

ESSAYS ON HYPOTHESIS TESTING AND FORECASTING WITH
HIGH-FREQUENCY FINANCIAL DATA

by

MINGMIAN CHENG

A dissertation submitted to the

School of Graduate Studies

Rutgers, The State University of New Jersey

In partial fulfillment of the requirements

For the degree of

Doctor of Philosophy

Graduate Program in Economics

Written under the direction of

Norman Rasmus Swanson

And approved by

New Brunswick, New Jersey

May, 2018

ABSTRACT OF THE DISSERTATION

Essays on Hypothesis Testing and Forecasting with High-frequency Financial Data

By MINGMIAN CHENG

Dissertation Director:

Norman Rasmus Swanson

This dissertation studies methodologies for hypothesis testing and forecasting in financial econometrics, and comprises two essays on these topics, respectively. The first essay mainly aims to shed light on the importance of the proper application of jump testing methods, while the second essay provides alternative methodological approaches to volatility forecasting.

In Chapter 2, we examine and compare a variety of jump tests in the financial econometrics literature. Numerous tests designed to detect realized jumps over a fixed time span have been proposed and extensively studied in recent years. These tests differ from “long time span” jump tests that detect jumps by examining the magnitude of the intensity parameter in the data generating process of an asset. In this chapter, a long time span jump test thereof called the *CSS* test, which is a variant of Corradi et al. (2018), is compared with a variety of fixed time span jump tests, in a series of Monte Carlo experiments. The *CSS* test is consistent against the null hypothesis of zero jump intensity, while the fixed time span tests are not designed to detect jumps in the data generating process, and instead detect realized jumps over a fixed time span. An empirical investigation of individual, sector specific and market level stock prices is

also carried out in order to contrast findings based on these different varieties of tests. The fixed time span tests examined include the higher order power variation *ASJ* test of Aït-Sahalia and Jacod (2009), the classic bipower variation *BNS* test of Barndorff-Nielsen and Shephard (2006), and the truncated power variation *PZ* test of Podolskij and Ziggel (2010). It is found that both the *ASJ* test and the *CSS* test exhibit good finite sample properties for time spans both short and long. The other tests suffer from finite sample distortions under long time spans. When applied to stock price and stock index data, the two aforementioned tests indicate that the prevalence of jumps is not as universal as might be expected. Various sector ETFs and individual stocks, for example, appear to exhibit no jumping behavior during a number of annual periods.

In Chapter 3, we use factor-augmented heterogeneous autoregressive (HAR)-type models to predict the daily integrated volatility of asset returns. Our approach is based on a proposed two-step dimension reduction procedure designed to extract latent common volatility factors from a large dimensional and high-frequency returns dataset with 267 constituents of the S&P 500 index. In the first step, we apply either least absolute shrinkage and selection operator (LASSO) or elastic net shrinkage on estimates of integrated volatility of all constituents in the dataset, in order to select a subset of asset return series for further processing. In the second step, we utilize (sparse) principal component analysis to estimate latent common asset return factors, from which latent integrated volatility factors are extracted. Although we find limited *in-sample* fit improvement, relative to a benchmark HAR model, all of our proposed factor-augmented models result in substantial *out-of-sample* predictive accuracy improvement. In particular, forecasting gains are observed at market, sector, and individual-stock levels, with the exception of the financial sector. Further investigation of the factor structures for non-financial assets shows that industrial and technology stocks are characterized by minimal exposure to financial assets, inasmuch as forecasting gains associated with factor-augmented models for these types of assets are largely attributable to the inclusion of non-financial stock price return volatility in our latent factors.

Acknowledgements

I am deeply indebted to my dissertation advisor, Professor Norman R. Swanson, not only for his invaluable guidance on my reserach, but also for his vigorous support during my entire doctoral studies at Rutgers University. Most importantly, he guided me through the transition from a student to a researcher, and led me to find the research fields in which I am motivated to spend my whole academic career. He is more than my advisor. He set a role model for me.

I am also extremely grateful to Professor Yuan Liao and Professor Xiye Yang for their generous advice and encouragement since the first day they joined the Department of Economics at Rutgers University. They both have very strong technical skills in statistics and econometrics, and are very patient and willing to instruct me on my research. Every time I had a meeting with them, I benefited significantly. Furthermore, they kept strengthening my courage and will to stay in academia. They are not only my advisors but also my good friends.

I would like to sincerely thank Professor John Landon-Lane, Professor Roger Klein, Professor Roberto Chang and Professor Zhifeng Cai for their instructive comments, suggestions and tips during my practice job talks and mock interviews. They helped me to find a better way of marketing myself on the job market.

As a job market candidate, I owe lots of thanks to Professor Hilary Sigman, Professor Oriol Carbonell-Nicolau and Linda Zullinger for their help, advice, distributing information and organizing meetings and mock interviews over the past year.

