



Time series momentum: Is it there?☆

Dashan Huang^{a,*}, Jiangyuan Li^a, Liyao Wang^b, Guofu Zhou^{c,d}^a Lee Kong Chian School of Business, Singapore Management University, 50 Stamford Road, 178899, Singapore^b School of Economics, Singapore Management University, 90 Stamford Road, 178903, Singapore^c Olin School of Business, Washington University in St. Louis, 1 Brookings Drive, St. Louis, MO 63130, USA^d China Academy of Financial Research (CAFR), 211 West Huaihai Road, Shanghai 200030, China

ARTICLE INFO

Article history:

Received 1 June 2018

Revised 5 December 2018

Accepted 7 January 2019

Available online 9 August 2019

JEL classification:

G12

G14

G15

F37

Keywords:

Time series momentum

Risk premium

Return predictability

Pooled regression

ABSTRACT

Time series momentum (TSM) refers to the predictability of the past 12-month return on the next one-month return and is the focus of several recent influential studies. This paper shows that asset-by-asset time series regressions reveal little evidence of TSM, both in- and out-of-sample. While the *t*-statistic in a pooled regression appears large, it is not statistically reliable as it is less than the critical values of parametric and nonparametric bootstraps. From an investment perspective, the TSM strategy is profitable, but its performance is virtually the same as that of a similar strategy that is based on historical sample mean and does not require predictability. Overall, the evidence on TSM is weak, particularly for the large cross section of assets.

© 2019 Elsevier B.V. All rights reserved.

1. Introduction

Whether past returns predict future returns is a central topic in finance. Fama (1965), French and Roll (1986), Lo and MacKinlay (1988), and Conrad and Kaul (1988), among others, show that past returns can positively predict future returns at a short horizon, but the magnitude is too small to be exploitable. Moskowitz et al. (MOP, 2012) conclude with a much greater degree of predictability that time series momentum (TSM) is everywhere: The past 12-month return positively predicts the next one- to 12-month return for a comprehensive set of approximately 55 assets. In addition, MOP show that a TSM trading strategy, which buys assets if their past 12-month returns are positive and sells them otherwise, earns significant average and risk-adjusted returns. Hurst et al. (2017) and Georgopoulou and Wang (2017) examine the TSM on a broader range of asset classes and longer sample periods, and they find similar results. Koijen et al. (2018) use TSM portfolio returns as a risk

☆ We are grateful to William Schwert (the editor) and an anonymous referee for very insightful and helpful comments that significantly improved the paper. We also thank Lauren H. Cohen, Eugene Fama, Amit Goyal, Wesley R. Gray, Campbell R. Harvey, Johan Hombert, Narasimhan Jegadeesh, Raymond Kan, Yan Liu, Roger Loh, David Rapach, Yiuman Tse, John Wald, Dacheng Xiu, and seminar and conference participants at Baruch College, Rutgers University, Washington University in St. Louis, 2017 Conference on Financial Predictability and Data Science, 2017 Workshop on Advanced Econometrics at the University of Kansas, and 2018 National Bureau of Economic Research-National Science Foundation (NBER-NSF) Time Series Conference for valuable comments. Dashan Huang acknowledges that this study was partially funded at the Singapore Management University through a research grant (C207MSS17B009) from the Ministry of Education Academic Research Fund Tier 1.

* Corresponding author.

E-mail addresses: dashanhuang@smu.edu.sg (D. Huang),

jiangyuanli.2015@pbs.smu.edu.sg (J. Li),

liyao.wang.2015@phdecons.smu.edu.sg (L. Wang), zhou@wustl.edu (G. Zhou).

factor to analyze carry trade. Kim et al. (2016) find that the profits of the TSM strategy are driven by volatility scaling and that its performance without volatility scaling is no better than that of a buy-and-hold strategy. Moreover, in a concurrent study, Goyal and Jegadeesh (2018) show that the traditional cross-sectional momentum strategy is more profitable than the TSM once the leverage ratio is properly adjusted. Despite an improved understanding, whether time series predictability is present at the 12-month frequency remains an open question.¹

In this paper, using the same data set as MOP (Moskowitz et al., 2012), we reexamine the statistical and economic evidence of TSM. From both time series and cross-sectional analyses, we find that the evidence on TSM is weak. Hence, concluding that TSM exists across the global asset classes appears questionable.

We conduct our study in three stages. In the first stage, we run a time series regression of monthly return for each asset on its past 12-month return. At the 10% significance level, 47 of the 55 assets have a t -statistic of less than 1.65, suggesting that the in-sample evidence of TSM is weak.² Since Welch and Goyal (2008), studies on return predictability have shifted the focus to out-of-sample performance. We compute the standard (Campbell and Thompson, 2008) out-of-sample R_{OS}^2 for each asset and find that only three assets deliver significant R_{OS}^2 at the 10% level. Univariate time series regressions thus indicate that the evidence on time series predictability is weak among the assets.

In the second stage, we follow MOP's approach and run a pooled regression by stacking all asset returns together. Consistent with MOP, we find a t -statistic of 4.34 in the regression of predicting the next one-month return using the past 12-month return. At the conventional critical level of 2, one could interpret this t -statistic of 4.34 as strong evidence against no predictability. We argue that the pooled regression is likely to over-reject the null hypothesis for three reasons. First, if assets have different mean returns, the slope estimate from the pooled regression without controlling for fixed effects tends to be biased upward (Hjalmarsson, 2010). Second, as a predictor, the past 12-month return is persistent and can generate substantial size distortions (Hodrick, 1992; Stambaugh, 1999; Campbell and Yogo, 2006; Ang and Bekaert, 2007; Boudoukh et al., 2008; Li and Yu, 2012). Third, because volatility varies dramatically across assets, volatility scaling in the pooled regression without controlling for fixed effects further exacerbates the upward bias.

To assess the degree of over-rejection, we use two bootstrap methods. The first is a parametric wild bootstrap that simulates samples based on the fitted pooled regression residuals, and the second is a nonparametric pairs bootstrap that resamples the predictor and the dependent variable simultaneously. Both methods accommodate conditional heteroskedasticity, but the latter allows for more general data-generating processes. We find that the 5%

critical values of the bootstraps are 12.53 and 4.83, respectively. They are larger than 4.34, the t -statistic from the pooled regression with real data. This finding is robust to all alternative cases, such as within each asset class, with different sample periods, and without volatility scaling. Hence, a high t -statistic found by MOP is not statistically significant in supporting the existence of TSM.

In the third stage, we examine why the TSM strategy is profitable even though the statistical evidence on time series predictability is weak. In their study, MOP (Moskowitz et al., 2012) construct the TSM strategy by buying assets with positive past 12-month return and selling assets with negative past 12-month return. At the same time, MOP assign a portfolio weight equal to 40% divided by the asset volatility, so that for each asset the ex ante annualized volatility is 40%. This volatility scaling can make the attribution of the performance of the TSM strategy complex. To separate the volatility effect, we follow (Kim et al., 2016) and Goyal and Jegadeesh (2018) and focus on the TSM strategy with the simple equal weighting without volatility scaling as the benchmark. Volatility scaling is not an issue if all strategies are based on volatility-scaled returns when comparing their performances.

We examine the performance of the TSM strategy in four ways. First, we investigate the performance of an alternative trading strategy that does not require predictability. Based on the observation that a high mean asset is more likely to have a positive past 12-month return and therefore is more likely to be bought by the TSM strategy, we propose a times series history (TSH) strategy that buys assets if their historical mean returns are positive and sells them otherwise, which is theoretically profitable even if asset returns are independent over time, but some have significantly higher means than others. We find that the TSM and TSH strategies perform virtually the same and their differences in average returns, as well as in risk-adjusted ones, are indifferent from zero. Also, we find that the performances of the TSM and TSH strategies mainly stem from their long legs, and their short legs have insignificant average and risk-adjusted returns. This result suggests that both the TSM and TSH tend to long assets with greater mean returns and short those with lower means. Because the TSH strategy is defined without requiring time series predictability, it seems questionable to attribute the performance of the TSM strategy to predictability.

Second, we report the results with alternative portfolio weighting schemes, such as volatility weighting as in MOP (Moskowitz et al., 2012), past 12-month return weighting, and equal weighting with a zero-investment constraint as in Goyal and Jegadeesh (2018). The results are similar, and the alpha differential between the TSM and TSH strategies is always indifferent from zero. In short, the profitable performance of the TSM strategy is similar to that of the TSH strategy that requires no predictability, suggesting that the performance of the TSM does not necessarily support predictability at the 12-month frequency across asset classes.

Third, based on the predictive slope of Lewellen (2015), we examine the overall predictability of TSM across assets. The slope measures how realized returns are explained by predicted returns. If the past 12-month return perfectly

¹ In this paper, the terms "TSM" and "time series predictability" are used interchangeably.

² The multiple test framework of Harvey et al. (2016) could suggest even weaker evidence on TSM.

predicts the next one-month return, the slope should have a value of one. We find that for the TSM forecasts the slope has a value close to zero, suggesting that the TSM forecasts have little predictive power.³ When regressing the TSM forecasts on the TSH forecasts, the slope is very close to one, irrespective of whether the TSM forecasts use volatility scaling or not. This result indicates little difference in predictability between the two forecasts, suggesting, again, no evidence of TSM across the assets.

Fourth, of interest is examining under what conditions the TSM is a better trading strategy than the TSH. Based on one thousand simulated samples by using pooled regression with varying assumed degrees of time series momentum (i.e., the slope varies from 0.1 to 0.4), we find that the TSM and TSH strategies perform similarly when the slope is 0.1 (with real data, the slope of the pooled regression is 0.08, controlling for fixed effects). When the slope is 0.2, the TSM outperforms the TSH, but the difference is statistically insignificant. This means that if there is genuine time series predictability, the advantage of the TSM strategy is not apparent as long as the slope is small. When the slope is 0.4, the TSM dominates the TSH in the sense that it does better in almost all the simulated samples. Because the two strategies generate similar performance using real data, our simulation indicates that the evidence of TSM is likely weak if it exists. Combined with other results, the TSM is unlikely to be statistically significant for all the assets. In short, a lack of empirical evidence exists to support the hypothesis that the TSM is everywhere.

As a final remark, our results do not claim in any way that there is no predictability in the asset classes, but that the predictability, if it exists, is not as simple as a constant 12-month return rule. The best example could be the stock market, with ample evidence that its risk premium can be predicted by a wide range of predictors such as macroeconomic variables and investor sentiment [see, e.g., Jiang et al. (2019) for the latest literature]. Cross-sectionally, since Jegadeesh and Titman (1993) discovery of momentum, there have been hundreds of potential anomalies [see, e.g., Hou et al. (2019) for the replication]. Recently, Gu et al. (2018) and Freyberger et al. (2019), among others, use machine learning tools to find even stronger predictability.⁴ However, none of them is related to the TSM.

The rest of this paper is organized as follows. Section 2 introduces data we use in this paper. Section 3 shows that asset-by-asset regressions suggest that the evidence of TSM, if any, is weak. Section 4 finds that the pooled regression overstates the presence of TSM and that bootstrap-corrected *t*-statistics cannot reject the null hypothesis of no predictability. Section 5 shows that the TSM strategy performs the same as an alternative trading strategy that does not require predictability. Section 6 concludes.

2. Data

We collect futures prices for 24 commodities, nine developed country equity indexes, 13 developed government bonds, and nine currency forwards from the same data sources as MOP (Moskowitz et al., 2012). These 55 instruments are the same as those in MOP's Table 1.⁵ The sample period is from January 1985 to December 2015. For each day, we calculate the daily excess return of each futures contract with the nearest- or next-nearest-to-delivery contract and compound the daily returns to a cumulative month return index. For brevity, returns in this paper always refer to excess returns, unless otherwise stated.

Table 1 reports the sample mean (arithmetic), volatility (standard deviation), and first-order autocorrelation of the returns on the 55 futures contracts. The mean and volatility are annualized and represented in percentage. Significant variations exist in mean return and volatility across different contracts. Within the commodity asset class, 24 contracts yield positive, zero, and negative mean returns, from −11.33% for natural gas to 12.31% for unleaded gasoline. The volatility ranges from 13.56% for cattle to 50.79% for natural gas. On average, the mean return and volatility are 2.59% and 28.50%, respectively. The nine equity index futures contracts are more homogenous, with mean return from 3.09% for TOPIX (Tokyo Stock Price Index) to 9.69% for DAX (German Stock Index) and volatility from 15.87% for FTSE 100 (Financial Times Stock Exchange 100 Index) to 23.21% for FTSE/MIB (Italian National Stock Exchange Index). On average, the return and volatility are 7.24% and 19.26%, respectively. Finally, bond futures and currency forwards earn lower mean returns with lower volatilities. Within each asset class, the average mean and volatility are 4.54% and 7.85% for bond futures and 1.22% and 10.90% for currency forwards, respectively. Table 1 also highlights a well-known fact that the past one-month return cannot predict the next one-month return, because the first-order correlation is generally close to zero.

3. Univariate time series regression

In this section, we run univariate time series regressions to explore the predictability of the past 12-month return for individual assets. These regressions clearly tell which asset return can be predicted by its past 12-month return and which cannot, thereby providing direct evidence on whether the finding in MOP (Moskowitz et al., 2012) is common across asset classes.

3.1. In-sample performance

Time series regression is standard for identifying return predictability, but pooled regression is seldom used in the literature. Ang and Bekaert (2007) and Hjalmarsson (2010) are exceptions, showing a heterogeneous pattern in

³ Lewellen (2015) shows in his footnote 3 that a slope of less than 0.5 under-performs a naive forecast. Han et al. (2019) discuss more properties of the slope.

⁴ Rapach et al. (2013) use LASSO, a major machine learning tool, to forecast international equity markets, which is the earliest study that we find in the finance literature applying LASSO to predict stock returns.

⁵ MOP use 12 cross-currency pairs in trading strategies but report only nine underlying currencies in their summary statistics table. To maximally replicate the results, we focus on the nine underlying currencies. As a consequence, we examine a total of 55 assets, not 58.

Table 1

Summary statistics of 55 assets across four asset classes.

This table reports the mean return, volatility (standard deviation), and first-order autocorrelation [$\rho(1)$], where the mean and volatility are annualized and represented in percentage. “Average” refers to the average value within asset class. The sample period is 1985:01–2015:12.

