and Shivakumar (2000), who consider cases where momentum is generated by time-varying expected returns, Conrad and Kaul claim that it is the cross-sectional dispersion in unconditional expected returns that generate momentum profits.

This article presents a direct test of the Conrad and Kaul hypothesis that momentum profits are due to cross-sectional differences in unconditional expected returns. Our results indicate that differences in unconditional expected returns explain very little, if any, of the momentum profits. As we show, the difference between the Conrad and Kaul results and our results is due to small sample biases in their empirical tests.

² Jegadeesh and Titman (1993) use the CAPM benchmark and Fama and French (1996) and Grundy and Martin (2001) use the Fama and French three-factor model benchmark to adjust for cross-sectional differences in risk and they both conclude that these benchmarks do not explain momentum profits.

Conrad and Kaul present some empirical evidence and support these tests with a number of simulation and bootstrap experiments. These experiments seemingly suggest that the magnitude of momentum profits found in the actual data can be obtained with randomly generated data constructed to have no time-series dependence. However, we show here that their bootstrap experiment and their simulations contain a small sample bias that is *identical* to the bias in their empirical tests. We present a variation of the Conrad and Kaul bootstrap that we analytically show is unbiased. In this unbiased bootstrap experiment we find that the momentum profits are virtually zero. Therefore our findings indicate that the Conrad and Kaul bootstrap results can be entirely attributed to small sample bias.

The rest of this article is organized as follows: Section 1 describes the trading strategy and reviews how the profits of the strategy can be decomposed into time-series and cross-sectional components. Section 2 presents our empirical tests. Section 3 describes the small sample bias in the Conrad and Kaul (1998) bootstrap experiments and presents an unbiased experiment. Section 4 presents a further assessment of the magnitude of cross-sectional variance in expected returns. Section 5 concludes.

1. The Components of Momentum Profits

1.1 The trading strategy

Momentum strategies attempt to exploit continuation in stock returns. The trading strategy in Jegadeesh and Titman buys the decile of stocks with the highest past returns and sells the decile of stocks with the lowest past returns. The stocks in the buy and sell portfolios are equally weighted in the Jegadeesh and Titman strategy. To understand the sources of momentum profits, however, it is analytically more convenient to consider the weighting scheme originally proposed by Lo and MacKinlay (1990), and also used by Conrad and Kaul. This strategy, which we label the weighted relative strength strategy (WRSS), buys stocks in proportion to their returns over the ranking period. Specifically, under this strategy each stock is assigned a weight at time t given by

$$w_{i,t} = \frac{1}{N}(r_{i,t-1} - \bar{r}_{t-1}),$$

where N is the number of stocks in the sample, $r_{i,t-1}$ is the return of stock i during the ranking period $t-1$ and $\bar{r}_{t-1}$ is the average return across all stocks in the sample at time $t-1$.³ Under this weighting scheme, the average weight across all stocks is zero, but the sum of the weights for the long and short positions vary month-to-month depending on past return realizations. Since our empirical test focuses on six-month momentum strategies the length of

³ The returns from the weighted relative strength strategy are very highly correlated with the returns associated with the decile strategies examined by Jegadeesh and Titman (1993) and others.

each period can be thought of as six months. The six-month trading strategy was the primary focus of the original Jegadeesh and Titman (1993) study and the results for this strategy are representative of that for other strategies with formation and holding periods ranging from 3 to 12 months.

The profit from this strategy, denoted as π_t , can be expressed as

$$\pi_t = \frac{1}{N} \sum_{i=1}^N r_{i,t} (r_{i,t-1} - \bar{r}_{t-1}). \quad (1)$$

1.2 The decomposition

To decompose the WRSS profit, the realized return for stock i is expressed as

$$r_{i,t} = \mu_i + u_{i,t}, \quad (2)$$

where μ_i is the unconditional expected return of stock i and $u_{i,t}$ is the unexpected return at time t . The momentum profits in Equation (1) can now be decomposed into components based on expected and unexpected components of returns as follows:

$$\pi_t = -\text{cov}(\bar{r}_t, \bar{r}_{t-1}) + \frac{1}{N} \sum_{i=1}^N \text{cov}(r_{i,t}, r_{i,t-1}) + \sigma_\mu^2, \quad (3)$$

where σ_μ^2 is the cross-sectional variance of expected returns. Lo and MacKinlay (1990) originally proposed this decomposition to investigate the source of short-horizon contrarian profits documented by Jegadeesh (1990) and Lehmann (1990). Their focus was largely on the first and second terms on the right-hand side of Equation (3).⁴