I would like to thank the administrative staff, Donna Ghilino, Janet Budge, Debra Holman, Linda Zullinger and Paula Seltzer, in our department. Their support year after year is also a key ingredient for my completion of the doctoral degree.

Special thanks are owed to a special group (nickname: “Pot Love”) of people, including Yixiao Jiang, Zhutong Gu, Shuyang Yang, Han Liu, Wukuang Cun, Xi Wang, Chen Yang, Zixin Lin, Siyu Wang, Long Feng, Chuan Liu and Yuanting Wu. They are my best friends that I made over the past six years. Without them, my everyday life

at Rutgers would be tedious, and I would not have the courage to finish my doctoral studies. I really appreciate the support from all my friends and remember the days we spent together.

Finally, I want to say many thanks to my parents, for their full financial support, constant encouragement and selfless love.

Dedication

To my parents and my grandmother.

Table of Contents

Abstract	ii
Acknowledgements	iv
Dedication	vi
List of Tables	ix
List of Figures	xi
1. Introduction	1
2. A Comparison of Fixed and Long Time Span Jump Tests	5
2.1. Introduction	5
2.2. Setup	9
2.3. Heuristic Discussion	11
2.4. Long Time Span Jump Intensity Test	14
2.5. Fixed Time Span Realized Jump Tests	17
2.5.1. Aït-Sahalia and Jacod (<i>ASJ</i> : 2009) Test	17
2.5.2. Barndorff-Nielsen and Shephard (<i>BNS</i> : 2006) Test	19
2.5.3. Podolskij and Ziggel (<i>PZ</i> : 2010) Test	19
2.6. Monte Carlo Simulations	20
2.7. Empirical Examination of Stock Market Data	24
2.7.1. Data	24
2.7.2. Empirical Findings	24
2.8. Concluding Remarks	27
3. Latent Common Volatility Factors: Capturing Elusive Predictive Accuracy Gains When Forecasting Volatility	43
3.1. Introduction	43
3.2. Setup	49

3.3. Dimension Reduction and Forecasting Methods	51
3.3.1. LASSO and Elastic Net	53
3.3.2. Principal Component Analysis	54
3.3.3. Sparse Principal Component Analysis	55
3.3.4. Forecasting Methods	56
3.4. Empirical Results	59
3.4.1. Data	59
3.4.2. Empirical Findings: Forecasting Performance	60
3.4.3. Empirical Findings: Latent Factor Structures	65
3.5. Concluding Remarks	66
Bibliography	83

List of Tables

2.1. Data Generating Processes in Monte Carlo Experiments	28
2.2. Empirical Size of Daily Fixed Time Span Jump Tests and Sequential Testing	29
2.3. Empirical Power of Daily Fixed Time Span Jump Tests	30
2.4. Empirical Size of Fixed Time Span Jump Tests When Utilized Using Long-span Samples	31
2.5. Empirical Power of the <i>ASJ</i> Jump Test When Utilized Using Long-span Samples	32
2.6. Empirical Power of the <i>BNS</i> Jump Test When Utilized Using Long-span Samples	33
2.7. Empirical Power of the <i>PZ</i> Jump Test When Utilized Using Long-span Samples	34
2.8. Empirical Size of $LTS_{T,\Delta}$ and $\widetilde{LTS}_{T,\Delta}$ Jump Tests	35
2.9. Empirical Power of $LTS_{T,\Delta}$ Jump Test	36
2.10. Empirical Power of $\widetilde{LTS}_{T,\Delta}$ Jump Test	37
2.11. <i>ASJ</i> Jump Test Results for ETFs	38
2.12. $LTS_{T,\Delta}$ Jump Test Results for ETFs	38
2.13. $\widetilde{LTS}_{T,\Delta}$ Jump Test Results for ETFs	39
2.14. <i>ASJ</i> Jump Test Results for Individual Stocks	39
2.15. $LTS_{T,\Delta}$ Jump Test Results for Individual Stocks	40
2.16. $\widetilde{LTS}_{T,\Delta}$ Jump Test Results for Individual Stocks	40
3.1. Experimental Setup	68
3.2. SPDR S&P 500 ETF (SPY)	69
3.3. Materials Sector ETF (XLB)	70
3.4. Energy Sector ETF (XLE)	70
3.5. Financial Sector ETF (XLF)	71

3.6. Industrial Sector ETF (XLI)	71
3.7. Technology Sector ETF (XLK)	72
3.8. Consumer Staples Sector ETF (XLP)	72
3.9. Utilities Sector ETF (XLU)	73
3.10. Health Care Sector ETF (XLV)	73
3.11. Consumer Discretionary Sector ETF (XLY)	74
3.12. General Electric Company (GE)	74
3.13. The Goldman Sachs Group, Inc. (GS)	75
3.14. International Business Machines Corporation (IBM)	75
3.15. Johnson & Johnson (JNJ)	76
3.16. JPMorgan Chase & Co. (JPM)	76
3.17. The Coca-Cola Company (KO)	77
3.18. McDonald's Corporation (MCD)	77
3.19. Merck & Co., Inc. (MRK)	78
3.20. Microsoft Corporation (MSFT)	78
3.21. Wal-Mart Stores, Inc. (WMT)	79
3.22. Exxon Mobil Corporation (XOM)	79
3.23. Factor Structure (SPY)	80
3.24. Factor Structure (XLE)	81
3.25. Factor Structure (IBM)	82