Asset	Start date	Mean	Volatility	$\rho(1)$
Panel A: Commodity futures				
Aluminum	November 1986	−2.09	19.92	0.10
Brent oil	February 1992	7.77	30.62	0.12
Cattle	January 1985	1.64	13.56	0.02
Cocoa	January 1985	−2.54	28.11	−0.13
Coffee	January 1985	−2.06	37.95	−0.02
Copper	January 1985	11.48	26.80	0.15
Corn	January 1985	−4.91	26.65	0.00
Cotton	January 1985	1.36	25.83	0.03
Crude	January 1985	6.37	34.98	0.19
Gas oil	April 1989	8.21	33.23	0.12
Gold	January 1985	1.54	15.65	−0.12
Heating oil	January 1985	6.41	32.60	0.08
Hogs	January 1985	−3.70	24.23	−0.04
Natural gas	April 1990	−11.33	50.79	0.08
Nickel	February 1993	7.06	34.37	0.05
Platinum	January 1992	6.36	20.76	0.06
Silver	January 1985	2.02	27.73	−0.08
Soybeans	January 1985	4.01	23.17	−0.05
Soymeal	January 1985	6.23	28.75	−0.11
Soy oil	October 1990	4.33	26.12	−0.07
Sugar	January 1985	6.24	34.75	0.11
Unleaded gasoline	January 1985	12.31	35.71	0.10
Wheat	January 1985	−4.33	26.68	−0.03
Zinc	February 1991	−0.33	25.08	0.00
Average		2.59	28.50	0.02
Panel B: Equity index futures				
SPI 200	January 1985	7.41	16.11	0.00
DAX	January 1985	9.69	21.70	0.07
IBEX 35	January 1985	9.27	22.66	0.10
CAC 40	January 1985	6.73	19.69	0.09
FTSE/MIB	January 1985	6.29	23.21	0.05
TOPIX	January 1985	3.09	19.70	0.09
AEX	January 1985	6.96	19.16	0.08
FTSE 100	January 1985	6.51	15.87	−0.01
S&P 500	January 1985	9.21	15.20	0.04
Average		7.24	19.26	0.06
Panel C: Government bond futures				
Three-year Australian	July 1988	3.34	4.58	−0.05
Ten-year Australian	January 1985	5.57	6.90	0.08
Two-year European	January 1989	1.47	3.41	0.13
Five-year European	January 1989	1.83	4.26	0.09
Ten-year European	January 1985	4.16	9.66	0.03
Thirty-year European	January 1987	7.56	10.39	0.09
Ten-year Canadian	January 1985	6.44	10.79	−0.03
Ten-year Japanese	December 1986	3.66	13.78	0.07
Ten-year UK	January 1985	4.14	8.28	0.08
Two-year US	January 1989	1.49	1.67	0.22
Five-year US	January 1989	2.84	4.30	0.15
Ten-year US	January 1985	3.64	7.60	0.06
Thirty-year US	January 1985	12.59	16.44	0.07
Average		4.54	7.85	0.08
Panel D: Currency forwards				
AUD/USD	January 1985	1.10	12.03	0.06
EUR/USD	January 1985	2.06	11.02	0.01
CAD/USD	January 1985	0.43	7.44	−0.06
JPY/USD	January 1985	1.72	11.12	0.05
NOK/USD	January 1985	0.51	11.01	0.04
NZD/USD	January 1985	2.13	12.31	−0.02
SEK/USD	January 1985	0.39	11.25	0.10
CHF/USD	January 1985	2.92	11.77	−0.01
GBP/USD	January 1985	−0.28	10.14	0.07
Average		1.22	10.90	0.03

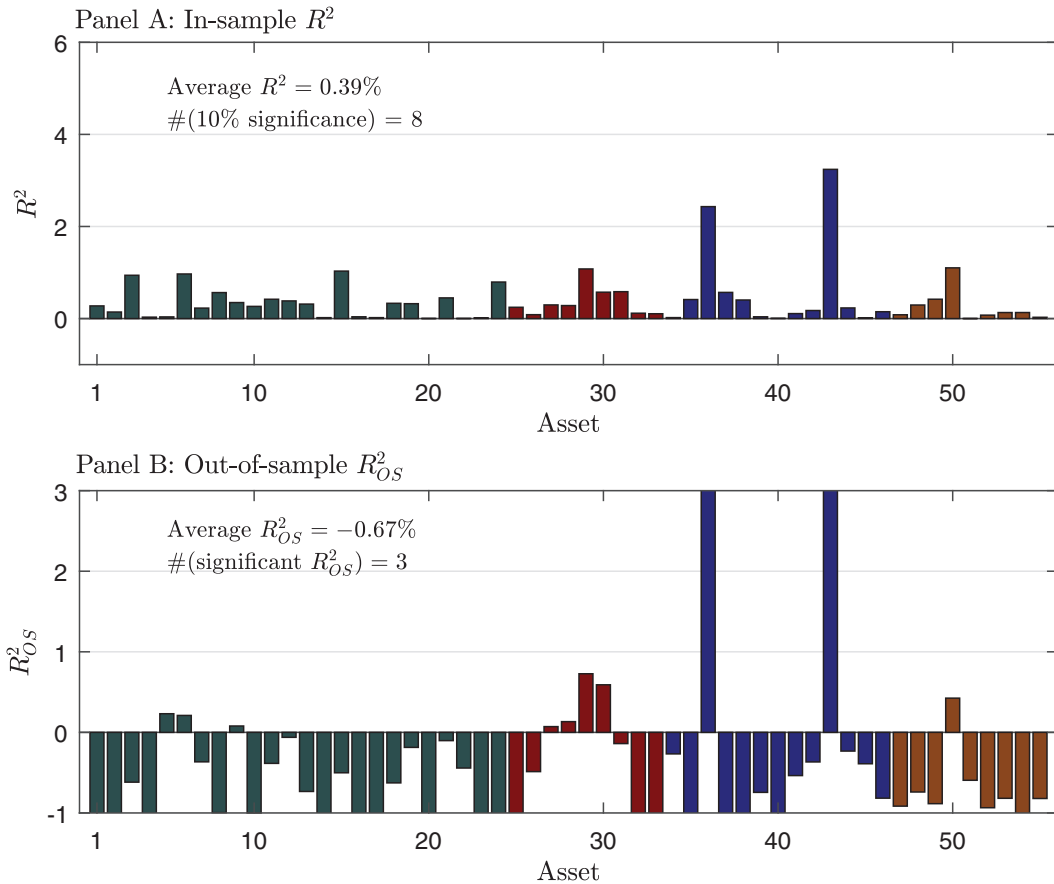


Fig. 1. Time series momentum (TSM) with asset-by-asset regression. This figure plots the in- and out-of-sample R^2 s of forecasting a futures contract return with time series regression as

$$r_{t+1}^i = \alpha_i + \beta_i r_{t-12,t}^i + \varepsilon_{t+1}^i,$$

where $r_{t-12,t}^i$ is asset i 's past 12-month return. The in- and out-of-sample periods are 1985:01–2015:12 and 2000:01–2015:12, respectively.

predictability. For example, [Ang and Bekaert \(2007, p.663\)](#) show that, in contrast to the US, the UK and Japan return predictability disappears when expected returns are constrained to be non-negative and conclude that “none of the [return predictability] patterns in other countries resembles the US pattern.”

For each asset, we run the predictive regression

$$r_{t+1}^i = \alpha + \beta r_{t-12,t}^i + \varepsilon_{t+1}^i, \quad (1)$$

where r_{t+1}^i is the return of asset i in month $t+1$ and $r_{t-12,t}^i$ is its past 12-month return (i.e., the return between months $t-12$ and t). The predictive power is based on either the regression slope β or the R^2 statistic. If the regression R^2 statistic is significantly larger than zero with a positive β , then asset i displays TSM, i.e., its past 12-month return predicts the next one-month return.

[Table 2](#) reports the regression slope, the Newey–West t -statistic, and R^2 . We have four observations. First, the presence of TSM is not prevalent. Of the 55 assets, only eight display significant regression slopes at the 10% level, representing 15% of assets (only three significant at the 5% level). Second, the significance is not concentrated but disperse among the four asset classes, including three com-

modities, two equity indexes, two government bonds, and one currency. Third, although not significant, 17 assets deliver negative slopes, amounting to 31% of the assets. Fourth, the R^2 s are small, with an average of 0.39%, and only five assets generate an R^2 larger than 1%. To have an intuitive understanding of the predictive performance, Panel A of [Fig. 1](#) plots the R^2 statistic for each asset. Only two assets have R^2 s that stand out above 2%.

3.2. Out-of-sample performance

Due to concerns of data mining and structural breaks, studies on return predictability have shifted the focus to out-of-sample performance since [Welch and Goyal \(2008\)](#). To investigate the out-of-sample performance of TSM, we use the [Campbell and Thompson \(2008\)](#) out-of-sample R^2_{OS} statistic as the assessment criterion, which is defined as

$$R^2_{OS} = 1 - \frac{\sum_{t=K}^{T-1} (r_{t+1}^i - \hat{r}_{t+1}^i)^2}{\sum_{t=K}^{T-1} (r_{t+1}^i - \bar{r}_{t+1}^i)^2}, \quad (2)$$

where K is the initial sample size for parameters training, $\hat{r}_{t+1}^i$ is the expected return estimated with information up to month t and calculated as $\hat{r}_{t+1}^i = \hat{\alpha}_t + \hat{\beta}_t r_{t-12,t}^i$, $\hat{\alpha}_t$ and

Table 2

In- and out-of-sample performance of time series momentum (TSM) with time series regression.

This table reports the slope, t -statistic, in-sample R^2 , and out-of-sample R_{OS}^2 of $r_{t+1}^i = \alpha_i + \beta_i r_{t-12,t}^i + \varepsilon_{t+1}^i$. “Average” refers to the average value within each asset class. # (10% significance) refers to the number of significant in-sample regression slopes or significant R_{OS}^2 s at the 10% level or stronger. ***, **, and * denote statistical significance at the 1%, 5%, and 10% level, respectively. The in- and out-of-sample periods are 1985:01–2015:12 and 2000:01–2015:12, respectively.

Asset	β_i	t -stat	R^2	R_{OS}^2
Panel A: Commodity futures				
Aluminum	0.30	0.88	0.28	−1.42
Brent oil	0.34	0.69	0.14	−1.29
Cattle	0.38**	2.23	0.94	−0.62
Cocoa	−0.14	−0.28	0.03	−1.51
Coffee	0.21	0.40	0.04	0.23
Copper	0.77*	1.69	0.97	0.21
Corn	−0.37	−0.94	0.23	−0.37
Cotton	0.57	1.23	0.56	−1.00
Crude	0.60	1.34	0.35	0.08
Gas oil	0.49	1.11	0.26	−1.00
Gold	0.29	1.43	0.42	−0.38
Heating oil	0.59	1.39	0.38	−0.06
Hogs	0.39	1.25	0.31	−0.73
Natural gas	−0.21	−0.30	0.02	−4.51
Nickel	1.01*	1.83	1.03	−0.50
Platinum	−0.12	−0.31	0.04	−1.83
Silver	−0.11	−0.25	0.02	−2.31
Soybeans	−0.39	−1.05	0.33	−0.63
Soymeal	−0.47	−1.06	0.32	−0.19
Soy oil	0.04	0.09	0.00	−1.93
Sugar	−0.65	−1.33	0.45	−0.10
Unleaded gasoline	0.05	0.12	0.00	−0.44
Wheat	−0.10	−0.26	0.02	−1.21
Zinc	0.65	1.24	0.79	−2.29
Average			0.33	−0.99
Panel B: Equity index futures				
SPI 200	−0.23	−0.49	0.25	−3.99
DAX	0.18	0.54	0.08	−0.49
IBEX 35	0.36	1.20	0.30	0.07
CAC 40	0.30	0.98	0.28	0.13
FTSE/MIB	0.70*	1.92	1.08	0.73
TOPIX	0.44*	1.69	0.57	0.59
AEX	0.43	1.22	0.58	−0.14
FTSE 100	−0.16	−0.48	0.12	−5.24
S&P 500	0.14	0.45	0.10	−1.78
Average			0.37	−1.12
Panel C: Government bond futures				
Three-year Australian	0.01	0.29	0.02	−0.27
Ten-year Australian	0.13	1.42	0.41	−1.74
Two-year European	0.15*	1.84	2.43	8.02**
Five-year European	0.09	1.11	0.57	−1.03
Ten-year European	0.17	1.25	0.40	−1.03
Thirty-year European	−0.06	−0.31	0.04	−0.74
Ten-year Canadian	−0.02	−0.17	0.01	−1.32
Ten-year Japanese	0.11	0.47	0.11	−0.54
Ten-year UK	0.10	1.06	0.18	−0.37
Two-year US	0.08***	3.57	3.24	4.26***
Five-year US	0.06	1.16	0.23	−0.23
Ten-year US	−0.03	−0.35	0.02	−0.39
Thirty-year US	0.17	0.60	0.15	−0.82
Average			0.60	0.29
Panel D: Currency forwards				
AUD/USD	−0.10	−0.47	0.08	−0.92
EUR/USD	0.17	1.08	0.29	−0.74
CAD/USD	0.14	1.18	0.42	−0.88
JPY/USD	0.33***	2.60	1.10	0.43*
NOK/USD	−0.01	−0.06	0.00	−0.59
NZD/USD	0.09	0.40	0.07	−0.94
SEK/USD	0.12	0.70	0.13	−0.82
CHF/USD	0.12	0.75	0.13	−1.42
GBP/USD	−0.05	−0.29	0.03	−0.82
Average			0.25	−0.75
Average across asset classes			0.39	−0.67
#(10% significance)	8			3

$\hat{\beta}_t$ are the coefficients of the time series regression Eq. (1), and $\bar{r}_{t+1}^i$ is the sample mean of asset i with data up to month t . The choice of K is ad hoc in the literature, which depends on the nature of the possible model instability and the timing of the possible breaks. Hansen and Timmermann (2012) theoretically show that a large K is preferable if the data-generating process is stationary, but it comes at the cost of low power as there are fewer observations for out-of-sample evaluation. A small K can give the out-of-sample test more desirable size properties, but it perhaps does not provide precise estimation. For these reasons, we select the first 15 years of data for in-sample training and the remaining 16 years of data for out-of-sample evaluation. That is, the full sample period is from January 1985 to December 2015 and the out-of-sample period is from January 2000 to December 2015.