The focus of Conrad and Kaul (1998) and this article is on the contribution of the cross-sectional variance of expected returns to momentum profits, which is captured by the last term on the right-hand side. This term indicates that any cross-sectional variation in expected returns contributes positively to momentum profits. This relation is intuitive since momentum strategies sort stocks based on realized past returns, which are positively correlated with expected returns. If a large part of realized returns is due to expected returns then past winners will tend to be stocks with higher than average expected returns and past losers will tend to be stocks with lower than average expected returns. Given this, past winners (losers) will on average continue to earn higher (lower) than average returns in the future.

If expected returns of individual stocks were observable, one could easily gauge the contribution of cross-sectional differences in expected returns to momentum profits. However, since expected returns cannot be directly

⁴ See Jegadeesh and Titman (1995) for a more detailed discussion of this decomposition and its economic interpretation.

Table 1
Weighted relative strength strategy (WRSS) profits (1965–1997)

	WRSS profit as percentage of long position (δ)	WRSS profit (π) ($\times 10^2$)	Variance of sample mean returns – (all firms) ^a ($\times 10^2$)	Variance of sample mean returns – (firms with at least five years of data) ^b ($\times 10^2$)
Average	4.06 (3.73)	.366 (3.30)	1.721	.471
Fraction of WRSS profits	—	—	4.72	1.29

The WRSS is a zero net investment strategy where each stock in the portfolio is assigned an equal weight proportional to the difference between its lagged six-month returns and cross-sectional mean returns. The momentum profit for month t (π_t) is computed as

$$\pi_t = \frac{1}{N} \sum_{i=1}^N r_{i,t}(r_{i,t-1} - \bar{r}_{t-1}),$$

where N is the number of stocks in the cross-section, r_{it} is the return of stock i in six-month period t to $t+5$ and r_{it-1} is the return in the six-month period $t-1$ to $t-6$. $\bar{r}_{t-1}$ is the average return across all stocks in the sample.

WRSS profit as percentage of long position (δ) is the WRSS profit scaled by the value of the long position each month and expressed in percentage per six-month holding period.

The sample includes all firms listed on the New York Stock Exchange and Amex. Cross-sectional variance of mean returns is the variance of sample average six-month returns. Autocorrelation consistent t -statistics are presented in parentheses.

^aCross-sectional variance of sample mean returns across all firms in the sample.

^bCross-sectional variance of sample mean returns across firms with at least five years of return data.

observed, it must be estimated from realized returns. Conrad and Kaul use the average realized return of each stock as their measure of the stock's expected return. Formally this estimate of expected return is

$$\hat{\mu}_i = \frac{1}{T_i} \sum_{t=1}^{T_i} r_{i,t}, \quad (4)$$

where T_i is the number of observations available for stock i . They use the cross-sectional variance of $\hat{\mu}_i$ as the estimator of σ_μ^2 .

2. Empirical Results

Table 1 presents the momentum profits with a sample of all stocks listed on the New York Stock Exchange or the American Stock Exchange over the 1965–1997 sample period.

The average WRSS profit is $.366 \times 10^{-2}$, which is reliably different from zero. This profit is earned by the zero investment WRSS, where the dollar investment in the long and short sides of the portfolio change over each holding period depending on the realized returns over the ranking period. To provide an economic perspective on this profit, we also compute the profit as a percentage of the long position over each holding period. The average profit is 4.06% (per six-month holding period) of the long side of the portfolio.

2.1 Components of momentum profits

The cross-sectional variance of $\hat{\mu}_i$, which we report in Table 1, is 1.721×10^{-2} , which is more than four times the WRSS profits. If the cross-sectional

variance of expected returns were really this high, we would expect to observe even larger momentum profits than we actually observe. Such inference, however, ignores the impact of the error in the estimates of $\hat{\mu}_i$ on the estimate of σ_μ^2 . To see this, let

$$\hat{\mu}_i = \mu_i + \varepsilon_i,$$

where ε_i represents estimation error. Since $\hat{\mu}_i$ is an unbiased estimator of expected returns, $E(\varepsilon_i) = 0$. However, since

$$\sigma_{\hat{\mu}_i}^2 = \sigma_{\mu_i}^2 + \sigma_{\varepsilon_i}^2, \quad (5)$$

the variance of the estimated expected returns overestimates the cross-sectional variance of true expected returns. The magnitude of this overestimation is exacerbated when we follow Conrad and Kaul and use *all* stocks in the sample period for the calculation of expected returns, regardless of the length of their return history. Indeed, a number of stocks in the CRSP sample have return series that are only 12 months long and hence have expected return estimates based on Equation (4) that are quite imprecise.