List of Figures

2.1. Annual Ratios of Jump Days for ETFs	41
2.2. Annual Ratios of Jump Days for Individual Stocks	42

Chapter 1

Introduction

Continuous-time financial econometrics and its application have received broad attention of both economists and practitioners in recent years. Asset prices (or returns) and exchange rates are frequently modeled as continuous-time processes, such as (Itô-) semimartingales (see e.g. Andersen et al. (2001), Chernov et al. (2003) and Aït-Sahalia and Jacod (2014)). A large variety of “in-fill” asymptotic theorems as the data sampling interval $\Delta \rightarrow 0$ have been developed and studied as well (see e.g. Jacod and Protter (2011)). In this dissertation, hypothesis tests of continuous-time modeling and methods for volatility forecasting are considered. They are of particular interests for the following reasons. On one hand, a clear understanding and a correct specification of the data generating process that governs asset price movements are crucial for consistent model estimation and statistical inference. On the other hand, accurate prediction of asset price volatility is a key ingredient for successful risk management and asset allocation.

In the second chapter entitled “*A Comparison of Fixed and Long Time Span Jump Tests*”, we compare a variety of “fixed time span” jump tests, including the higher order power variation *ASJ* test of Aït-Sahalia and Jacod (2009), the classic bipower variation *BNS* test of Barndorff-Nielsen and Shephard (2006), and the truncated power variation *PZ* test of Podolskij and Ziggel (2010), with a “long time span” jump test called the *CSS* test, which is a variant of Corradi et al. (2018), in a series of Monte Carlo experiments as well as empirical studies. Fixed time span jump tests are designed to detect realized jumps over a fixed time span, such as a day or a week, while the long time span jump tests detect jumps by examining the magnitude of the intensity parameter in the data generating process of an asset. As a result, the long time span *CSS* test is consistent against non-zero jump intensity, while the fixed time span jump tests are inconsistent. Heuristically, if researchers detect jumps in a particular sample path, they might conclude that the jump intensity is non-zero. However, if no jumps are found in a sample path, this does not mean there are no jumps in other sample paths, and

thus they cannot conclude that the jump intensity parameter in the data generating process is identically zero. Furthermore, fixed time span jump tests are sensitive to sequential testing bias since the probability of rejecting the null of zero jump intensity approaches unity when sequentially applying a growing numbers of fixed- T tests, even if the true data generating process is purely continuous. In addition, we also examine the performance of fixed- T tests on a long-span dataset, which is rarely studied in the literature.

Our findings from Monte Carlo experiments can be summarized as follows. First, we show that the finite sample power of daily jump tests against non-zero jump intensity is low. For instance, when the jump intensity is 0.4 and the jump size parameter is our largest, rejection rates of the *ASJ*, *BNS* and *PZ* tests at a 0.05 significance level are only around 0.26, 0.38, and 0.36, respectively. Second, sequential testing bias grows rapidly when we apply a growing sequence of fixed- T jump tests. Even for the most conservative test (i.e., the *ASJ* test), empirical size is over 0.95 after 50 days. Third, we show that the empirical sizes of fixed- T jump tests over samples with growing time spans also increase. But the size distortion accumulates much more slowly when using the *ASJ* test than when using the *BNS* and *PZ* tests. Moreover, the power of *ASJ* test is very good in general. When the sample is over 50 days, the *ASJ* test is powerful even for infrequent and weak jumps. Fourth, our *CSS* type test which is not robust to leverage has good size properties if there is no leverage in the DGP, while empirical size increases in T when the DGP is characterized by the presence of leverage, as expected. On the other hand, our “leverage-robust” *CSS* type test is conservatively sized as it has zero asymptotic size. Also, the power of the “leverage robust test” is found to be good when a simple rule-of-thumb is used to specify coarser Δ , say $\tilde{\Delta}$, as T is grows, when constructing bootstrap critical values.

In our empirical analysis, we examine 5-minute intraday observations between 2006 and 2013 on twelve individual stocks, nine sector ETFs, and the market ETF. Our main empirical findings are summarized as follows. First, using daily *ASJ*, *BNS* and *PZ* tests, jumps are widely detected in asset prices over almost all time periods considered. For instance, all three tests detect jumps on around 35%-40% of the days in 2006 for two

of the ETFs that we examine. Second, these jump percentages have diminished over time. Third, long-span jump tests tell a different story. Namely, the *ASJ* test, as well as our leverage robust *CSS* type test indicate far fewer jumps than are found when using daily tests.