Welch and Goyal (2008) show that the sample mean is a very stringent out-of-sample benchmark. If $R_{OS}^2 > 0$, the forecast $\hat{r}_{t+1}^i$ outperforms the sample mean in terms of mean squared forecast error (MSFE). Empirically, they show that the in-sample forecasting abilities of a variety of return predictors generally do not hold in out-of-sample tests. To ascertain whether a forecast delivers a statistically significant improvement in MSFE relative to the sample mean, we use the Clark and West (2007) statistic to test the null hypothesis that the MSFE of the sample mean forecast is less than or equal to the MSFE of the forecasted expected return, corresponding to $H_0: R_{OS}^2 \leq 0$ against $H_A: R_{OS}^2 > 0$.

Although no strict relation exists between the in-sample and out-of-sample performance (Inoue and Kilian, 2005), the last column of Table 2 and Panel B of Fig. 1 show that the R_{OS}^2 is smaller than the in-sample R^2 on average. Of the 55 assets, 45 have negative R_{OS}^2 , indicating no out-of-sample predictability. Of the remaining ten assets with positive R_{OS}^2 , only three are significant at the 10% level, which are the two-year European bond, the two-year US bond, and the JPY/USD (Japanese yen/US dollar) forward. As a result, the average R_{OS}^2 across the 55 assets is -0.67% , suggesting that there is no TSM out of sample.

To further explore the robustness, Fig. 2 plots the R^2 s and R_{OS}^2 s of TSM by regressing the next one-month return on the past one-, three-, and six-month return, respectively. The results are similar to the case with the past 12-month return in Fig. 1. In addition, we consider volatility scaling when running the asset-by-asset time series regressions. The results are still quantitatively the same: Only the same three assets have significant R_{OS}^2 s. In sum, based on the typical univariate time series regression, the evidence of TSM across all the assets is very weak.

4. Pooled regression

In this section, we first replicate the results in MOP (Moskowitz et al., 2012) and then show that the pooled regression tends to overstate the presence of TSM.

4.1. The t -statistic

By stacking all futures contracts' returns and dates, MOP (Moskowitz et al., 2012) run a pooled predictive re-

gression of monthly returns scaled by volatility on the scaled returns lagged h months,

$$r_{t+1}^i/\sigma_t^i = \alpha + \beta r_{t-h+1}^i/\sigma_{t-h}^i + \varepsilon_{t+1}^i, \quad (3)$$

where r_{t+1}^i is asset i 's return in month $t+1$ and σ_t^i is the ex ante annualized volatility estimated by its exponentially weighted lagged squared daily returns:

$$(\sigma_t^i)^2 = 261 \sum_{j=0}^{\infty} (1-\delta)\delta^j (r_{t-j}^i - \bar{r}_t^i)^2, \quad (4)$$

where $\bar{r}_t^i$ is the exponentially weighted average return and δ is chosen so that $\sum_{j=0}^{\infty} (1-\delta)\delta^j = 60$ days.

MOP (Moskowitz et al., 2012) also use an alternative specification with the sign of lagged returns as the regressor to examine the robustness of TSM,

$$r_{t+1}^i/\sigma_t^i = \alpha + \beta \text{sign}(r_{t-h+1}^i) + \varepsilon_{t+1}^i, \quad (5)$$

where sign is the sign function that equals $+1$ when $r_{t-h+1}^i \geq 0$ and -1 when $r_{t-h+1}^i < 0$.

As in MOP (Moskowitz et al., 2012), we calculate the t -statistics by clustering the standard errors by time (month) and plot the t -statistics of the pooled regression slopes with lagged returns from one month to 60 months in Fig. 3.⁶ Qualitatively, we confirm MOP (Moskowitz et al., 2012) that the past 12-month return of each asset is a positive predictor of its future returns for one month to 12 months in the pooled regression. After 12 months, the forecasting sign changes and the forecasting power decays. In Fig. 3, Panel A shows the results with Eq. (3) over all asset classes, and Panel B replaces the lagged return with its sign as Eq. (5). Both have a sizable t -statistic of about 4 at the 12-month horizon.

Panels C to F of Fig. 3 plot the t -statistics of Eq. (5) within each asset class and exhibit a similar pattern. The t -statistics appear to show a strong return continuation for the first 12 months and weak reversal for the following 48 months. Overall, Fig. 3 appears to provide strong evidence on TSM. However, the t -statistics at the conventional level tend to overstate the predictability of the past 12-month return.

4.2. Estimation bias

In Eq. (3), MOP (Moskowitz et al., 2012) make an implicit assumption that the mean returns of all assets are the same by imposing a common intercept. From Table 1, the sample means of individual assets vary dramatically across asset classes. In the literature, Jorion and Goetzmann (1999) show strong evidence that the equity premium varies across countries. Ang and Bekaert (2007) investigate return predictability with pooled regression but explicitly consider the variation in average returns. Menzly et al. (2004) analyze cross-sectional differences in time series return predictability.

To highlight fixed effects, a possible specification is

$$r_{t+1}^i/\sigma_t^i = \alpha + \beta r_{t-h+1}^i/\sigma_{t-h}^i + \mu_i/\sigma_i + \varepsilon_{t+1}^i, \quad (6)$$

⁶ The t -statistics that double-cluster the standard errors by time and asset are quantitatively similar.

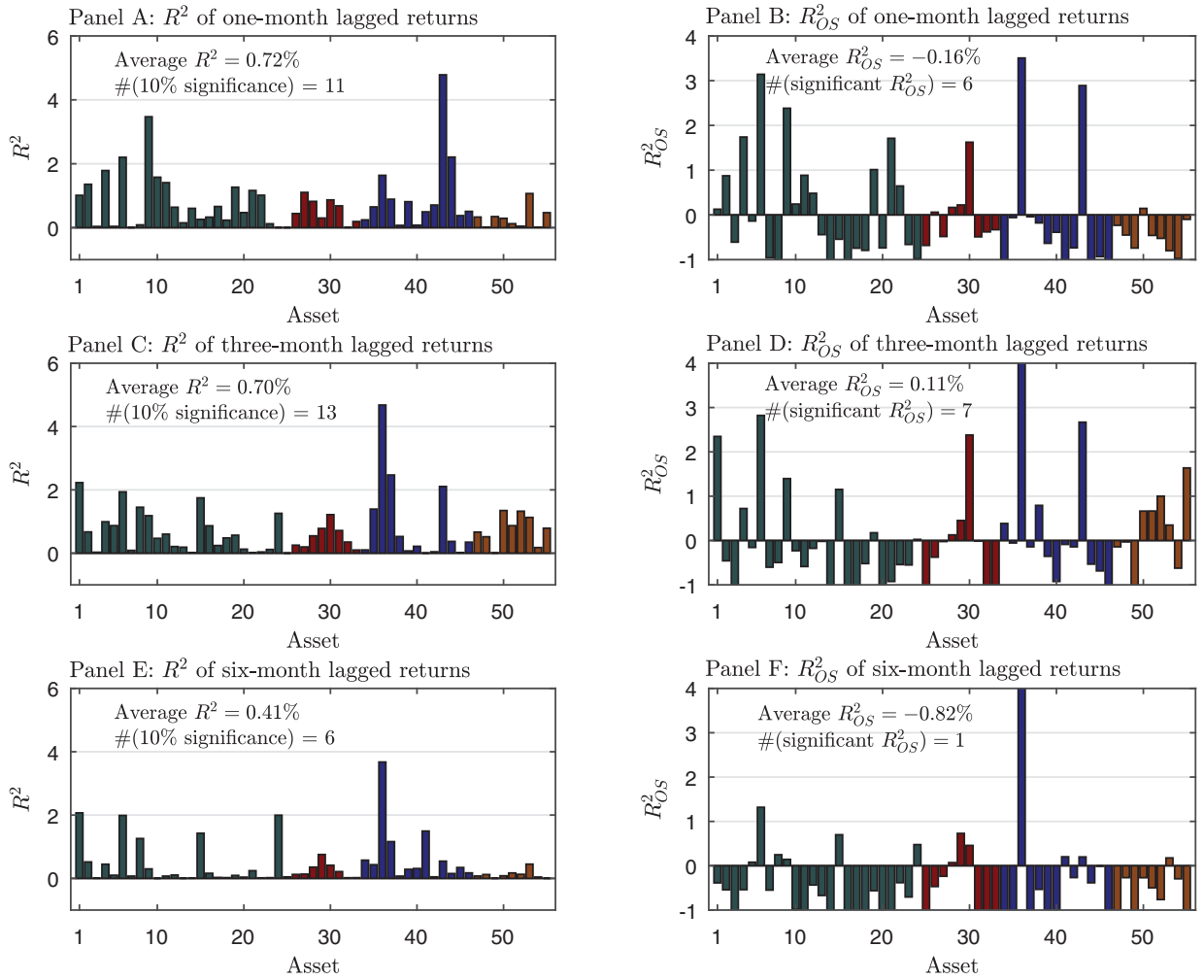


Fig. 2. Time series momentum (TSM) with asset-by-asset regression based on different lags. This figure plots the in- and out-of-sample R^2 s of forecasting a futures contract return with time series regression as

$$r_{t+1}^i = \alpha_i + \beta_i r_{t-h,t}^i + \varepsilon_{t+1}^i,$$

where $r_{t-h,t}^i$ is asset i 's past h -month return ($h = 1, 3$, and 6). The in- and out-of-sample periods are 1985:01–2015:12 and 2000:01–2015:12, respectively.

where μ_i and σ_i are the unconditional mean and volatility of asset i . Hence, the estimate of β from Eq. (3) should be

$$\hat{\beta} = \beta + \frac{\text{Cov}(r_{t-h+1}^i / \sigma_{t-h}^i, \mu_i / \sigma_i)}{\text{Var}(r_{t-h+1}^i / \sigma_{t-h}^i)}. \quad (7)$$

If all assets have the same Sharpe ratio (or mean, if volatilities are the same), the second term is zero. Otherwise, it would be significantly positive when the number of assets is large, as the correlation between realized returns and their means is mechanically positive. As a result, the slope estimate of Eq. (3) is biased upward.

The question then is whether the 55 assets have the same mean or Sharpe ratio. We perform four tests for this hypothesis. The first is analysis of variance (ANOVA), which was proposed by Fisher (1918) with two assumptions: normality and homoskedasticity. The second is B.L. Welch's ANOVA (Welch, 1951), which allows the variance

to be heteroskedastic. The third is the Kruskal–Wallis test, which relaxes both the normality and homoskedasticity assumptions (Kruskal and Wallis, 1952), and the fourth is a bootstrap test. When applying these four tests to real data, Table 3 shows that the null hypothesis that all 55 assets have the same mean is strongly rejected. In addition, we reject the null that they have the same Sharpe ratio. Therefore, the evidence of TSM shown in MOP (Moskowitz et al., 2012) is at least partially driven by the fixed effects.

The fixed effects are not easily corrected in the predictive regression framework. Statistically, Hjalmarsen (2010) shows that when different assets have different average returns, the pooled regression, after controlling for fixed effects, suffers from a looking-forward bias because the time series demeaning of the data requires information after month t , which induces a correlation between the lagged value of the demeaned regressor and the error term in the predictive regression.

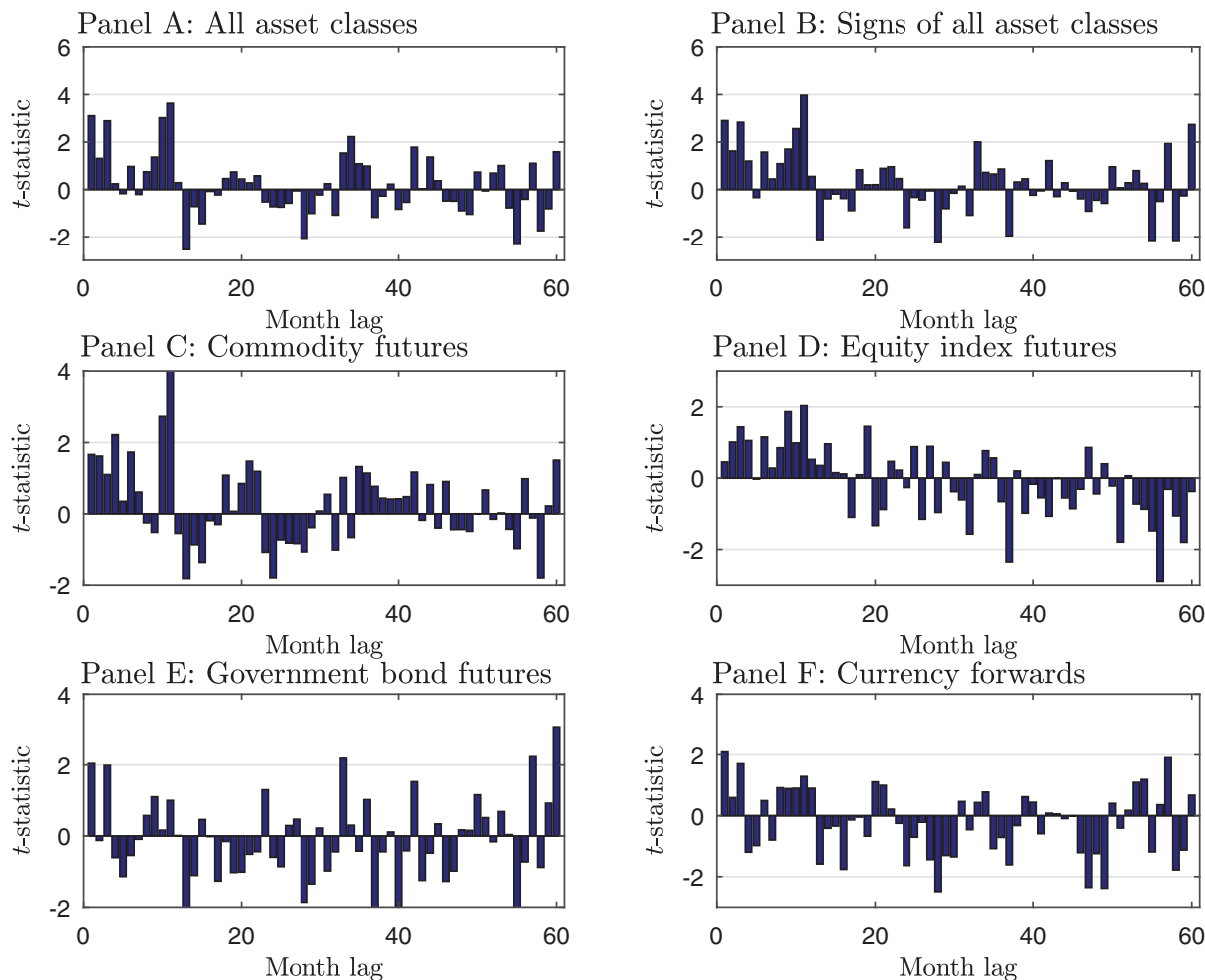


Fig. 3. Time series momentum (TSM) with pooled regression: in-sample performance. This figure plots the t -statistics of the pooled regression slopes that regress month t returns on month $t-h$ returns as

$$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h r_{t-h+1}^i/\sigma_{t-h}^i + \varepsilon_{t+1}^i,$$

for Panel A and

$$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h \text{sign}(r_{t-h+1}^i) + \varepsilon_{t+1}^i$$

for Panels B to F, where r_{t-h+1}^i is asset i 's return in month $t-h+1$ for $h = 1, 2, \dots, 60$. The t -statistics are clustered by time (month). The sample period is 1985:01–2015:12.