To obtain somewhat more precise estimates we restrict the sample to only firms that had at least five years of returns data. The cross-sectional variance of mean returns for this sample is $.471 \times 10^{-1}$, which is still larger than the momentum profits but somewhat closer to the cross-sectional variance of $.387 \times 10^{-2}$ that Conrad and Kaul report for the 1962–1989 sample period with all stocks (their Table 2). Conrad and Kaul obtain a smaller estimate than we do because they follow a slightly different approach. We compute a single cross-sectional variance of mean returns across all stocks. Conrad and Kaul, on the other hand, first compute a cross-sectional variance for each individual month across stocks in the sample for that month and then report the time-series average of these estimates.

The fact that different estimators lead to widely divergent point estimates of cross-sectional expected return variance is primarily because each estimator weighs measurement errors differently. The Conrad and Kaul method generates lower estimates since firms with fewer return observations, which typically have larger absolute measurement errors, enter the sample fewer times, and hence get a lower overall weight in their estimator. As we show later, however, all of these estimators vastly overstate the contribution of differences in expected returns to momentum profits.

To evaluate the extent to which the measurement error in the sample mean could potentially bias the estimates of dispersion in true expected returns, Figure 1 plots the distribution of sample mean returns. Estimated expected returns are negative for nearly 20% of the stocks in the sample, while for several others, the average returns exceed 100%, suggesting that these sample average returns are unlikely to capture the true *ex ante* expected returns. These measurement errors therefore bias upwards the Conrad and

Table 2
Momentum profits adjusted for sample average returns

Mean return estimator	WRSS profit (π) ($\times 10^2$)	WRSS profit as percentage of long position (δ)
No mean return adjustment	.303 (2.60)	3.33 (2.94)
Preranking period	.305 (2.54)	3.30 (2.82)
Postholding period	.321 (2.43)	3.46 (2.79)
Preranking and post- holding periods	.359 (2.93)	3.91 (3.30)

This table presents WRSS profit after adjusting for expected returns estimated in different sample periods. The row titled "Preranking period" computes expected returns for each stock as the sample average returns in the period from first month when returns data are available for a particular firm to the month prior to the ranking period. The row titled "Postholding period" computes expected returns for each stock as the sample average returns in the period from 13 months after the holding period to the last month when returns data are available for a particular firm. The last row computes expected returns for each stock as the sample average returns in the preranking and postholding periods. See Table 1 for description of the WRSS. Firms are included in the sample only when sufficient returns are available for computation of expected returns in both pre- and postranking periods. Because of data availability requirement for computing postholding period average returns, the sample period in this table is 1965–1996. For direct comparison the WRSS profits with no mean return adjustment for the sample of stocks in this table is also presented. Autocorrelation consistent *t*-statistic is presented in parentheses. WRSS profit as percentage of long position (δ) is the WRSS profit scaled by the value of the long position each month and expressed in percentage per six-month holding period.

Kaul estimate of the extent to which cross-sectional differences in returns contribute to momentum profits.

2.2 Direct tests of the risk hypothesis

Given the difficulties associated with obtaining an accurate estimate of the cross-sectional variance of expected returns, it is probably impossible to directly measure the different components of momentum profits based on the decomposition given by Equation (3). However, it is still possible to

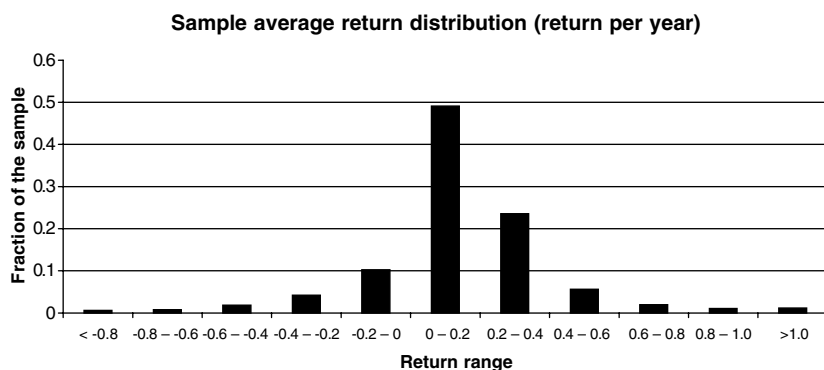


Figure 1
Distribution of annual mean returns for all NYSE and AMEX stocks in the 1965–1997 period