In the third chapter entitled “*Latent Common Volatility Factors: Capturing Elusive Predictive Accuracy Gains When Forecasting Volatility*”, we use factor-augmented heterogeneous autoregressive (HAR)-type models to predict the daily integrated volatility of asset returns. Our approach is based on a proposed two-step dimension reduction procedure designed to extract latent common volatility factors from a large dimensional and high-frequency returns dataset in extensive forecasting experiments. In the first step, we apply either least absolute shrinkage and selection operator (LASSO) or elastic net shrinkage on estimates of integrated volatility of all constituents in the dataset, in order to select a subset of asset return series for further processing. In the second step, we utilize (sparse) principal component analysis to estimate latent common asset return factors, from which latent integrated volatility factors are extracted.

Our dataset consists of intra-day observations on 267 constituents of the S&P 500 index, 9 sector ETFs, and one market EFT (SPY). Data were analyzed for the sample period from January 3, 2006 to December 31, 2010. We report the results based on prediction of SPY, 9 sector ETFs, and 11 individual stocks, for the period of July 1, 2009 to December 31, 2010. Our key findings are summarized below. First, *in-sample* fit is surprisingly stable across different models, with most R^2 values ranging rather tightly between 0.35 and 0.55. Second, for ex ante prediction, data frequency is crucial. Our factor augmented HAR models generally yield the “best” predictions using 5-minute frequency data. So we recommend using the 5-minute frequency, as a general rule-of-thumb. Third, and perhaps most importantly, predictive accuracy improves appreciably when latent common volatility factors are included in benchmark HAR-type models. For example, for Johnson & Johnson, the benchmark model using 5-minute frequency data achieves an *out-of-sample* R^2 value of only 0.14. This is approximately one-third of the *out-of-sample* R^2 value associated with our “best” factor-augmented model. This pattern occurs for many firms and sectors, as well as for the market ETF. Fourth, there

is an important wrinkle to the above story. Namely, for financial assets, out-of sample R^2 values are approximately 0 in some cases. Finally, financial stocks are frequently selected in our first variable selection step. However, they are often assigned small weights in the second step. For instance, when we forecast the volatility of our energy sector ETF using 1-minute frequency data, over 33% of the most frequently selected stocks in the first step are in financial sector. However, the average weight assigned by PCA to Goldman Sachs is only around 0.09, while the corresponding weight assigned to Texas Instruments is around double that. As a result, it is very likely that the marginal predictive content of common volatility factors is largely accounted for by information in sectors other than the financial sector.

Chapter 2

A Comparison of Fixed and Long Time Span Jump Tests

2.1 Introduction

In risk management and financial engineering, investors and researchers often require knowledge of the data generating process (DGP) that governs asset price movements. For example, asset prices are frequently modeled as continuous-time processes, such as (Itô-) semimartingales (see, e.g. Aït-Sahalia (2002a,b), Chernov et al. (2003), and Andersen et al. (2007b)). At the same time, investors and researchers are also interested in nonparametrically estimable quantities such as spot/integrated volatility (see, e.g. Barndorff-Nielsen (2002), Barndorff-Nielsen and Shephard (2003), Todorov and Tauchen (2011), and Patton and Sheppard (2015)), jump variation (see, e.g. Barndorff-Nielsen and Shephard (2004), Andersen et al. (2007a), and Corsi et al. (2009)), leverage effects (see, e.g. Kalnina and Xiu (2017) and Aït-Sahalia et al. (2016)), and jump activity (see e.g. Aït-Sahalia and Jacod (2011) and Todorov (2015)). In this paper, we add to the jump testing literature by carrying out an extensive Monte Carlo and empirical analysis of jump detection using so-called “fixed time span” jump tests (see, e.g. Barndorff-Nielsen and Shephard (2006), Lee and Mykland (2008), Aït-Sahalia and Jacod (2009), Corsi et al. (2010), and Podolskij and Ziggel (2010)) and “long time span” jump tests (see e.g. Corradi et al. (2014, 2018) (*CSS*)). The reason why a “horse-race” between alternative jump tests of these varieties is of interest is because it is well known that tests constructed using observed sample paths of asset returns on a “fixed time span”, such as a day or a week, are not consistent, and are sensitive to sequential testing bias. On the other hand, the *CSS* type test that we examine, which is based on direct evaluation of the data generating process, is consistent and asymptotically correctly sized when the time span, $T \rightarrow \infty$, and the sampling interval, $\Delta \rightarrow 0$.