	ANOVA	Welch's ANOVA	Kruskal-Wallis	Bootstrap
Mean	0.08	$< 10^{-3}$	$< 10^{-10}$	0
Sharpe ratio	$< 10^{-5}$	$< 10^{-5}$	$< 10^{-15}$	0

In addition to the fixed effects, two more reasons can lead to overstating the evidence on time series predictability. First, as a predictor, the past 12-month cumulative return is persistent and can generate substantial size distortions (Hodrick, 1992; Stambaugh, 1999; Valkanov, 2003; Campbell and Yogo, 2006; Ang and Bekaert, 2007; Boudoukh et al., 2008; Li and Yu, 2012). For ex-

ample, Ang and Bekaert (2007) show substantial size distortions with the Newey–West t -statistic when predicting stock returns with persistent variables. Second, because volatility varies dramatically across assets, volatility scaling in the pooled regression without controlling for fixed effects can further exacerbate the upward bias. For example, in Eq. (6), even when all assets have the same mean,

volatility scaling generates the fixed effects as σ_i varies dramatically across assets. MOP (Moskowitz et al., 2012) also explore the pooled regression in Eq. (5) by using the sign of the past 12-month return as the predictor, which can distort and change the true statistical significance as the sign of the past 12-month return is highly skewed.

4.3. Bootstrap tests

Due to the concerns discussed above, the t -statistic from the pooled regression is questionable. To correctly evaluate the statistic significance, we use bootstrap to simulate the distribution of the t -statistic and define its 97.5% quantile as the simulated t -statistic for significance at the 5% level. If the t -statistic from the real data is larger than 1.96 but smaller than the simulated t -statistic, we can conclude that the pooled regression tends to overreject the null hypothesis and no significant evidence supports TSM.

We use two standard bootstrap approaches. The first is a more restrictive parametric wild bootstrap that samples data based on the pooled regression residuals, and the second is a more general nonparametric pairs bootstrap that resamples the predictor and the dependent variable simultaneously. Both approaches accommodate conditional heteroskedasticity, but the second allows for a wider range

of data-generating processes. Pairs bootstrap is considered the most general and applicable method of bootstrapping.

Wild bootstrap. Suppose the true data-generating process is as Eq. (3). Let $\hat{\alpha}$ and $\hat{\beta}$ be the estimates from the full sample of real data. Then, the residuals are

$$\varepsilon_{t+1}^i = r_{t+1}^i / \sigma_t^i - \hat{\alpha} - \hat{\beta} r_{t-h+1}^i / \sigma_{t-h}^i. \quad (8)$$

We simulate a pseudo sample path with T observations as

$$r_{t+1}^{i*} / \sigma_t^{i*} = \hat{\alpha} + \hat{\beta} r_{t-h+1}^i / \sigma_{t-h}^i + \varepsilon_{t+1}^i v_{t+1}^i, \quad (9)$$

where $*$ indicates that the value is a bootstrapped observation and v_t^i is a random draw from a two-point Rademacher distribution with mean 0 and variance 1:

$$v_t^i = \begin{cases} 1 & \text{with probability } 1/2, \\ -1 & \text{with probability } 1/2. \end{cases} \quad (10)$$

This distribution has an appealing property that the error-in-rejection probability is minimal when the sample size is small, and it is robust to other distributions such as the Mammen distribution and standard normal distribution (Davidson and Flachaire, 2008). After constructing a pseudo sample path, we run pooled regression Eq. (3). We repeat this procedure one thousand times to calculate the simulated t -statistic of $\hat{\beta}$.

Table 4

t -statistic of pooled regression without controlling for fixed effects.

This table reports the t -statistic of pooled regression with real data and the simulated t -statistics with wild and pairs bootstrap. For each asset, we bootstrap a path with T observations and run pooled regression without controlling for fixed effects to calculate the t -statistic. We repeat this procedure one thousand times and obtain the distribution of the t -statistic for testing the null hypothesis that there is no time series momentum. The bootstrapped t -statistic is defined as the 97.5% percentile of the simulated t -statistics. The sample period is 1985:01–2015:12.

h	t -statistic	Bootstrapped t -statistic		t -statistic	Bootstrapped t -statistic	
		Wild	Pairs		Wild	Pairs
Panel A: Forecast with return lagged h months						
	$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h r_{t-h+1}^i/\sigma_{t-h}^i + \varepsilon_{t+1}^i$			$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h \text{sign}(r_{t-h+1}^i) + \varepsilon_{t+1}^i$		
1	3.11	9.26	3.63	2.90	8.18	3.41
2	1.31	4.98	1.98	1.62	4.44	2.31
3	2.89	8.61	3.45	2.83	6.84	3.45
4	0.24	2.46	1.06	1.20	2.12	1.99
5	-0.17	1.88	0.60	-0.34	1.83	0.54
6	0.97	4.18	1.71	1.58	3.62	2.28
7	-0.21	1.52	0.65	0.44	1.55	1.29
8	0.75	3.81	1.49	1.09	3.20	1.84
9	1.36	4.76	2.10	1.70	4.20	2.47
10	3.02	8.12	3.60	2.56	6.46	3.30
11	3.63	10.34	4.13	3.97	7.61	4.39
12	0.29	2.52	1.00	0.55	2.35	1.26
Panel B: Forecast with past h -month returns						
	$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h r_{t-h,t}^i/\sigma_{t-1}^i + \varepsilon_{t+1}^i$			$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h \text{sign}(r_{t-h,t}^i) + \varepsilon_{t+1}^i$		
1	3.11	9.26	3.63	2.90	8.18	3.41
2	2.92	9.46	3.46	3.07	8.32	3.61
3	3.74	11.45	4.22	4.15	10.20	4.61
4	3.49	10.71	3.97	4.57	9.49	4.96
5	3.11	9.58	3.63	4.24	8.85	4.72
6	3.29	9.65	3.80	3.88	8.88	4.39
7	3.03	9.30	3.62	3.93	8.31	4.40
8	3.05	9.44	3.62	3.74	8.33	4.24
9	3.38	9.85	3.95	4.44	8.99	4.78
10	3.94	11.38	4.46	5.27	10.22	5.63
11	4.64	13.12	5.08	5.71	11.73	6.05
12	4.34	12.53	4.83	5.14	11.05	5.53

Table 5

t-statistic of pooled regression within each asset class without controlling for fixed effects.

This table reports the *t*-statistic of pooled regression with real data and the simulated *t*-statistics with wild and pairs bootstrap. For each asset, we bootstrap a path of *T* observations and run pooled regression without controlling for fixed effects to calculate the *t*-statistic. We repeat this procedure one thousand times and obtain the distribution of the *t*-statistic for testing the null hypothesis that there is no time series momentum. The bootstrapped *t*-statistic is defined as the 97.5% percentile of the simulated *t*-statistics. The sample period is 1985:01–2015:12.

<i>h</i>	$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h r_{t-h,t}^i/\sigma_{t-1}^i + \varepsilon_{t+1}^i$			$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h \text{sign}(r_{t-h,t}^i) + \varepsilon_{t+1}^i$		
	Bootstrapped <i>t</i> -statistic			Bootstrapped <i>t</i> -statistic		
	<i>t</i> -statistic	Wild	Pairs	<i>t</i> -statistic	Wild	Pairs
Panel A: Commodity futures						
1	1.74	5.09	2.49	1.66	4.86	2.40
3	2.32	6.28	2.97	2.95	5.82	3.52
6	2.98	7.08	3.65	2.89	6.59	3.54
12	3.46	8.20	4.13	4.20	7.43	4.68
Panel B: Equity index futures						
1	1.77	5.56	2.41	0.46	4.85	1.27
3	1.97	6.07	2.68	1.03	5.33	1.79
6	1.92	5.88	2.57	2.29	5.37	2.89
12	2.20	6.49	2.92	3.00	6.05	3.60
Panel C: Government bond futures						
1	2.35	6.58	3.04	2.04	5.78	2.75
3	2.37	6.68	2.99	2.87	5.75	3.41
6	0.60	3.22	1.53	0.73	2.93	1.61
12	1.68	5.39	2.44	1.69	4.77	2.42
Panel D: Currency forwards						
1	1.67	4.97	2.46	2.09	4.32	2.88
3	2.46	6.95	3.09	2.46	5.97	3.12
6	1.40	4.73	2.19	2.23	4.27	2.94
12	1.73	5.49	2.61	1.96	5.08	2.74

Pairs bootstrap. Pairs bootstrap resamples *T* pairs of $(r_{t+1}^i/\sigma_t^i, r_{t-h+1}^i/\sigma_{t-h}^i)$ with replacement from the real data and uses these pairs to run the pooled regression

$$r_{t+1}^{i*}/\sigma_t^{i*} = \alpha_h + \beta_h r_{t-h+1}^{i*}/\sigma_{t-h}^{i*} + \varepsilon_{t+1}^i. \quad (11)$$

Hence, the simulated *t*-statistic of $\hat{\beta}$ can be calculated after repeating the procedure one thousand times. This bootstrap allows not only more general data-generating processes, but also potential model misspecification [e.g., $E(r)$ is not a linear function of r_{t-h+1}].

Table 4 reports the *t*-statistics with real data and the bootstrapped *t*-statistics. Consistent with Ang and Bekaert (2007), the *t*-statistics that cluster by time tend to over-reject the null hypothesis. For example, when forecasting the next one-month return with the past one-month return, the *t*-statistic from the real data is 3.11, suggesting strong evidence of TSM. However, this is not the case because the bootstrapped *t*-statistics are 9.26 and 3.63, respectively. Similarly, when forecasting the next one-month return with the past 12-month return, the *t*-statistic from the real data is 4.34, and the simulated *t*-statistics are 12.53 and 4.83, respectively, suggesting that the evidence is weak in support of TSM. Moreover, forecasting with the sign of lagged returns does not support TSM either.

Table 5 presents the simulated *t*-statistics within each asset class. For brevity, we report the results of predicting the next one-month return with the past one-, three-, six-, and 12-month return, respectively. Consistent with

Tables 2 and 4, the results reveal that TSM is unlikely to be present in any of the four asset classes.

Does volatility scaling play a role in estimating the regression slopes and detecting the existence of TSM because MOP (Moskowitz et al., 2012) run Eqs. (3) and (5) with volatility scaling while volatility varies across assets and is predictable by its lagged values (Paye, 2012)? Table 6 reports the *t*-statistics with real data and bootstrapped *t*-statistics from Eqs. (3) and (5) without volatility scaling. The results display two empirical facts. First, the *t*-statistics without volatility scaling are much smaller than those with volatility scaling. For example, when predicting the next one-month return with the past one- and 12-month return without volatility scaling, the *t*-statistics of the regression slope are 1.80 and 1.68, respectively, which are much smaller than the values with volatility scaling in Table 4 (3.11 and 4.34), lending little support to TSM. Second, when predicting the next one-month return with the signs of the past one- and 12-month returns, the *t*-statistics are 2.20 and 3.72, respectively, and they are smaller than that with volatility scaling. Therefore, volatility scaling plays a role. In fact, it seems at least partially responsible for the performance of the TSM trading strategy.

This paper extends the sample ending period of MOP (Moskowitz et al., 2012) from 2009 to 2015 and raises a possibility that TSM exists before 2009 and disappears thereafter. Table 7 reports the results for the 1985 to 2009 sample period and rules out the possibility. The *t*-statistics from real data are still smaller than the simulated

Table 6

t-statistic of pooled regression without volatility scaling and without controlling for fixed effects.

This table reports the *t*-statistic of pooled regression with real data and the simulated *t*-statistics with wild and pairs bootstrap. For each asset, we bootstrap a path with *T* observations and run pooled regression without volatility scaling and without controlling for fixed effects to calculate the *t*-statistic. We repeat this procedure one thousand times and obtain the distribution of the *t*-statistic for testing the null hypothesis that there is no time series momentum. The bootstrapped *t*-statistic is defined as the 97.5% percentile of the simulated *t*-statistics. The sample period is 1985:01–2015:12.

h	t -statistic	Bootstrapped t -statistic		t -statistic	Bootstrapped t -statistic	
		Wild	Pairs		Wild	Pairs
Panel A: Forecast with return lagged h months						
	$r_{t+1}^i = \alpha_h + \beta_h r_{t-h+1}^i + \varepsilon_{t+1}^i$				$r_{t+1}^i = \alpha_h + \beta_h \text{sign}(r_{t-h+1}^i) + \varepsilon_{t+1}^i$	
1	1.80	5.49	2.51	2.20	6.13	2.85
2	0.52	2.58	1.47	1.65	2.67	2.45
3	1.43	4.57	2.19	1.84	4.58	2.58
4	0.67	3.21	1.58	1.47	3.21	2.26
5	-1.33	-0.10	-0.14	-0.89	-0.08	0.28
6	1.03	3.37	1.92	1.77	3.45	2.48
7	-1.21	-0.47	-0.18	-0.18	-0.50	0.77
8	-0.64	0.60	0.42	0.17	0.54	1.12
9	-0.97	0.23	0.28	0.22	0.19	1.30
10	2.52	6.11	3.21	2.72	5.87	3.42
11	5.04	9.88	5.30	5.17	9.89	5.51
12	-1.04	-0.17	0.08	-0.11	-0.06	0.85
Panel B: Forecast with past h -month returns						
	$r_{t+1}^i = \alpha_h + \beta_h r_{t-h,t}^i + \varepsilon_{t+1}^i$				$r_{t+1}^i = \alpha_h + \beta_h \text{sign}(r_{t-h,t}^i) + \varepsilon_{t+1}^i$	
1	1.80	5.49	2.51	2.20	6.13	2.85
2	1.39	4.56	2.21	2.57	5.10	3.18
3	1.71	5.26	2.45	3.06	5.81	3.62
4	1.82	5.30	2.59	3.75	5.94	4.25
5	1.27	4.27	2.09	3.23	4.75	3.77
6	1.55	4.85	2.39	2.71	5.38	3.32
7	1.04	3.89	1.90	2.54	4.26	3.19
8	0.78	3.49	1.58	2.24	3.64	2.94
9	0.62	3.12	1.50	2.57	3.31	3.29
10	1.08	3.75	1.96	3.62	4.23	4.14
11	2.09	5.61	2.84	4.17	6.41	4.72
12	1.68	4.96	2.50	3.72	5.64	4.16

t-statistics, regardless of which bootstrap approach is employed. For example, when forecasting the next one-month return with the past 12-month return, the *t*-statistic is 4.48 with real data, but it is 12.76 and 4.96 with the two bootstrap approaches, respectively. The TSM is also insignificant for each asset class. The results are reported in the Online Appendix.