This figure presents the fraction of the sample average returns that fall within the annualized return range on the x-axis.

directly investigate the extent to which the returns of the momentum strategy can potentially be due to cross-sectional differences in expected return. As we mentioned in the introduction, a number of authors have concluded that traditional single-factor and multi factor models fail to explain the observed momentum returns. As Conrad and Kaul correctly point out, these methods do not account for all risk factors. Fortunately there is no need to use a specific asset pricing model to obtain unbiased estimates of expected returns if we are willing to impose the assumption made by Conrad and Kaul that expected returns are constant over time. In this case we can use sample average returns, as described in Equation (4), as an unbiased estimate of expected returns. As we will see below, our application here only requires that the expected return estimator be unbiased and we do not require the cross-sectional variance of sample mean returns to be an unbiased estimator of the cross-sectional variance of true expected returns.

We compute sample average returns for each firm over three different sample periods. The first measure is the ex ante sample means computed over the preranking period. Specifically the ex ante sample mean for the ranking period ending in month t is

$$\hat{\mu}_{i,t}(1) = \frac{1}{(t - 6 - T_{i,pre} + 1)} \sum_{j=T_{i,pre}}^{t-6} r_{i,j},$$

where $T_{i,pre}$ is the first date on or after January 1963 for which returns data for firm i can be obtained from the CRSP database. The ex ante sample period excludes the ranking period, since by construction the returns are low for losers and high for winners in this period.

The advantage of using the ex ante ranking period returns is that this estimate is available to the investor at the time of portfolio formation. However, since losers experience decreases in their equity values and winners experience gains over the ranking period, it is likely that their financial leverage changes during this period, resulting in changes in their future risk exposures. Since ex ante means do not account for the effect of the changes in risk, our second measure provides an estimate of expected returns using sample means in the postholding periods. Specifically, for each stock we compute the postholding period sample mean as follows for ranking periods ending in month t :

$$\hat{\mu}_{i,t}(2) = \frac{1}{(T_{i,post} - (t + 13) + 1)} \sum_{j=t+13}^{T_{i,post}} r_{i,j},$$

where $T_{i,post}$ is the last month for which return data for firm i is available. Since we are trying to explain the returns in the holding period, we exclude the 12-month period following the ranking period to calculate these averages.

Our final estimate of expected returns is computed over both preranking and postholding periods. Specifically,

$$\hat{\mu}_{i,t}(3) = \frac{1}{(T_{i,post} - (t+13) + 1) + (t - 6 - T_{i,pre} + 1)} \left(\sum_{j=T_{i,pre}}^{t-6} r_{i,j} + \sum_{j=t+13}^{T_{i,post}} r_{i,j} \right).$$

Using these estimators, we compute the abnormal returns $r_{i,t}^{ab}(k)$ for each stock during the holding period as

$$r_{i,t}^{ab}(k) = r_{i,t} - \hat{\mu}_{i,t}(k), k = 1, 2, 3.$$

Table 2 presents the momentum profits adjusted for the different measures of expected returns. To make the various expected return adjustment procedures directly comparable, we now use the sample of stocks for which data to compute all three measures of expected returns are available. The sample period now ends in December 1996, since we need postholding period returns to estimate $\hat{\mu}_{i,t}(2)$. The average unadjusted momentum profit with this sample of stocks over the 1965–1996 sample period is 3.33% over the six-month holding period. The momentum profits after adjusting for various measures of expected returns ranges from 3.30% to 3.91%.⁵ If anything, the momentum profits *increase* after adjusting for various measures of expected return. Of interest, consistent with our results, Jegadeesh and Titman (1993, 2000) and Fama and French (1996) also find that the abnormal momentum returns increase marginally after adjusting for risk under the CAPM and the Fama and French three-factor model. Overall the results indicate that virtually none of the momentum profits can be attributed to compensation for risk.