One reason why detecting jumps using long time span tests is of potential interest is that empirical researchers often estimate DGPs after testing for jumps using fixed time

span tests. However, when estimating jump diffusions, drift, volatility, jump intensity and jump size parameters are usually jointly estimated. This is problematic if the jump intensity is identically zero, since parameters characterizing jump size are unidentified, and consistent estimation of the rest of the parameters is thus no longer feasible (see Andrews and Cheng (2012)). As a result, testing for jumps via pretesting for zero jump intensity is a natural alternative to the use of nonparametric fixed time span jump tests. In addition to the issue of identification, if researchers detect jumps in a particular sample path, they might conclude that the jump intensity is non-zero. However, if no jumps are found in a sample path, this does not mean there are no jumps in other sample paths, and hence that a DGP should be estimated with no jump component.

We consider two *CSS* type jump tests. Both of these build on earlier work of Aït-Sahalia (2002a,b), and are special cases of the leverage robust test introduced in Corradi et al. (2018).¹ One test assumes no leverage. The other test is robust to leverage, and is a rescaled version of the no leverage test. Both tests are derived under the assumption that $E\left((Y_{k\Delta} - Y_{(k-1)\Delta})^3\right) = 0$, whenever there are no jumps, where $Y_{k\Delta} = \ln X_{k\Delta} - \frac{\Delta}{T} \sum_{k=2}^n \ln X_{k\Delta}$ and $Y_{(k-1)\Delta} = \ln X_{(k-1)\Delta} - \frac{\Delta}{T} \sum_{k=2}^n \ln X_{(k-1)\Delta}$, and where the X are asset prices.

In the Monte Carlo and empirical analyses discussed in the sequel, the finite sample properties of three representative fixed time span tests as well as the *CSS* type tests are investigated. The three “fixed- T ” tests include the higher order power variation test of Aït-Sahalia and Jacod (2009) (*ASJ*), the classic bipower variation test of Barndorff-Nielsen and Shephard (2006) (*BNS*), and the truncated power variation test of Podolskij and Ziggel (2010) (*PZ*). These tests are chosen to be representative of three broader classes of fixed- T tests that utilize multipower variation, higher order power variation, and truncation. For a detailed comparison of more fixed- T tests, refer to Theodosiou and Zikes (2009) and Dumitru and Urga (2012). These authors concisely summarize and compare a large group of existing jump tests via extensive Monte Carlo experiments.

¹The tests are variants of the test discussed in Corradi et al. (2018) that originally appeared in Corradi et al. (2014).

However, it is worth noting that very few researchers examine the performance of fixed- T tests on a long-span dataset. A key exception is Huang and Tauchen (2005), who examine a “long time span” *BNS* type test. In this interesting paper, the authors find that the empirical size of the *BNS* test deviates from the nominal size more significantly as the time span increases; and that size distortion becomes even more substantial if the sample path is more volatile but still continuous. As a consequence, a null of zero jump intensity is more likely to be overrejected and it is possible to falsely identify a jump diffusion process when it is purely continuous. They suggest that an appropriate way to solve both inconsistency and size distortion problems involves using test statistics that are asymptotically valid under a double asymptotic scheme where both $T \rightarrow \infty$ and $\Delta \rightarrow 0$.

Our findings can be summarized as follows. First, we show that the finite sample power of daily jump tests against non-zero jump intensity is low, particularly when jumps are infrequent or jump magnitudes are “weak”. For instance, when the jump intensity is 0.4 and the jump size parameter is our largest, rejection rates of the *ASJ*, *BNS* and *PZ* tests at a 0.05 significance level are only around 0.26, 0.38, and 0.36, respectively. Second, sequential testing bias grows rapidly as the time span increases. The size of a joint test based on the strategy of sequentially performing many fixed- T daily tests approaches unity very quickly. Even for the most conservative test (i.e., the *ASJ* test), empirical size is over 0.95 after 50 days. Importantly, we also show that the empirical sizes of fixed- T jump tests over samples with growing time spans also increase in T . Specifically, the size of the *PZ* test over a sample of 300 days is close to one. The empirical size of the *BNS* test is twice as large as the nominal size, when the sample is over 300 days. For the *ASJ* test, empirical size also increases as the time span increases. However, as long as the sample is not too long, say more than 150 days, the *ASJ* test is surprisingly well sized. More generally, size distortion accumulates much more slowly when using the *ASJ* test than when using the *BNS* and *PZ* tests. Moreover, the power of *ASJ* test is very good for all long time spans, as long as jumps are not too rare and too weak. When the sample is over 50 days, the *ASJ* test is powerful even for infrequent and weak jumps. Fourth, our *CSS* type test which is not robust to leverage has good

size properties, even when the sample is very long, such as 500 days, if there is no leverage in the DGP, while empirical size increases in T when the DGP is characterized by the presence of leverage, as expected. On the other hand, our “leverage-robust” *CSS* type test is conservatively sized. This is not surprising, since the test has zero asymptotic size.² Also, the power of the “leverage robust test”, while not as good as that of the “non leverage-robust” test, is found to be good when a simple rule-of-thumb is used to specify coarser Δ , say $\tilde{\Delta}$, as T is grows, when constructing bootstrap critical values, in order to mitigate the effect on finite sample power of the use of adjustment term accounting for leverage.