4.4. Controlling for fixed effects

Earlier evidence shows that the assets do not have the same mean, implying that fixed effects should be controlled in the pooled regression. In so doing, one can run the pooled regression by removing the asset means. The bootstrap procedures can also make a similar modification. The question is whether controlling for fixed effects can alter substantially the evidence on TSM.

Following Gonçalves and Kaffo (2015), we now compute the *t*-statistic from the pooled regression

$$r_{t+1}^i / \sigma_t^i - \overline{r^i / \sigma^i} = \beta (r_{t-h+1}^i / \sigma_{t-h}^i - \overline{r_{t-h+1}^i / \sigma_{t-h}^i}) + \varepsilon_{t+1}^i, \quad (12)$$

where $\overline{r^i / \sigma^i}$ and $\overline{r_{t-h+1}^i / \sigma_{t-h}^i}$ denote the time series averages of r_{t+1}^i / σ_t^i and $r_{t-h+1}^i / \sigma_{t-h}^i$, respectively. Suppose

the estimate of β is $\hat{\beta}_{FE}$, then the residual $\hat{\varepsilon}_{t+1}^i$ can be calculated as

$$\hat{\varepsilon}_{t+1}^i = r_{t+1}^i / \sigma_t^i - \overline{r^i / \sigma^i} - \hat{\beta}_{FE} (r_{t-h+1}^i / \sigma_{t-h}^i - \overline{r_{t-h+1}^i / \sigma_{t-h}^i}). \quad (13)$$

Then, we simulate a pseudo sample path with *T* observations as

$$(r_{t+1}^i / \sigma_t^i - \overline{r^i / \sigma^i})^* = \hat{\beta}_{FE} (r_{t-h+1}^i / \sigma_{t-h}^i - \overline{r_{t-h+1}^i / \sigma_{t-h}^i}) + \hat{\varepsilon}_{t+1}^i v_{t+1}^i, \quad (14)$$

where v_{t+1}^i follows the Rademacher distribution. We can then estimate the model with the simulated samples and repeat the procedure one thousand times, to obtain the critical value of the wild bootstrap *t*-statistic.

Regarding the pairs bootstrap, we resample *T* pairs of the predictor and the dependent variable in Eq. (12) with replacement from real data after de-meaning and then use these pairs to rerun the pooled regression. The pairs bootstrap *t*-statistic is naturally obtained after repeating the procedure one thousand times. Both the wild and pairs bootstraps do not suffer from the incidental parameter bias emphasized in Gonçalves and Kaffo (2015), because the sample size is relatively large here.

Table 7

t-statistic of pooled regression without controlling for fixed effects over 1985:01–2009:12.

This table reports the *t*-statistic of pooled regression with real data and the bootstrapped *t*-statistics with wild and pairs bootstrap. For each asset, we bootstrap a path with *T* observations and run pooled regression without controlling for fixed effects to calculate the *t*-statistic. We repeat this procedure one thousand times and obtain the distribution of the *t*-statistic for testing the null hypothesis that there is no time series momentum. The bootstrapped *t*-statistic is defined as the 97.5% percentile of the simulated *t*-statistics.

h	t -statistic	Bootstrapped t -statistic		t -statistic	Bootstrapped t -statistic	
		Wild	Pairs		Wild	Pairs
Panel A: Forecast with return lagged h months						
	$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h r_{t-h+1}^i/\sigma_{t-h}^i + \varepsilon_{t+1}^i$				$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h \text{sign}(r_{t-h+1}^i) + \varepsilon_{t+1}^i$	
1	3.71	10.68	4.20	3.75	9.31	4.19
2	0.97	4.07	1.68	1.34	3.65	2.02
3	2.48	7.43	3.11	2.44	6.09	3.01
4	0.22	2.40	1.14	0.65	2.28	1.59
5	-0.15	1.53	0.67	-0.38	1.56	0.66
6	0.52	3.08	1.30	1.35	2.78	2.15
7	0.39	3.07	1.24	0.95	2.74	1.84
8	0.59	3.32	1.37	1.20	2.84	2.06
9	1.68	5.26	2.42	1.96	4.59	2.66
10	2.70	7.37	3.32	2.11	5.83	2.84
11	3.70	10.37	4.23	4.04	7.61	4.54
12	0.37	2.74	1.14	0.54	2.39	1.37
Panel B: Forecast with past h -month returns						
	$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h r_{t-h,t}^i/\sigma_{t-1}^i + \varepsilon_{t+1}^i$				$r_{t+1}^i/\sigma_t^i = \alpha_h + \beta_h \text{sign}(r_{t-h,t}^i) + \varepsilon_{t+1}^i$	
1	3.71	10.68	4.20	3.75	9.31	4.19
2	3.09	9.53	3.54	3.19	8.39	3.70
3	3.74	11.53	4.27	4.43	9.96	4.94
4	3.45	10.37	3.98	4.78	9.19	5.19
5	3.06	9.27	3.63	4.36	8.39	4.85
6	3.17	9.31	3.73	4.03	8.52	4.46
7	3.05	9.09	3.60	4.19	8.32	4.62
8	3.12	9.06	3.69	4.13	8.11	4.58
9	3.59	10.31	4.13	4.70	9.27	5.13
10	4.00	11.68	4.54	5.46	10.44	5.85
11	4.69	13.16	5.14	5.64	11.77	6.07
12	4.48	12.76	4.96	5.24	11.17	5.62

Table 8 reports the *t*-statistics from the pooled regression and the bootstrapped *t*-statistics. Compared with Table 4, after controlling for the fixed effects, the *t*-statistic is smaller than that without controlling for fixed effects. For example, when predicting the next one-month return with the past 12-month return, the *t*-statistic is 4.34 in Table 4 and 3.37 in Table 8. The most important result is that the *t*-statistic of 3.37 when controlling for fixed effects does not affect the conclusion that insufficient evidence exists in support of TSM. In the Online Appendix, we show that this finding is robust to the cases within each asset class and without volatility scaling.

4.5. Out-of-sample performance

A further implicit assumption of a pooled regression is that all 55 futures contracts are homogenous with the same slope in Eqs. (3) and (5). If the individual slopes are all identical, the pooled estimate converges to the common slope, and pooling data leads to a more precise estimate than the individual time series regression estimate. Whether the slopes of all individual assets are identical or not, there is no guarantee that pooling the data will help. Nevertheless, whether pooling the data improves out-

of-sample performance is of interest to examine empirically.

Fig. 4 plots the out-of-sample R_{OS}^2 for each individual asset. In Panel A, we present the results with volatility scaling when running the pooled regression. To predict returns in month $t + 1$ at the end of month t , we run pooled regression Eq. (3) with returns up to month t as MOP (Moskowitz et al., 2012). Let $\hat{\alpha}_t$ and $\hat{\beta}_t$ be the estimated intercept and slope. We calculate the expected return of asset i for month $t + 1$ as

$$E_t(r_{t+1}^i) = \hat{\alpha}_t \sigma_t^i + \hat{\beta}_t \frac{r_{t-12,t}^i}{\sigma_{t-1}^i} \sigma_t^i, \quad (15)$$

which can be plugged into Eq. (2) to calculate the R_{OS}^2 accordingly.

In comparison with earlier univariate regressions (Table 2), the pooled regression does improve the out-of-sample forecasting performance in some of the markets. The R_{OS}^2 is significantly positive for three commodity futures contracts: cocoa, copper, and gold. Of the nine international equity markets, six are significant at the 10% level, with the remaining three positive but not significant. The average R_{OS}^2 in the equity markets is

Table 8

t-statistic of pooled regression controlling for fixed effects.

This table reports the *t*-statistic of pooled regression with real data and the bootstrapped *t*-statistics with wild and pairs bootstrap. For each asset, we bootstrap a path with *T* observations and run pooled regression controlling for fixed effects to calculate the *t*-statistic. We repeat this procedure one thousand times and obtain the distribution of the *t*-statistic for testing the null hypothesis that there is no time series momentum. The bootstrapped *t*-statistic is defined as the 97.5% percentile of the simulated *t*-statistics. The sample period is 1985:01–2015:12.

h	t -statistic	Bootstrapped t -statistic		t -statistic	Bootstrapped t -statistic	
		Wild	Pairs		Wild	Pairs
Panel A: Forecast with return lagged h months						
	$r_{t+1}^i/\sigma_t^i = \alpha_h^i + \beta_h r_{t-h+1}^i/\sigma_{t-h}^i + \varepsilon_{t+1}^i$				$r_{t+1}^i/\sigma_t^i = \alpha_h^i + \beta_h \text{sign}(r_{t-h+1}^i) + \varepsilon_{t+1}^i$	
1	2.80	8.51	3.39	2.66	7.60	3.19
2	0.96	4.17	1.66	0.94	3.85	1.67
3	2.53	7.77	3.12	2.17	6.36	2.83
4	−0.19	1.56	0.70	0.36	1.56	1.27
5	−0.56	1.00	0.25	−0.94	1.03	0.02
6	0.58	3.26	1.36	1.07	2.90	1.79
7	−0.62	0.59	0.27	0.20	1.01	1.04
8	0.37	2.90	1.14	0.80	2.64	1.53
9	0.94	3.84	1.73	0.94	3.53	1.79
10	2.57	7.21	3.22	1.87	5.71	2.61
11	3.22	9.40	3.75	3.53	7.03	4.17
12	−0.12	1.63	0.65	0.37	1.70	1.13
Panel B: Forecast with past h -month returns						
	$r_{t+1}^i/\sigma_t^i = \alpha_h^i + \beta_h r_{t-h,t}^i/\sigma_{t-1}^i + \varepsilon_{t+1}^i$				$r_{t+1}^i/\sigma_t^i = \alpha_h^i + \beta_h \text{sign}(r_{t-h,t}^i) + \varepsilon_{t+1}^i$	
1	2.80	8.51	3.39	2.66	7.60	3.19
2	2.51	8.43	3.07	2.62	7.41	3.19
3	3.23	10.17	3.74	3.56	9.08	4.17
4	2.89	9.24	3.46	3.60	8.36	4.11
5	2.44	7.89	2.99	3.17	7.45	3.66
6	2.53	7.97	3.12	3.15	7.27	3.70
7	2.22	7.32	2.86	2.97	6.86	3.43
8	2.19	7.35	2.80	2.55	6.67	3.24
9	2.49	7.75	3.09	3.43	7.19	3.92
10	3.00	9.23	3.58	3.94	8.34	4.38
11	3.68	10.80	4.20	4.49	9.71	4.94
12	3.37	10.13	3.93	4.04	9.14	4.53

2.08%, indicating potential economic significance as well (Campbell and Thompson, 2008). The R_{OS}^2 s in the bond and currency markets are generally negative or slightly positive. The two exceptions are the two-year European bond and two-year US bond. Overall, if there is any TSM, it appears to show up in the international equity markets only, not present in the entire cross section of assets.

Panel B of Fig. 4 plots the R_{OS}^2 for each asset without volatility scaling when running regression Eq. (3). The out-of-sample performance does not change significantly. Untabulated results show that the value of R_{OS}^2 with volatility scaling is generally similar to that without volatility scaling. Two exceptions are the two-year European bond and two-year US bond, which have extreme positive R_{OS}^2 in the case with volatility scaling (15.18% and 3.48%) but extreme negative R_{OS}^2 in the case without volatility scaling (−19.54% and −16.71%). As such, the average R_{OS}^2 using volatility scaling is −0.06%, which is larger than −0.35% without volatility scaling.

Overall, to a certain extent, a pooled regression can improve the out-of-sample forecasting power relative to the asset-by-asset time series regression, but such improvement is restricted to some specific assets. For the entire cross section of assets, it does little to improve their out-

of-sample forecasting performances, and it cannot provide significant support for TSM either.

5. Trading strategy

In this section, we examine the source of profitability of the TSM strategy proposed by MOP (Moskowitz et al., 2012). We show in various ways that its performance does not necessarily indicate that TSM exists across assets.

5.1. TSM versus TSH at asset level

The early univariate regressions show that time series predictability is not a common feature across assets, which suggests that the performance of the TSM strategy perhaps is not attributed to predictability, at least not entirely. Furthermore, it raises the possibility that some strategies that do not require predictability can perform as well as the TSM strategy. As it turns out, this is the case.