3. The Bootstrap and Simulation Puzzle

One puzzle that still remains is the bootstrap and simulation evidence presented by Conrad and Kaul, which leads them to conclude that the “main determinant of the profits of return-based trading strategies is the cross-sectional variation in mean returns.” The Conrad and Kaul bootstrap results are indeed in direct conflict with the empirical evidence provided in the last section. Bootstrap experiments are quite commonly used in the finance literature and some variations of the experiments in Conrad and Kaul have also been used in different contexts by Brock, Lakonishok, and Le Baron (1992), Karoyli and Cho (1993), and Conrad and Kaul (1999). It is therefore important to understand the properties of these bootstrap experiments and reconcile

⁵ Note that the profits after adjusting for combined period mean returns ($\hat{\mu}_{i,t}(3)$) in an unweighted average of the momentum profits after adjusting separately for preranking period and postholding period mean returns.

our findings in the last section with the results implied by the Conrad and Kaul tests.

The Conrad and Kaul bootstrap experiment is designed as follows: They first scramble the monthly returns for each stock by randomly drawing a return for each month, *with replacement*, from the observed distribution of the stock's returns. This scrambling process should eliminate any serial correlation in the return series.

3.1 A small sample bias

Conrad and Kaul report that their bootstrap experiment generates the same momentum profits as the actual return series. This evidence appears to be inconsistent with the idea that momentum profits are generated by the time-series pattern of returns since any time-series dependence is eliminated when the return data are scrambled. However, as we show below, the observed momentum returns in the scrambled time series arises because of a small sample bias.⁶ This bias arises because when returns are drawn *with replacement*, the *same* return observation for a stock can be drawn in both the ranking and the holding periods. This bias can be easily illustrated with a simple example where one stock realizes a particular high return in one month (say 1000%). In any particular month that this return is drawn, because of its extreme return, the stock will fall in the winner portfolio. If this particular observation is drawn in an adjacent six-month period, which is not unlikely given the length of the return time series, the simulation will spuriously show high returns for the momentum strategy in the holding period.

To illustrate this bias, let $r_{i,t}$ be the return realization for a stock i at time t . Under the null hypothesis the return series is serially uncorrelated. Now consider the Conrad and Kaul bootstrap experiment that draws returns *with replacement* and uses the bootstrapped returns to simulate a momentum strategy that forms portfolios based on $t - 1$ returns, and buys winners and sells losers that are held in period t . Denote the returns drawn in this bootstrap experiment by $r_{i,t}^*$. When returns are drawn at random with replacement, the probability of drawing any return observation from the time series is $\frac{1}{T_i}$, where T_i is the number of time-series observations for stock i . Therefore, given the actual return data, the expected value of the return drawn for the holding period is

$$E_{rep,t}(r_{i,t}^* | r_{i,t}, t = 1, \dots, T_i) = \frac{1}{T_i} \sum_{j=1}^{T_i} r_{i,j}^* = \hat{\mu}_i,$$

where, as defined in Equation (4), $\hat{\mu}_i$ is the *sample* average returns.

⁶ The potential for biases in bootstrap estimators has been discussed extensively in the statistics literature [see, e.g., Efron and Tibshirani (1993)].

Let π_{rep}^* denote the momentum profit in the bootstrap experiment where returns are drawn with replacement. The expectation of π_{rep}^* given the original data is

$$\begin{aligned} E(\pi_{rep,t}^* \mid r_{i,j}, i = 1, \dots, N, j = 1, \dots, T_i) &= E\left(\frac{1}{N} \sum_{i=1}^N r_{i,t}^* (r_{i,t-1}^* - \bar{r}_{t-1}^*)\right) \\ &= \hat{\mu}_i^2 - \bar{\mu}^2 = \sigma_{\hat{\mu}}^2, \end{aligned} \quad (6)$$

where $\bar{\mu} = \frac{1}{N} \sum_{i=1}^N \hat{\mu}_i$. We have imposed the null hypothesis here that the average sample serial covariance of return equal zero. As we discussed earlier, in small samples, $\sigma_{\hat{\mu}}^2$ overestimates the cross-sectional variance in true expected returns. This is of course a small sample bias, since as $T \rightarrow \infty$, the cross-sectional variance of the sample mean tends to its population counterpart.