In our empirical analysis, we examine 5-minute intraday observations between 2006 and 2013 on twelve individual stocks, nine sector ETFs, and the market (SPDR S&P500) ETF. Our main empirical findings are summarized as follows. First, using daily *ASJ*, *BNS* and *PZ* tests, jumps are widely detected in asset prices over almost all time periods considered. Moreover, in some cases, the annual percentage of jump days seems inconceivably large. For instance, all three tests detect jumps on around 35%-40% of the days in 2006 for two of the ETFs that we examine. Second, these jump percentages have diminished over time. Third, long-span jump tests tell a different story. Namely, the *ASJ* test, as well as our leverage robust *CSS* type test indicate far fewer jumps than are found when using daily tests. This finding has important implications for both specification and estimation of asset price models.

The rest of the paper is organized as follows. Section 2.2 outlines the theoretical framework and introduces notation. Section 2.3 discusses statistical issues associated with testing for jumps. Section 2.4 discusses the long time span jump tests that we examine, and Section 2.5 briefly discusses the extant fixed time span tests examined in the sequel. Section 2.6 reports results from our Monte Carlo experiments. Section 2.7 presents the results of our empirical analysis of various stock price and price index data. Finally, Section 2.8 contains concluding remarks.

²The long span test in Corradi et al. (2018) is robust to leverage and is correctly asymptotically sized. They achieve this by introducing a threshold variance estimator with which to scale their test, rather than relying on the bootstrap, as is done in the variant of their test examined in this paper.

2.2 Setup

We use the setup of Corradi et al. (2018). Namely, assume that (log-)asset prices are recorded at an equally spaced discrete interval, $\Delta = \frac{1}{m}$, where m is the total number of observations on each trading day. In our model, we assume that $\Delta \rightarrow 0$; or equivalently that $m \rightarrow \infty$. Log-prices follow a jump diffusion model defined on some filtered probability space $(\Omega, \mathbb{F}, (\mathcal{F}_t)_{t \geq 0}, \mathbb{P})$, with

$$d \ln X_t = \mu dt + \sqrt{V_t} dW_{1,t} + Z_t dN_t, \quad (2.1)$$

where μ is the drift term, $\sqrt{V_t}$ is the spot volatility, and $W_{1,t}$ is an adapted standard Brownian motion (i.e., it is $\mathcal{F}_t$ -measurable for each $t \geq 0$). Here, V_t is defined according to either (i), (ii), (iii), or (iv), as follows:

(i) a constant:

$$V_t = v \text{ for all } t; \quad (2.2)$$

(ii) a measurable function of the state variable:

$$V_t \text{ is } X_t\text{-measurable}; \quad (2.3)$$

(iii) a stochastic volatility process without leverage:

$$dV_t = \mu_{V,t}(\theta)dt + g(V_t, \theta) dW_{2,t}, \quad \mathbb{E}(W_{1,t}W_{2,t}) = 0; \quad (2.4)$$

(iv) a stochastic volatility process with leverage:

$$dV_t = \mu_{V,t}(\theta)dt + g(V_t, \theta) dW_{2,t}, \quad \mathbb{E}(W_{1,t}W_{2,t}) = \rho \neq 0. \quad (2.5)$$

Evidently, the volatility process is quite general, although we do not consider jumps in volatility.

Now, define,

$$\Pr(N_{t+\Delta} - N_t = 1 | \mathcal{F}_t) = \lambda_t \Delta + o(\Delta), \quad (2.6)$$

$$\Pr(N_{t+\Delta} - N_t = 0 | \mathcal{F}_t) = 1 - \lambda_t \Delta + o(\Delta), \quad (2.7)$$

and

$$\Pr(N_{t+\Delta} - N_t > 1 | \mathcal{F}_t) = o(\Delta), \quad (2.8)$$

where λ_t characterizes the jump intensity. The jump size, Z_t , is identically and independently distributed with density $f(z; \gamma)$. Equation (2.8) implies that we rule out infinite-activity jumps.

When constructing the fixed time span realized jump tests discussed in the sequel, we remain agnostic about the jump generating process. However, for the case of our long time span jump intensity tests, we must provide a moderate amount of additional structure. This is one of the key trade-offs associated with using either variant of test. In particular, following Corradi et al. (2018), we consider two cases. First, N_t , which is the number of jump arrivals up to t , follows a counting process, such as the widely used Poisson process. In this case, $\lambda_t = \lambda$, for all t . Second, jumps may be “self-exciting”, in the sense that the jump intensity follows Hawkes diffusion (see Bowsher (2007) and Aït-Sahalia et al. (2015)), with

$$d\lambda_t = a(\lambda_\infty - \lambda_t)dt + \beta dN_s, \quad (2.9)$$

where $\lambda_\infty \geq 0$, $\beta \geq 0$, $a > 0$, and $a > \beta$, so that the process is mean reverting with $E(\lambda_t) = \frac{a\lambda_\infty}{a-\beta}$. As noted in Corradi et al. (2018), if $\lambda_\infty = 0$, then $E(\lambda_t) = 0$, which implies that $\lambda_t = 0$, a.s. for all t . This implies that $N_t = 0$, a.s., for all t . As a result, β , a and γ in this case are all unidentified. On the other hand, if $\lambda_\infty > 0$, then β and γ are identified. But if $\beta = 0$, a is not identified. These observations highlight the importance of pretesting for $\lambda_\infty = 0$ against $\lambda_\infty > 0$, in order to obtain consistent estimation of parameters in the case of Hawkes diffusions.