Suppose the return of asset *i* follows an independent and identically distributed normal distribution with mean μ^i and volatility σ^i . Then, the probability of the past

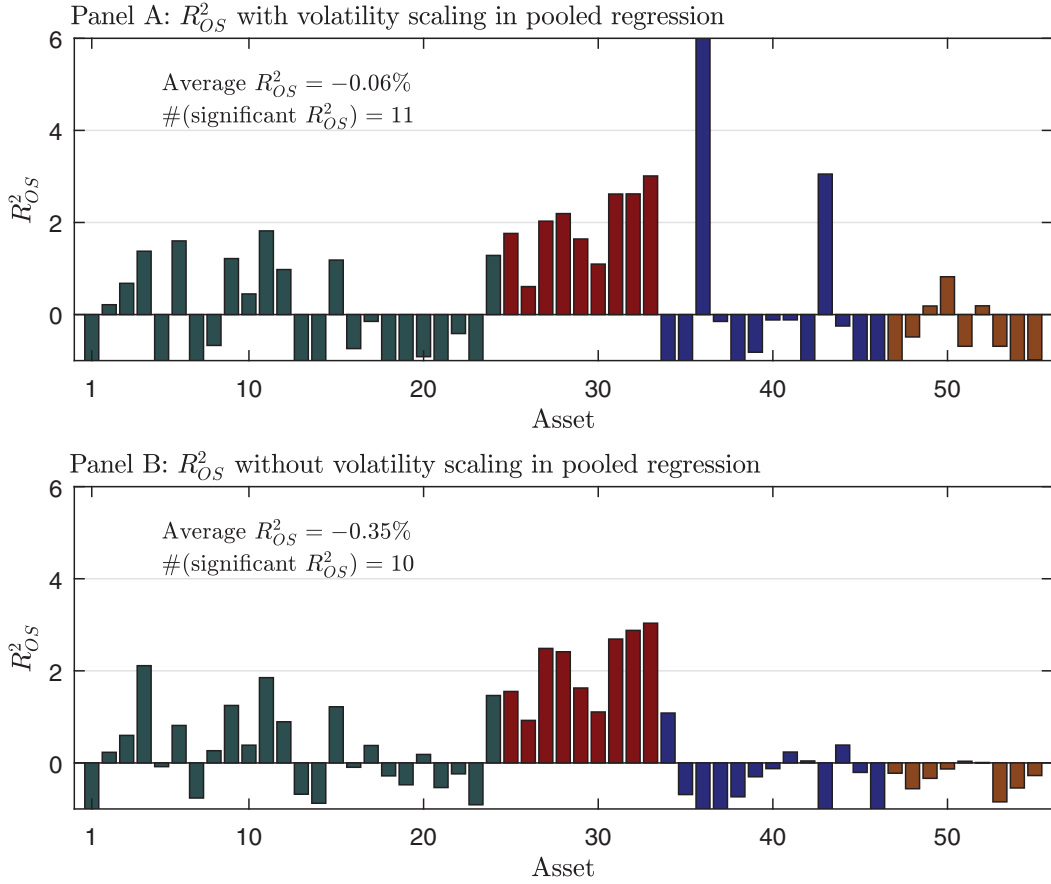


Fig. 4. Time series momentum (TSM) with pooled regression: out-of-sample performance. This figure plots the out-of-sample R^2_{OS} of forecasting a futures contract return with pooled regression

$$r_{t+1}^i/\sigma_t^i = \alpha + \beta r_{t-12,t}^i/\sigma_{t-1}^i + \varepsilon_{t+1}^i,$$

for Panel A and

$$r_{t+1}^i = \alpha + \beta r_{t-12,t}^i + \varepsilon_{t+1}^i$$

for Panel B, where $r_{t-12,t}^i$ is asset i 's past 12-month return. We calculate asset i 's out-of-sample R^2_{OS} by applying the same $\hat{\alpha}_t$ and $\hat{\beta}_t$ to all assets in estimating the expected return as $E_t(r_{t+1}^i) = \hat{\alpha}_t\sigma_t^i + \hat{\beta}_t\frac{r_{t-12,t}^i}{\sigma_{t-1}^i}\sigma_t^i$ in Panel A and $E_t(r_{t+1}^i) = \hat{\alpha}_t + \hat{\beta}_t r_{t-12,t}^i$ in Panel B, where $\hat{\alpha}_t$ and $\hat{\beta}_t$ are the pooled regression estimates with data up to month t . The out-of-sample period is 2000:01–2015:12.

12-month return being positive is

$$\begin{aligned} \Pr(r_{t-12,t}^i > 0) &= 1 - \Pr\left(\frac{r_{t-12,t}^i - 12\mu^i}{\sqrt{12}\sigma^i} \leq -\sqrt{12}\frac{\mu^i}{\sigma^i}\right) \\ &= \Phi(\sqrt{12}\mu^i/\sigma^i), \end{aligned} \quad (16)$$

where $\Phi(\cdot)$ is the $N(0, 1)$ cumulative distribution function. Hence, without time series predictability, the TSM strategy tends to buy an asset with high mean return (i.e., Sharpe ratio). Based on this observation, we consider an alternative strategy based on the time series history of asset i 's return

$$r_{t+1}^{\text{TSH},i} = \text{sign}(r_{1,t}^i)r_{t+1}^i, \quad (17)$$

where $r_{1,t}^i$ is the accumulative return of asset i from month 1 to month t or the historical sample mean multiplied by t .

Volatility scaling on the TSH strategy is unnecessary because we compare the TSM and TSH strategies at the asset

level in this section. Without volatility scaling, the corresponding return of the TSM strategy in month $t + 1$ for asset i is

$$r_{t+1}^{\text{TSM},i} = \text{sign}(r_{t-12,t}^i)r_{t+1}^i. \quad (18)$$

Comparing Eqs. (17) and (18), the TSM strategy attempts to exploit possible predictability of the past 12-month return, and the TSH does not rely on any predictability at all. Our goal here is to examine their performances across assets.

Table 9 reports the average returns and Sharpe ratios, and their differences, of the TSM and TSH strategies based on Eqs. (18) and (17), respectively. The results show that the TSM strategy generally performs the same as the TSH strategy. Of the 55 assets, only five show that the TSM strategy generates a higher average return than the TSH strategy. When we use the Sharpe ratio as the performance measure, the results remain unchanged. Thus, the TSM strategy does not significantly outperform at the asset level the TSH strategy that does not require predictability.

Table 9

Time series momentum (TSM) versus time series history (TSH) at the asset level.

This table reports the mean returns and Sharpe ratios of the time series momentum and time series history strategies, as well as their difference, on the basis of individual assets. TSM refers to the strategy that buys the future contract if its past 12-month return is non-negative and sells it if its past 12-month return is negative, and TSH refers to the strategy that buys the futures contract if its historical sample mean is non-negative and sells it if its historical sample mean is negative. ***, **, and * denote statistical significance at the 1%, 5%, and 10% level, respectively. The investment period is 1986:01–2015:12.

Asset	TSM return	TSH return	TSM Sharpe ratio	TSH Sharpe ratio	Return difference	p-value of return difference	Sharpe ratio difference	p-value of Sharpe ratio difference
Aluminum	0.27	−0.47	0.05	−0.08	0.74**	0.04	0.13**	0.04
Brent oil	0.80	0.32	0.09	0.04	0.48	0.44	0.05	0.44
Cattle	0.28	0.08	0.07	0.02	0.20	0.48	0.05	0.48
Cocoa	−0.46	0.13	−0.06	0.02	−0.59	0.32	−0.08	0.31
Coffee	0.18	−0.55	0.02	−0.05	0.73	0.30	0.07	0.30
Copper	0.77	0.94	0.10	0.12	−0.17	0.74	−0.02	0.73
Corn	0.11	0.14	0.01	0.02	−0.04	0.94	−0.01	0.94
Cotton	1.04	−0.07	0.14	−0.01	1.11**	0.04	0.15**	0.04
Crude	1.07	0.21	0.11	0.02	0.86	0.19	0.09	0.19
Gas oil	0.98	0.53	0.10	0.05	0.45	0.49	0.05	0.49
Gold	0.55	0.05	0.12	0.01	0.50*	0.10	0.11*	0.10
Heating oil	1.09	0.30	0.12	0.03	0.79	0.21	0.09	0.21
Hogs	0.29	−0.17	0.04	−0.02	0.46	0.37	0.06	0.37
Natural gas	1.26	0.05	0.09	0.01	1.21	0.25	0.08	0.25
Nickel	0.72	0.43	0.07	0.04	0.29	0.70	0.03	0.70
Platinum	0.30	0.53	0.05	0.09	−0.23	0.61	−0.04	0.61
Silver	0.33	−0.09	0.04	−0.01	0.42	0.46	0.05	0.47
Soybean	−0.15	0.02	−0.02	0.01	−0.17	0.68	−0.03	0.67
Soymeal	0.20	0.51	0.02	0.06	−0.31	0.58	−0.04	0.58
Soy oil	0.42	0.14	0.06	0.02	0.28	0.57	0.04	0.57
Sugar	0.05	0.49	0.01	0.05	−0.44	0.50	−0.04	0.50
Unleaded gasoline	1.00	0.83	0.10	0.08	0.17	0.77	0.02	0.77
Wheat	0.38	0.26	0.05	0.03	0.12	0.81	0.02	0.81
Zinc	0.67	−0.22	0.09	−0.03	0.89*	0.08	0.12*	0.08
SPI 200	0.18	0.55	0.04	0.12	−0.37	0.17	−0.08	0.17
DAX	0.79	0.68	0.13	0.11	0.11	0.80	0.02	0.79
IBEX 35	0.71	0.72	0.11	0.11	−0.01	0.98	0.00	0.98
CAC 40	0.43	0.49	0.08	0.09	−0.06	0.89	−0.01	0.88
FTSE/MIB	0.90	0.36	0.14	0.05	0.54	0.27	0.09	0.27
TOPIX	0.84	0.25	0.15	0.04	0.59	0.18	0.11	0.18
AEX	0.73	0.55	0.13	0.10	0.18	0.66	0.03	0.66
FTSE 100	0.27	0.52	0.06	0.11	−0.25	0.40	−0.05	0.40
S&P 500	0.67	0.73	0.15	0.16	−0.06	0.84	−0.01	0.83
Three-year Australian bond	0.24	0.28	0.24	0.28	−0.04	0.39	−0.04	0.37
Ten-year Australian bond	0.32	0.42	0.15	0.21	−0.10	0.27	−0.06	0.26
Two-year European bond	0.13	0.10	0.13	0.10	0.03	0.57	0.03	0.56
Five-year European bond	0.08	0.13	0.06	0.11	−0.05	0.46	−0.05	0.45
Ten-year European bond	0.05	0.25	0.02	0.09	−0.20	0.18	−0.07	0.18
Thirty-year European bond	0.14	0.56	0.05	0.19	−0.42***	0.01	−0.14***	0.01
Ten-year Canadian bond	0.34	0.52	0.11	0.17	−0.18	0.32	−0.06	0.31
Ten-year Japanese bond	0.21	0.08	0.06	0.02	0.13	0.55	0.04	0.56
Ten-year UK bond	0.34	0.28	0.14	0.12	0.06	0.68	0.02	0.68
Two-year US bond	0.13	0.12	0.28	0.26	0.01	0.65	0.02	0.63
Five-year US bond	0.19	0.23	0.15	0.19	−0.04	0.41	−0.04	0.40
Ten-year US bond	0.19	0.28	0.09	0.13	−0.09	0.40	−0.04	0.40
Thirty-year US bond	0.54	0.89	0.12	0.20	−0.35*	0.07	−0.08*	0.06
AUD/USD	0.06	−0.16	0.02	−0.05	0.22	0.37	0.07	0.37
EUR/USD	0.11	0.11	0.03	0.03	0.00	1.00	0.00	1.00
CAD/USD	0.23	−0.08	0.11	−0.04	0.31**	0.05	0.15**	0.05
JPY/USD	0.46	0.09	0.14	0.03	0.37	0.11	0.11	0.11
NOK/USD	0.06	−0.06	0.02	−0.02	0.12	0.58	0.04	0.58
NZD/USD	0.24	0.02	0.07	0.01	0.22	0.38	0.06	0.38
SEK/USD	0.04	−0.05	0.01	−0.02	0.09	0.70	0.03	0.70
CHF/USD	0.18	0.19	0.05	0.06	−0.01	0.96	−0.01	0.97
GBP/USD	0.00	−0.03	0.00	−0.01	0.03	0.87	0.01	0.87
#(significance)					7		7	

5.2. TSM versus TSH at portfolio level

Even though no time series predictability exists, the TSM strategy could still be profitable. Conrad and Kaul (1988), Jegadeesh (1990), and Jegadeesh and Titman

(1993) note that if there are differences in mean returns, a strategy that buys high-mean assets using the proceeds from selling low-mean assets has a natural tilt toward high-mean assets. To see this, consider just two assets. If the mean return of the first asset far exceeds that of the

second, buying the first and shorting the second is profitable. Because the past 12-month return can be viewed as an estimate of the mean return, the TSM strategy could profit from the differences in mean returns. If this is the case, it will unlikely outperform the TSH strategy at the portfolio level.

To make the two strategies comparable, we consider the following equal-weighting scheme for the TSM and TSH:

$$r_{t+1}^{\text{TSM}} = \frac{1}{N_t} \sum_{i=1}^{N_t} \text{sign}(r_{t-12,t}^i) r_{t+1}^i \quad (19)$$

and

$$r_{t+1}^{\text{TSH}} = \frac{1}{N_t} \sum_{i=1}^{N_t} \text{sign}(r_{1,t}^i) r_{t+1}^i, \quad (20)$$

where N_t is the number of assets investable at time t .⁷ By doing so, the two strategies differ only in how the past information is used to select the assets. So, differences in performance stem from the differences in asset selection, not from differences in scaling the portfolio weights. Because our goal here is not to improve the performance, the concern that the TSM strategy tilts weighting toward low volatility assets is not an issue, because the TSH strategy has the same tilts. Nevertheless, we will examine volatility weighting.

Panel A of Table 10 presents the average and risk-adjusted returns of the two strategies, the second of which is computed from two benchmark asset pricing models as in MOP (Moskowitz et al., 2012). The first model is the Fama–French four-factor model that uses the MSCI World Index as the market factor, and the second is the Asness et al. (2013) three-factor model with the MSCI World Index, the value everywhere factor, and the momentum everywhere factor. Over the 1986 to 2015 investment period, the average return differential between the two strategies is as small as 0.14%, not significant with a p -value of 0.19. The results suggest that the TSM strategy generates virtually the same average portfolio return as the TSH strategy, consistent with earlier comparison at the asset level.

When turning to the risk-adjusted returns, the TSM and TSH alphas are 0.15% (t -statistic = 1.94) and 0.05% (t -statistic = 0.80) with the Fama–French four-factor model and 0.07% (t -statistic = 1.01) and 0.09% (t -statistic = 1.14) with the Asness et al. (2013) three-factor model, respectively. As for the case with average return, the two strategies' alpha differentials with the two models are 0.10% (p -value = 0.29) and -0.02% (p -value = 0.84) and, therefore, are not statistically significant from zero. Thus, the TSM and TSH strategies do not generate sizable abnormal returns. Moreover, their difference in alpha is even smaller than the difference in average return and is also insignificant.

Also reported in Panel A of Table 10 are the average and risk-adjusted returns of the long- and short-leg portfolios

of the TSM and TSH strategies. The results show that the performance of the two strategies mainly stems from the long legs and that the performance of their short legs is always indifferent from zero. This new finding, not shown by MOP (Moskowitz et al., 2012) or Goyal and Jegadeesh (2018), is consistent with our argument that the strong TSM performance is due to the difference in mean returns.