It is noteworthy that the expected value of the momentum profit that is generated by the Conrad and Kaul bootstrap experiment, described in Equation (6), is *identical* to the estimate of the first component of momentum profits one would obtain using the cross-sectional variance of sample average returns in place of the cross-sectional variance of true expected returns. Conrad and Kaul present the bootstrap experiment with replacement as a robustness check for their empirical tests and state that they use this experiment “to address the potentially serious effects of measurement errors in in-sample mean returns.” However, as the results here indicate, the Conrad and Kaul bootstrap experiment is actually subject to the identical measurement error problem. This explains why the cross-sectional variance of sample returns that Conrad and Kaul report in their Table 2, $.387 \times 10^{-2}$, is so close to the six-month momentum profit of $.378 \times 10^{-2}$ in their bootstrap experiment that they report in their Table 3. These numbers are not exactly equal because of sampling variations. However, as we have shown here, the fact

Table 3
Weighted relative strength strategy (WRSS) profits and cross-sectional dispersion in mean returns: simulation evidence

	WRSS profits ($\times 10^2$)	<i>t</i> -statistics	Percentage of momentum profits in actual data
Actual data (1965–1997)	0.366	3.30	—
Simulation, <i>without</i> replacement	0.002	0.05	.54

The WRSS is a zero net investment strategy where each stock in the portfolio is assigned a weight equal proportional to the difference between its lagged six-month returns and cross-sectional mean returns. See Table 6 for further details. Autocorrelation consistent *t*-statistic is presented for profits based on actual returns data. The simulation with replacement computes the WRSS profits based on 500 simulations where in each simulation the time series of returns for each stock is scrambled with replacement. The last column reports simulated WRSS profits when returns are scrambled without replacement.

that they are so close is not a mere coincidence—both these estimates are mathematically identical except for sampling variations.⁷

3.2 An unbiased bootstrap experiment

Consider now a bootstrap experiment, which is identical to the CK experiment except that now returns are drawn *without* replacement. In this experiment, let $\tilde{r}_{i,t}$ denote the time series of returns and let $\pi_{no_rep}^*$ denote the momentum profit. The expectation of $\pi_{no_rep}^*$ given the original data is

$$\begin{aligned} E(\pi_{no_rep,t}^* | r_{i,j}, i=1, \dots, N, j=1, \dots, T_i) &= E\left(\frac{1}{N} \sum_{i=1}^N (\tilde{r}_{i,t}(\tilde{r}_{i,t-1} - \tilde{r}_{t-1}))\right) \\ &= E\left(\frac{1}{N} \sum_{i=1}^N (r_{i,j} r_{i,k \neq j} - r_{i,j} \tilde{r}_{t-1})\right) \\ &= \frac{1}{N} \sum_{i=1}^N \mu_i^2 - \bar{\mu}^2 = \sigma_\mu^2. \end{aligned}$$

The first equality on the second line follows by imposing the null hypothesis that the average sample serial covariance equals zero. Therefore momentum profits in this bootstrap experiment are an unbiased estimator of the contribution of σ_μ^2 to momentum profits.

We carry out 500 replications of the without replacement bootstrap experiment. Table 3, which reports the result of this experiment, indicates that the average momentum profit is virtually zero. Therefore the cross-sectional differences in expected returns contribute very little to momentum profits.

3.3 Simulation experiment

Conrad and Kaul also carry out a Monte Carlo simulation experiment as a further robustness check of their results. In these simulations they generate returns from “independent and identical normal distributions that have means and variances that are identical to those observed in the real data.” They show that the momentum profits with this simulated data are close to their estimate of cross-sectional variance of sample average returns.

The inherent bias in this experiment is perhaps less subtle than that in their bootstrap experiment. The simulation generates serially uncorrelated returns. Therefore, from Equation (3) it is clear that the only source of profit is the cross-sectional variance of returns that is built into the simulation. Since

⁷ The Conrad and Kaul estimates of cross-sectional variance of three-month and nine-month returns are $.098 \times 10^{-2}$ and $.849 \times 10^{-2}$, respectively (see their Table 2). The momentum profits in their bootstrap experiments for these horizons are $.099 \times 10^{-2}$ and $.841 \times 10^{-2}$, respectively (see their Table 3). Given our results, one can expect ex ante that the momentum profits in the bootstrap experiment will be very close to the cross-sectional variance of sample average returns for the corresponding horizons.

the cross-sectional variance of returns in their simulation equals the cross-sectional variance of *sample* mean returns, the momentum profits will turn out to be equal to the cross-sectional variance of *sample* mean returns. In fact, this is close to what they find in their Table 3, and any difference between the momentum profits in this table and the cross-sectional variance of returns used in their simulations (see their Table 2 for cross-sectional variances at different horizons) can only be attributed to sampling variation.