In light of the above discussion, we are interested in testing

$$H_0 : \lambda = 0$$

versus

$$H_A : \lambda > 0,$$

where λ is the constant jump intensity, in the case of Poisson-type jumps; and is the expectation of the stochastic jump intensity (i.e. $\lambda = E(\lambda_t)$), in the case of self-exciting jumps.³ This is a nonstandard inference problem because, under H_0 , some parameters

³Note that $\lambda_\infty = (>) 0$ if and only if $E(\lambda_t) = (>) 0$.

are not identified, and a parameter lies on the boundary of the null parameter space.

Before discussing the tests that are compared in our Monte Carlo experiments, we first provide some heuristic motivation for long time span jump testing. This discussion follows Corradi et al. (2014), who also provide complete details on the *CSS* type tests that we consider in the sequel.

2.3 Heuristic Discussion

In recent years, a large variety of tests for realized jumps have been proposed and studied. One common feature of the preponderance of these tests is that they are all carried out using high-frequency observations over a fixed time span and justified by in-fill asymptotic theorems. Therefore, they have power against realized jumps over fixed time spans; and none are consistent against the alternative of $\lambda > 0$. Many of the tests can be considered as Hausman-type tests in which a comparison of two realized measures of the integrated volatility is made, where one is robust to jumps, and one is not. Under the null of no jumps, both consistently estimate the integrated volatility. Under the alternative of jumps, however, the consistency of the non-robust realized measure fails. Instead, it estimates the total quadratic variation that contains the contribution from jump components. As a result, these two realized measures differ in the presence of jumps. In general, Hausman-type tests are designed to detect whether $\sum_{j=N_t}^{N_{t+1}} c_j^2 = 0$ or $\sum_{j=N_t}^{N_{t+1}} c_j^2 > 0$, where N_t denotes the number of jumps up to time t , and c_j is the (random) size of the jumps. However, $\lambda > 0$ does not imply that $\sum_{j=N_t}^{N_{t+1}} c_j^2 > 0$, given that $\Pr(N_{t+1} - N_t > 0) < 1$. Therefore, such tests have power against realized jumps, but not necessarily against positive jump intensity.

Two techniques are often employed in practice to construct the jump-robust realized measures. The first uses multipower variations, such as bipower variation or tripower variation. Under these measures, the effect of jumps is asymptotically “removed” by using the product of consecutive high-frequency observations. The second uses jump thresholding that allows for the separation of jump and continuous components, based on the difference between their orders of magnitude. (see, e.g. Mancini (2009) and Corsi

et al. (2010)). Recent higher order power variation tests are motivated by the fact that for $p > 2$, $\sum_{i=1}^{n-1} |\ln X_{(t+(i+1)\Delta)} - \ln X_{t+i\Delta}|^p$ converges to $\sum_{t \leq s \leq t+1} |\ln X_s - \ln X_{s-}|^p$, where $\sum_{t \leq s \leq t+1} |\ln X_s - \ln X_{s-}|^p$ is strictly positive if there are jumps, and zero otherwise (see, e.g. Aït-Sahalia and Jacod (2009) and Aït-Sahalia et al. (2012)). In this case, however, test power obtains still because of realized jumps, and not because of positive jump probability.

Additionally, other recent tests related to those discussed above have been proposed that are based on comparisons using pre-averaged volatility measures, in order to obtain tests that are robust to microstructure noise (see, e.g. Podolskij and Vetter (2009a,b) and Aït-Sahalia et al. (2012)).

In the Monte Carlo and empirical experiments reported in this paper, we consider three fixed time span tests based on multipower variation, jump thresholding and higher order power variation, respectively.