For robustness, we also consider three alternative portfolio weighting schemes: volatility weighting as in MOP (Moskowitz et al., 2012), past-12-month-return weighting, and equal weighting with a zero-investment constraint as in Goyal and Jegadeesh (2018). The results do not change qualitatively, and the alpha differential between the TSM and TSH strategies is always indifferent from zero. The On-line Appendix considers the case of constructing the TSM with the past six-month return, instead of the past 12-month return, and the results remain unchanged. In sum, the performance of the TSM strategy in MOP (Moskowitz et al., 2012) seems mainly stemming from the difference in mean returns, not from the times series predictability of the past 12-month return.

5.3. TSM and TSH forecast comparison: predictive slope

Lewellen (2015) proposes an interesting predictive slope that assesses the degree of predictability of cross-sectional forecasts in an elegant way, in which one simply runs a cross-sectional regression of the realized returns on the forecasts. If the forecasts are perfect, the slope should be one. Generally, a value less than 0.5 indicates no predictability as the forecasts under-perform naive forecasts.

At the end of month t , we calculate the expected return of asset i as $\hat{r}_{t+1}^{\text{TSM},i}$, which is estimated by pooled regression Eq. (3) with data up to month t , and then we regress month $t+1$ return r_{t+1}^i on $\hat{r}_{t+1}^{\text{TSM},i}$. The first three columns of Table 11 report the results. Consistent with the earlier no predictability results, the regression slope is close to zero and its t -statistic is less than one standard deviation, suggesting that the in-sample performance with pooled regression Eq. (3) is not reliable and that the TSM estimates do not line up with the true expected returns out of sample.

We also explore whether the TSM and TSH forecasts have the same mean. The last three columns of Table 11 report the results of regressing the TSM forecasts of expected returns on the TSH forecasts. Because these two estimates are potentially unbiased, we do not include an intercept to improve the estimate efficiency. Consistent with earlier results indicating little difference between the two strategies, the slope is close to one and the t -statistic is much larger than two standard deviations. For robustness, we also perform the tests for each asset class, and the results are the same as the overall case. In short, the TSM strategy has little predictive power and behaves in a very similar manner to the TSH strategy.

5.4. When does the TSM outperform the TSH?

In the previous sections, we have shown that the predictability of the past 12-month return is weak with real data and that the TSM strategy performs similarly as the

⁷ TSH and TSM have similar expressions here, in parallel with the asset-level comparison. At the portfolio level, we also explore two alternative TSH strategies that buy (sell) half or one-third of the assets with high (low) historical sample means and find that their performances are quantitatively similar.

Table 10

Time series momentum (TSM) versus time series history (TSH) at the portfolio level.

This table reports the average and risk-adjusted returns of the TSM and TSH strategies, where we restrict portfolio weights on individual assets to be the same for comparison when constructing these two strategies. TSM refers to the strategy that buys futures contracts with non-negative past 12-month return and sells futures contracts with negative past 12-month return, and TSH refers to the strategy that buys futures contracts with non-negative historical sample mean and sells futures contracts with negative historical sample mean. The benchmarks are the Fama-French four-factor model that includes MSCI World Index, SMB (small minus big), HML (high minus low), and UMD (momentum), and the [Asness et al. \(2013\)](#) three-factor model. Newey-West *t*-statistics and *p*-values are reported in parentheses and brackets, respectively. ***, **, and * denote statistical significance at the 1%, 5%, and 10% level, respectively. The investment period is 1986:01–2015:12.

Fama–French four-factor model													Asness-Moskowitz-Pedersen three-factor model				
	Mean	Alpha	MSCI World Index	SMB	HML	UMD	R ²	Alpha	MSCI World Index	Value everywhere	Momentum everywhere	R ²					
Panel A: Equal weighting, i.e., portfolio weight = $\frac{1}{N}$																	
TSM strategy																	
Long leg	0.34*** (4.92)	0.12** (2.29)	0.16*** (7.73)	0.05** (2.24)	0.09*** (2.62)	0.32*** (7.19)	42.57%	0.09 (1.60)	0.16*** (7.21)	0.14*** (2.99)	0.38*** (7.33)	41.92%					
Short leg	−0.05 (−0.72)	−0.03 (−0.46)	0.14*** (5.30)	0.11*** (4.01)	0.03 (0.96)	−0.28*** (−8.34)	48.31%	0.02 (0.23)	0.13*** (5.08)	−0.09 (−1.42)	−0.32*** (−6.89)	45.85%					
Long - short	0.39*** (4.73)	0.15* (1.94)	0.02 (0.61)	−0.06* (−1.83)	0.06 (1.01)	0.60*** (9.99)	46.03%	0.07 (1.01)	0.03 (0.93)	0.23** (2.50)	0.70*** (8.79)	47.39%					
TSH strategy																	
Long leg	0.27*** (2.56)	0.07 (0.95)	0.28*** (8.79)	0.14*** (4.41)	0.10*** (3.90)	0.08* (1.93)	49.07%	0.10 (1.13)	0.26*** (7.73)	0.01 (0.24)	0.08* (1.65)	45.30%					
Short leg	0.02 (0.61)	0.02 (0.52)	0.03*** (3.51)	0.03* (1.83)	0.02 (1.47)	−0.04** (−2.01)	8.73%	0.01 (0.25)	0.03*** (3.17)	0.03 (1.15)	−0.03 (−1.04)	8.04%					
Long - short	0.25*** (2.70)	0.05 (0.80)	0.25*** (8.54)	0.11*** (3.70)	0.08*** (2.96)	0.13*** (2.76)	44.83%	0.09 (1.14)	0.23*** (7.59)	−0.02 (−0.29)	0.11* (1.94)	42.01%					
TSM versus TSH																	
Mean difference	0.14 [0.19]																
Alpha difference		0.10 [0.26]						−0.02 [0.84]									
Panel B: Volatility weighting, i.e., portfolio weight = $\frac{1}{N} \frac{40\%}{\sigma_t^2}$																	
TSM strategy																	
Long leg	1.02*** (7.66)	0.58*** (5.31)	0.28*** (7.34)	−0.04 (−0.67)	0.12 (1.85)	0.68*** (8.02)	36.04%	0.48*** (4.29)	0.28*** (7.96)	0.36*** (3.21)	0.84*** (8.41)	37.46%					
Short leg	−0.14 (−1.07)	−0.06 (−0.53)	0.22*** (5.75)	0.19*** (3.06)	0.04 (0.79)	−0.52*** (−8.67)	44.46%	0.02 (0.13)	0.20*** (5.52)	−0.17 (−1.48)	−0.59*** (−9.19)	42.42%					
Long - short	1.16*** (6.31)	0.64*** (3.68)	0.05 (0.80)	−0.22** (−2.24)	0.08 (0.85)	1.19*** (10.31)	40.62%	0.46*** (2.90)	0.08 (1.25)	0.53*** (2.75)	1.43*** (10.81)	41.58%					
TSH strategy																	
Long leg	0.81*** (4.72)	0.42*** (3.06)	0.48*** (12.17)	0.12** (2.29)	0.16** (2.48)	0.25*** (2.97)	40.10%	0.43*** (2.71)	0.46*** (10.96)	0.13 (1.03)	0.29*** (3.05)	38.80%					
Short leg	0.07 (1.50)	0.10* (1.93)	0.02 (1.28)	0.03 (1.33)	−0.01 (−0.11)	−0.08*** (−2.73)	4.27%	0.07 (1.39)	0.02 (1.42)	0.06 (1.07)	−0.04 (−1.06)	4.35%					
Long - short	0.74*** (4.36)	0.32** (2.18)	0.46*** (9.81)	0.09 (1.46)	0.16** (2.17)	0.33*** (3.76)	36.23%	0.36** (2.25)	0.44*** (8.85)	0.06 (0.43)	0.33*** (3.21)	35.01%					
TSM versus TSH																	
Mean difference	0.42* [0.07]																
Alpha difference		0.32* [0.08]						0.10 [0.56]									
Panel C: Past 12-month return weighting, i.e., portfolio weight = $\frac{r_{t-12,t}^j}{\sum_{i=1}^N r_{t-12,t}^i }$																	
TSM strategy																	
Long leg	0.50*** (3.54)	0.08 (0.72)	0.27*** (5.85)	0.17*** (3.62)	0.14** (2.06)	0.66*** (8.65)	32.83%	0.06 (0.52)	0.25*** (5.09)	0.15* (1.65)	0.73*** (7.71)	30.95%					
Short leg	−0.07 (−0.65)	−0.01 (−0.08)	0.19*** (5.06)	0.16*** (3.61)	0.01 (0.18)	−0.43*** (−7.74)	39.93%	0.05 (0.49)	0.18*** (5.06)	−0.15 (−1.36)	−0.49*** (−6.97)	38.05%					
Long - short	0.57*** (3.79)	0.09 (0.69)	0.08 (1.48)	0.01 (0.14)	0.13 (1.35)	1.09*** (12.55)	40.17%	0.01 (0.08)	0.08 (1.57)	0.30** (2.04)	1.22*** (11.33)	40.55%					
TSH strategy																	
Long leg	0.30* (1.78)	−0.03 (−0.23)	0.43*** (8.10)	0.23*** (4.71)	0.14*** (3.22)	0.19*** (2.77)	48.23%	0.03 (0.21)	0.41*** (6.70)	−0.03 (−0.29)	0.17** (1.96)	44.26%					
Short leg	0.02 (0.77)	0.02 (0.83)	0.02*** (2.82)	0.01 (1.09)	0.01 (1.06)	−0.04* (−1.70)	5.23%	0.01 (0.50)	0.02*** (2.64)	0.03 (1.08)	−0.03 (−0.86)	5.18%					
Long - short	0.28* (1.71)	−0.05 (−0.42)	0.41*** (7.87)	0.22*** (4.39)	0.13*** (2.93)	0.23*** (3.10)	46.18%	0.02 (0.10)	0.39*** (6.85)	−0.06 (−0.56)	0.20** (1.99)	42.65%					

(continued on next page)

Table 10 (continued)

	Mean	Alpha	Fama–French four-factor model				R ²	Asness-Moskowitz-Pedersen three-factor model				
			MSCI World Index	SMB	HML	UMD		Alpha	MSCI World Index	Value everywhere	Momentum everywhere	R ²
TSM versus TSH												
Mean difference	0.29 [0.12]											
Alpha difference		0.14 [0.34]						-0.01 [0.98]				
Panel D: Zero investment, i.e., long = $\frac{1}{N^{\text{long}}}$ and short = $\frac{1}{N^{\text{short}}}$												
TSM strategy												
Long leg	0.60*** (5.29)	0.24*** (2.65)	0.25*** (6.93)	0.08** (2.02)	0.11** (2.26)	0.53*** (7.60)	39.74%	0.20** (2.14)	0.24*** (6.41)	0.17** (2.19)	0.61*** (7.17)	39.10%
Short leg	-0.12 (-0.74)	-0.06 (-0.42)	0.28*** (5.76)	0.23*** (3.94)	0.05 (0.79)	-0.56*** (-7.40)	42.77%	0.02 (0.10)	0.26*** (5.33)	-0.17 (-1.24)	-0.64*** (-6.73)	40.27%
Long - short	0.72*** (4.32)	0.30** (1.96)	-0.03 (-0.65)	-0.16** (-2.43)	0.06 (0.69)	1.10*** (10.26)	45.91%	0.18 (1.25)	-0.02 (-0.40)	0.34** (1.97)	1.25*** (8.52)	46.18%
TSH strategy												
Long leg	0.35*** (2.64)	0.10 (1.06)	0.35*** (8.98)	0.17*** (4.39)	0.12*** (3.85)	0.10* (1.90)	49.49%	0.13 (1.25)	0.33*** (7.92)	0.02 (0.21)	0.10 (1.56)	45.78%
Short leg	0.08 (0.51)	0.04 (0.30)	0.15*** (3.66)	0.14** (2.13)	0.09 (1.62)	-0.18* (-1.86)	8.16%	-0.01 (-0.05)	0.14*** (3.28)	0.17 (1.48)	-0.09 (0.75)	7.37%
Long - short	0.27** (2.10)	0.06 (0.40)	0.20*** (5.30)	0.03 (0.45)	0.03 (0.49)	0.28*** (2.75)	11.89%	0.14 (1.02)	0.19*** (4.83)	-0.16 (-1.24)	0.19 (1.44)	12.27%
TSM versus TSH												
Mean difference	0.45** [0.03]											
Alpha difference		0.24 [0.13]						0.04 [0.76]				

Table 11

Time series momentum (TSM) and time series history (TSH) forecast comparison.

This table reports the results of regressing r_{t+1}^i on the expected return ($\hat{r}_{t+1}^{\text{TSM},i}$) estimated at time t with the TSM pooled regression Eq. (3) and regressing $\hat{r}_{t+1}^{\text{TSM},i}$ on the expected return ($\hat{r}_{t+1}^{\text{TSH},i}$) estimated with the TSH approach (i.e., historical sample mean). ***, **, and * denote statistical significance at the 1%, 5%, and 10% level, respectively.