These identities apart, it is also interesting to consider the mean returns used in the simulation. Figure 1 presents the distribution of the means observed in real data, and in the Conrad and Kaul simulation the expected stock returns match this distribution. As this figure illustrates, some stocks in the Conrad and Kaul simulations have “expected” returns of less than -80% per year and some stocks have expected returns in excess of 100% per year. The variance of this “expected return” distribution is clearly much too large. Nevertheless, these simulations are of interest because they illustrate how large the variations in expected returns would need to be to account for the observed magnitude of momentum profits.

4. How Large is the Cross-Sectional Variance of Expected Returns?

Before concluding the article we provide some rough estimates of the cross-sectional variance of expected returns and consider their possible contribution to observed momentum profits. These estimates are calculated from the evidence provided in Fama and French (1992). To obtain our first estimate we examine the beta estimates of the 12 beta-sorted portfolios in Table II of Fama and French. The cross-sectional standard deviation of these betas is .31. Suppose the CAPM holds and suppose the market risk premium is 6% per year. These parameters imply that for six-month expected returns the cross-sectional variance, σ_μ^2 , is 8.6×10^{-5} , which is 2.3% of the momentum profits reported in Table 1. This may, however, be an overestimate, since empirically the cross-sectional dispersion in expected returns is smaller than that implied by the CAPM. In fact, the six-month cross-sectional variance of average returns of the 12 beta-sorted portfolios in Table II of Fama and French is only 0.05% the momentum profits.

Alternatively, suppose that a large proportion of the cross-sectional variance of average returns is determined by differences in book-to-market ratios and market capitalizations. Then σ_μ^2 can be estimated from the average returns of size and book-to-market sorted portfolios. For example, the σ_μ^2 across the 100 book-to-market and size portfolios reported in Table V of Fama and French (1992) is 5.8×10^{-5} , which is 1.6% of momentum profits in Table 1. This estimate is slightly larger than our estimate of 0.05% in Table 3, probably due to sampling errors in average portfolio returns. In any event, these estimates also indicate that the contribution of cross-sectional differences in expected returns to momentum profits is likely to be trivial.

5. Conclusion

The observed profitability of momentum strategies has attracted considerable attention because it appears to be inconsistent with market efficiency. However, a number of authors have noted that momentum profits can be generated in an efficient market as long as the cross-sectional variation in unconditional expected returns is large relative to the variation in unexpected returns. If this were the case, then past winners are likely to consist primarily of stocks with high expected returns and past losers are likely to consist of stocks with low expected returns, implying positive expected returns for the momentum strategy.

Although the cross-sectional variation in expected returns can, in theory, account for the observed momentum profits, we conclude that its contribution is likely to be quite small in practice. Intuitively this is because the cross-sectional variation in unconditional expected returns is small relative to the variation in realized returns and a stock's realized return over any six-month period provides very little information about the stock's unconditional expected return. Hence the unconditional expected return of past winners is unlikely to be significantly different from that of past losers. Our empirical tests support this intuition.

More generally, given the growing list of return anomalies, it has become increasingly important to calibrate various models and perform simulations that allow us to gauge the magnitudes of different factors that can potentially be responsible for any apparent excess returns. As one of the first to seriously consider such a calibration exercise, Conrad and Kaul (1998) make an important contribution to this literature. However, they overlook some important small sample biases in their estimates and as a result draw erroneous inferences. As we illustrate in this article, simulation experiments can potentially be quite fragile; relatively minor errors in the experimental design can often have profound implications on the conclusions.

References

- Barberis, N., A. Shleifer, and R. Vishny, 1998, "A Model of Investor Sentiment," *Journal of Financial Economics*, 49, 307–343.
- Berk, J., R. Green, and V. Naik, 1999, "Optimal Investment, Growth Options, and Security Returns," *Journal of Finance*, 54, 1553–1607.
- Brock, W., J. Lakonishok, and B. LeBaron, 1992, "Simple Trading Rules and Stochastic Properties of Stock Returns," *Journal of Finance*, 47, 1731–1764.
- Chan, L. K., N. Jegadeesh, and J. Lakonishok, 1996, "Momentum Strategies," *Journal of Finance*, 51, 1681–1713.
- Chan, L. K., N. Jegadeesh, and J. Lakonishok, 1999, "The Profitability of Momentum Strategies," *Financial Analyst Journal*, 55, 80–90.
- Chordia, T., and L. Shivakumar, 2000, "Momentum, Business Cycle and Time Varying Expected Returns," working paper, London Business School.

- Conrad, J., and G. Kaul, 1998, "An Anatomy of Trading Strategies," *Review of Financial Studies*, 11, 489–519.
- Conrad, J., and G. Kaul, 1999, "Value vs. Growth," working paper, University of Michigan.
- Daniel, K., D. Hirshleifer, and A. Subrahmanyam, 1998, "Investor Psychology and Security Market Under- and Overreactions," *Journal of Finance*, 53, 1839–1886.
- Daniel, K., and S. Titman, 1999, "Market Efficiency in an Irrational World," *Financial Analyst Journal*, 55, 28–40.
- Efron, B., and R. J. Tibshirani, 1993, *An Introduction to the Bootstrap*, Chapman & Hall, New York.
- Fama, E., and K. French, 1992, "The Cross-Section of Expected Stock Returns," *Journal of Finance*, 47, 427–465.
- Fama, E., and K. French, 1993, "Common Risk Factors in the Returns on Stocks and Bonds," *Journal of Financial Economics*, 33, 3–56.
- Fama, E., and K. French, 1996, "Multifactor Explanations of Asset Pricing Anomalies," *Journal of Financial Economics*, 51, 55–84.
- Grundy, B. D., and S. J. Martin, 2001, "Understanding the Nature of Risks and the Sources of Rewards to Momentum Investing," *Review of Financial Studies*, 14, 29–78.
- Haugen, R. A., 1999, *The New Finance*, Prentice Hall, Upper Saddle River, NJ.
- Hong, H., and J. Stein, 1999, "A Unified Theory of Underreaction, Momentum Trading and Overreaction in Asset Markets," *Journal of Finance*, 54, 2143–2184.
- Hong, H., T. Lim, and J. C. Stein, 2000, "Bad News Travels Slowly: Size, Analyst Coverage, and the Profitability of Momentum Strategies," *Journal of Finance*, 55, 265–295.
- Jegadeesh, N., 1990, "Evidence of Predictable Behavior of Security Returns," *Journal of Finance*, 45, 881–898.
- Jegadeesh, N., and S. Titman, 1993, "Returns to Buying Winners and Selling Losers: Implications for Stock Market Efficiency," *Journal of Finance*, 48, 65–91.
- Jegadeesh, N., and S. Titman, 1995, "Overreaction, Delayed Reaction and Contrarian Profits," *Review of Financial Studies*, 8, 973–993.
- Jegadeesh, N., and S. Titman, 2001, "Profitability of Momentum Strategies: An Evaluation of Alternative Explanations," *Journal of Finance*, 56, 699–720.
- Karolyi, G., and B. Cho, 1993, "Time-Varying Risk Premia and the Returns to Buying Winners and Selling Losers: Caveat Emptor et Veditor," working paper, Ohio State University.
- Lakonishok, J., A. Shleifer, and R. W. Vishny, 1994, "Contrarian Investment, Extrapolation and Risk," *Journal of Finance*, 49, 1541–1578.
- Lee, C., and B. Swaminathan, 2000, "Price Momentum and Trading Volume," forthcoming in *Journal of Finance*.
- Lehmann, B., 1990, "Fads, Martingales and Market Efficiency," *Quarterly Journal of Economics*, 105, 1–28.
- Lo, A., and A. C. MacKinlay, 1990, "When are Contrarian Profits Due to Stock Market Overreaction?," *Review of Financial Studies*, 3, 175–208.
- Merton, R. C., 1980, "On Estimating the Expected Return on the Market: An Exploratory Investigation," *Journal of Financial and Quantitative Analysis*, 7, 323–361.
- Moskowitz, T., and M. Grinblatt, 1999, "Does Industry Explain Momentum?" *Journal of Finance*.
- Rouwenhorst, K. G., 1998, "International Momentum Strategies," *Journal of Finance*, 53, 267–284.

Copyright of Review of Financial Studies is the property of Society for Financial Studies and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.