Asset class	$r_{t+1}^i = \alpha + \beta \hat{r}_{t+1}^{\text{TSM},i} + \varepsilon_{t+1}^i$			$\hat{r}_{t+1}^{\text{TSM},i} = d \hat{r}_{t+1}^{\text{TSH},i} + u_t^i$		
	β	t -statistic	R ²	d	t -statistic	R ²
Panel A: $\hat{r}_{t+1}^{\text{TSM},i}$ is estimated with volatility scaling						
Overall	0.19	0.61	0.04	1.09***	18.56	40.33
Commodity	0.15	0.42	0.02	1.24***	11.62	23.53
Equity	0.07	0.10	0.01	0.84***	14.90	45.06
Bond	0.23	0.60	0.08	0.99***	68.75	92.27
Currency	-0.08	-0.12	0.01	1.01***	14.95	4.45
Panel B: $\hat{r}_{t+1}^{\text{TSM},i}$ is estimated without volatility scaling						
Overall	0.30	0.45	0.03	1.04***	41.89	54.96
Commodity	0.09	0.11	0.00	1.01***	26.53	37.65
Equity	-0.37	-0.32	0.07	0.93***	35.34	77.93
Bond	-0.49	-0.52	0.07	1.00***	72.40	91.38
Currency	0.03	0.03	0.00	1.64***	19.78	14.32

TSH strategy that does not require any predictability. This section attempts to answer another question: If the predictability is strong, to what extent does the TSM strategy outperform the TSH strategy?⁸

Consistent with the pooled regression of MOP, we assume the following data-generating process:

$$r_{t+1}^i = \alpha^i + \beta \frac{r_{t-12,t}^i}{12} + \varepsilon_{t+1}^i, \quad (21)$$

⁸ We thank the anonymous referee for this intriguing research question.

where $\beta = 0.1, 0.2$, and 0.4 (we divide the past 12-month return $r_{t-12,t}^i$ by 12 to make β easy to interpret; β is 0.08 for the real data). There are two ways to draw the residuals. The first is to ignore the off-diagonal elements and draw the residuals asset by asset, thereby assuming no cross-asset predictability, and the second is to keep the covariance matrix structure and draw the residuals across assets. For a given specification and beta, we simulate a path for the 55 assets with $T = 372$ observations and construct the TSM and TSH strategies accordingly. We repeat this procedure one thousand times to test whether these two strategies generate the same mean returns. Empirically, we

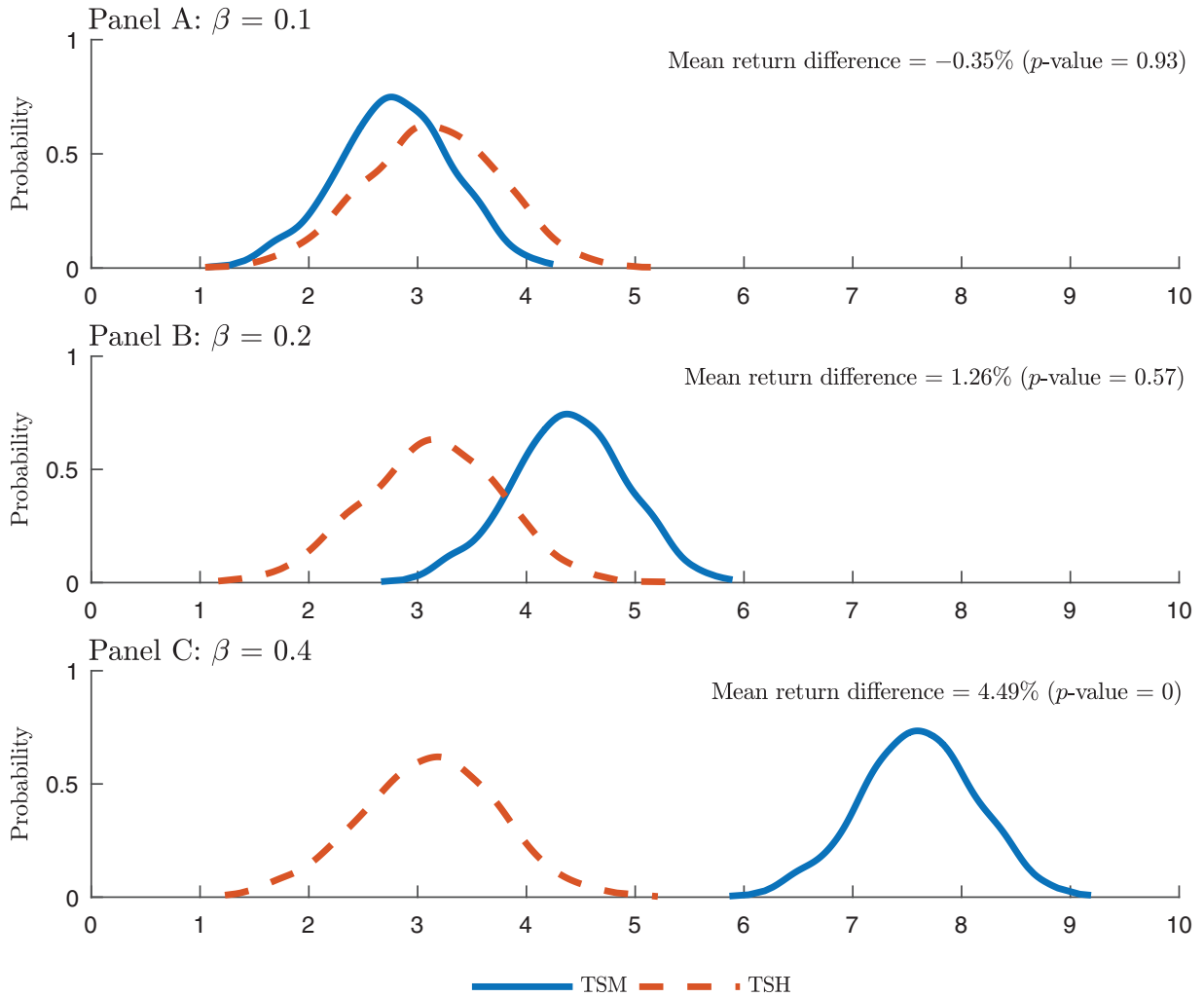


Fig. 5. Annualized mean return difference between time series momentum (TSM) and time series history (TSH). This figure plots the distributions of simulated annualized mean returns of the TSM and TSH strategies, where each asset is assumed to follow

$$r_{t+1}^i = \alpha^i + \beta \frac{r_{t-12,t}^i}{12} + \varepsilon_{t+1}^i,$$

where β equals 0.1, 0.2, and 0.4. For each asset, we assume it has the same mean and variance as that in Table 1. Then, given the common slope β , α^i is estimated with asset i 's real returns. We simulate a path of $T = 372$ observations, construct the TSM and TSH strategies, and calculate their mean returns. We repeat this procedure one thousand times to test whether these two strategies generate the same mean returns.

find that the two specifications generate almost the same results and therefore focus on the first specification.

Fig. 5 reports the results. When the slope is 0.1, the two strategies perform almost the same. When the slope is 0.2, the TSM outperforms the TSH, but the difference is not significant. Thus, if there is genuine time series predictability, the advantage of the TSM strategy is not apparent as long as the slope is small. When the slope is 0.4, the TSM dominates the TSH in the sense that it does better in almost all the simulated data sets. Because the two strategies generate similar performance using the real data, our simulation indicates that the evidence of time series predictability is weak if it exists. Combined with other results, the TSM is unlikely to be statistically significant for all the assets. In short, a lack of empirical evidence exists to support that the TSM is everywhere.

Because the TSM and TSH use overlapping data at the beginning of the investment period, one could expect that this explains the statistically indistinguishable difference even when the slope is 0.2. In fact, it is not the case. To see this, we simulate a path of $T + 240$ observations for the 55 assets, construct the TSM and TSH strategies starting from the 241th observation (i.e., the TSM is based on the past 12-month return and the TSH is based on the historical sample mean), and calculate their mean returns. We repeat this procedure one thousand times to test whether these two strategies generate the same mean returns. Fig. A1 of the Online Appendix summarizes the results. Generally, the basic patterns are similar to Fig. 5; that is, the TSM strategy performs similarly as the TSH when the data exhibit weak or intermediate time series predictability, and it outperforms the TSH when the data

exhibit strong time series predictability. One observation is that with longer historical data in calculating the historical mean, the TSM generates 1% more annualized return relative to the case of Fig. 5, which is simply due to the fact that observations over a longer time horizon generate a better estimate of the sample mean (Merton, 1980).

6. Conclusion

In their influential study, MOP (Moskowitz et al., 2012) assert a surprising time series momentum that the past 12-month return can positively predict the next one-month to 12-month return everywhere. They also show that a trading strategy based on TSM generates significant average and risk-adjusted returns. As TSM is in stark contrast to the previous literature and strongly challenges the weak form market efficiency hypothesis, we revisit TSM in this paper. Employing the same data as MOP but extending the sample period to 2015, we show that, statistically, the evidence for TSM is weak in asset-by-asset time series regressions and a pooled regression accounting for size distortions. Economically, we show that the performance of the TSM strategy is likely driven by differences in mean returns, not predictability. A predictive slope analysis following the approach of Lewellen (2015) further confirms weak evidence on TSM.

A number of topics are of interest for future research. First, while the TSM strategy focuses on the 12-month return predictability, examining such a strategy at other time horizons and an optimal combination of all would be worthwhile. Second, the predictability horizon can be time-varying and could be different across assets (instead of the same horizon), and developing a test and a new trading strategy for this possibility would be important. Finally, considering the conditions under which time series predictability can exist in an equilibrium model and developing an empirical test for the implications would be desirable.

References

- Ang, A., Bekaert, G., 2007. Stock return predictability: is it there? *Rev. Financ. Stud.* 20, 651–707.
- Asness, C., Moskowitz, T., Pedersen, L., 2013. Value and momentum everywhere. *Rev. Financ. Stud.* 20, 651–707.
- Boudoukh, J., Richardson, M., Whitelaw, R., 2008. The myth of long-horizon predictability. *Rev. Financ. Stud.* 21, 1577–1605.
- Campbell, J., Thompson, S., 2008. Predicting excess stock returns out of sample: can anything beat the historical average? *Rev. Financ. Stud.* 21, 1509–1531.
- Campbell, J., Yogo, M., 2006. Efficient tests of stock return predictability. *J. Financ. Econ.* 81, 27–60.
- Clark, T., West, K., 2007. Approximately normal tests for equal predictive accuracy in nested models. *J. Econom.* 138, 291–311.
- Conrad, J., Kaul, G., 1988. An anatomy of trading strategies. *Rev. Financ. Stud.* 11, 489–519.
- Davidson, R., Flachaire, E., 2008. The wild bootstrap, tamed at last. *J. Econom.* 146, 162–169.
- Fama, E., 1965. The behavior of stock market prices. *J. Bus.* 38, 34–105.
- Fisher, R.A., 1918. The correlation between relatives on the supposition of Mendelian inheritance. *Philos. Trans. R. Soc. Edinb.* 52, 399–433.
- French, K., Roll, R., 1986. Stock return variances: the arrival of information and the reaction of traders. *J. Financ. Econ.* 17, 5–26.
- Freyberger, J., Neuhierl, A., Weber, M., 2019. Dissecting characteristics nonparametrically. *Rev. Financ. Stud.*, forthcoming.
- Georgopoulou, A., Wang, J., 2017. The trend is your friend: time series momentum strategies across equity and commodity markets. *Rev. Finance* 21, 1557–1592.
- Gonçalves, S., Kaffo, M., 2015. Bootstrap inference for linear dynamic panel data models with individual fixed effects. *J. Econom.* 186, 407–426.
- Goyal, A., Jegadeesh, N., 2018. Cross-sectional and time series tests of return predictability: what is the difference? *Rev. Financ. Stud.* 31, 1784–1824.
- Gu, S., Kelly, B. T., Xiu, D., 2018. Empirical asset pricing via machine learning. Unpublished working paper, University of Chicago.
- Han, Y., He, A., Rapach, D., Zhou, G., 2019. What firm characteristics drive US stock returns? Unpublished working paper, Washington University in St. Louis.
- Hansen, P., Timmermann, A., 2012. Choice of sample split in out-of-sample forecast evaluation. Unpublished working paper, University of California, San Diego.
- Harvey, C.R., Liu, Y., Zhu, H., 2016. ... and the cross section of expected returns. *Rev. Financ. Stud.* 29 (1), 5–68.
- Hjalmarsson, E., 2010. Predicting global stock returns. *J. Financ. Quant. Anal.* 45, 49–80.
- Hodrick, J., 1992. Dividend yields and expected stock returns: alternative procedures for inference and measurement. *Rev. Financ. Stud.* 5, 357–386.
- Hou, K., Xue, C., Zhang, L., 2019. Replicating anomalies. *Rev. Financ. Stud.*, forthcoming.
- Hurst, B., Ooi, Y., Pedersen, L., 2017. A century of evidence on trend-following investing. *J. Portf. Manag.* 44, 15–29.
- Inoue, A., Kilian, L., 2005. In-sample or out-of-sample tests of predictability: which one should we use? *Econom. Rev.* 23 (4), 371–402.
- Jegadeesh, N., 1990. Evidence of predictable behavior of security returns. *J. Finance* 45 (3), 881–898.
- Jegadeesh, N., Titman, S., 1993. Returns to buying winners and selling losers: implications for stock market efficiency. *J. Finance* 48 (1), 65–91.
- Jiang, F., Lee, J., Martin, X., Zhou, G., 2019. Manager sentiment and stock returns. *J. Financ. Econ.* 132 (1), 126–149.
- Jorion, P., Goetzmann, W., 1999. Global stock markets in the 20th century. *J. Finance* 54, 953–980.
- Kim, A., Tse, Y., Wald, J., 2016. Time series momentum and volatility scaling. *J. Financ. Mark.* 30, 102–124.
- Kojien, R., Moskowitz, T., Pedersen, L., Vrugt, E., 2018. Carry. *J. Financ. Econ.* 127, 197–225.
- Kruskal, W.H., Wallis, W.A., 1952. Use of ranks in one-criterion variance analysis. *J. Am. Stat. Assoc.* 47, 583–621.
- Lewellen, J., 2015. The cross section of expected stock returns. *Crit. Finance Rev.* 4, 1–44.
- Li, J., Yu, J., 2012. Investor attention, psychological anchors, and stock return predictability. *J. Financ. Econ.* 104, 401–419.
- Lo, A., MacKinlay, A., 1988. Stock market prices do not follow random walks: evidence from a simple specification test. *Rev. Financ. Stud.* 1, 41–66.
- Menzly, L., Santos, T., Veronesi, P., 2004. Understanding predictability. *J. Polit. Econ.* 112, 1–47.
- Merton, R.C., 1980. On estimating the expected return on the market: an exploratory investigation. *J. Financ. Econ.* 8 (4), 323–361.
- Moskowitz, T., Ooi, Y., Pedersen, L., 2012. Time series momentum. *J. Finance* 104, 228–250.
- Paye, B.S., 2012. 'Déjà vol': predictive regressions for aggregate stock market volatility using macroeconomic variables. *J. Financ. Econ.* 106 (3), 527–546.
- Rapach, D., Strauss, J., Zhou, G., 2013. International stock return predictability: what is the role of the United States? *J. Finance* 68, 1633–1662.
- Stambaugh, R.F., 1999. Predictive regressions. *J. Financ. Econ.* 54 (3), 375–421.
- Valkanov, R., 2003. Long-horizon regressions: theoretical results and applications. *J. Financ. Econ.* 68, 201–232.
- Welch, B.L., 1951. On the comparison of several mean values: an alternative approach. *Biometrika* 38, 330–336.
- Welch, I., Goyal, A., 2008. A comprehensive look at the empirical performance of equity premium performance. *Rev. Financ. Stud.* 21, 1455–1508.