Generally, jump tests performed over a fixed time span are designed to distinguish between:⁴

$$\Omega_{t,l}^C = \{\omega : s \rightarrow \ln X_s(\omega) \mid \Delta \ln X_s = \ln X_s(\omega) - \ln X_{s-}(\omega) = 0, \forall s \in [t, t+l]\}$$

and

$$\Omega_{t,l}^J = \{\omega : s \rightarrow \ln X_s(\omega) \mid \Delta \ln X_s = \ln X_s(\omega) - \ln X_{s-}(\omega) \neq 0, \forall s \in [t, t+l]\},$$

where l indicates a fixed time span. Hence, all of the tests discussed above are dependent upon pathwise behavior. Clearly, one might decide in favor of $\Omega_{t,l}^C$, even if $\lambda > 0$, simply because jumps are by coincidence absent over the interval $[t, t+l]$. Now, in order to carry out a consistent test against positive jump intensity, two approaches may be used. First, one may consider the following joint hypothesis:

$$\Omega_T^C = \cap_{t=0}^{T-1} \Omega_{t,l}^C, \text{ as } T \rightarrow \infty,$$

versus its negation. Here, the objective is to test the joint null hypothesis that none of the fixed-span sample paths contain jumps. In fact, under mild conditions on the degree of heterogeneity of the process, failure to reject $\Omega_\infty^C = \lim_{T \rightarrow \infty} \cap_{t=1}^{T-1} \Omega_{t,l}^C$ implies

⁴Jump test inconsistency has been pointed out by Huang and Tauchen (2005) and Aït-Sahalia and Jacod (2009), among others.

failure to reject $\lambda = 0$. The difficulty lies in how to implement a test for Ω_T^C , when T gets large. Needless to say, sequential application of fixed time span jump tests leads to sequential test bias, and for T large, Ω_T^C is rejected with probability approaching unity. This is because the empirical size of the joint hypothesis test based on the sequential strategy is $\hat{\alpha}_T = 1 - \prod_{i=1}^T (1 - \hat{\alpha}_i)$, where $\hat{\alpha}_i$ is the empirical size of the i^{th} individual fixed time span test. As a result,

$$\begin{aligned} \lim_{T \rightarrow \infty} \hat{\alpha}_T &= \lim_{T \rightarrow \infty} 1 - \prod_{i=1}^T (1 - \hat{\alpha}_i) \\ &= 1 - \lim_{T \rightarrow \infty} \prod_{i=1}^T (1 - \hat{\alpha}_i) \\ &= 1. \end{aligned}$$

In our Monte Carlo simulations, we illustrate this issue under a set of experiments conducted with an increasing time span. One common approach to this problem is based on controlling the overall Family-Wise Error-Rate (FWER), which ensures that no single hypothesis is rejected at a level larger than a fixed value, say α . This is typically accomplished by sorting individual p -values, and using a rejection rule which depends on the overall number of hypotheses. For further discussion, see Holm (1979), who develops modified Bonferroni bounds, White (2000), who develops the so-called “reality check”, and Romano and Wolf (2005), who provide a refinement of the reality check. However, when the number of hypotheses in the composite grows with the sample size, the null will (almost) never be rejected. In other words, approaches based on the FWER are quite conservative.

An alternative approach, which allows for the number of hypotheses in the composite to grow to infinity, is based on the Expected False Discovery Rate (E-FDR). When using this approach, one controls the expected number of false discoveries (rejections). For further discussion, see Benjamini and Hochberg (1995) and Storey (2003). Although the E-FDR approach applies to the case of a growing number of hypotheses, it is very hard to implement in the presence of generic dependence across p -values, as in our context.

The above discussion, when coupled with issues of identification and test consistency,

provides ample impetus for using long time span jump tests of the variety discussed in the sequel. Still, it should be noted that researchers have shown good performance of fixed time span tests over a day or a week, and we provide further evidence on this front in our Monte Carlo experiments. However, almost no one considers performing the tests over a year or even a decade. The only exception that we are aware of is Huang and Tauchen (2005). They propose using “full-sample statistics” based on *BNS* test statistics. They show that when the time span is long, the *BNS* test over-rejects the null of no realized jumps, since the approximation error on a short interval accumulates as the time span increases. Consequently, the empirical size is biased upwards.

2.4 Long Time Span Jump Intensity Test

Assume the existence of a sample of n observations over an increasing time span, T , and a shrinking discrete interval Δ , so that $n = \frac{T}{\Delta}$, with $T \rightarrow \infty$ and $\Delta \rightarrow 0$. Our interest lies in the following hypotheses:

$$H_0 : \lambda = 0$$

versus

$$H_A = H_A^{(1)} \cup H_A^{(2)} : \left(\lambda > 0 \text{ and } E\left((Z_t - E(Z_t))^3\right) \neq 0 \right) \\ \cup \left(\lambda > 0 \text{ and } E\left((Z_t - E(Z_t))^3\right) = 0 \right).$$

Notice that the alternative hypothesis is the union of two different alternatives, designed to allow for both symmetric and asymmetric jump size density. This property also characterizes the closely related jump test discussed in Corradi et al. (2018), although their test differs in a key respect. Namely, their test is dependent on jump thresholding.

Let $Y_{k\Delta} = \ln X_{k\Delta} - \frac{\Delta}{T} \sum_{k=2}^n \ln X_{k\Delta}$, and $Y_{(k-1)\Delta} = \ln X_{(k-1)\Delta} - \frac{\Delta}{T} \sum_{k=2}^n \ln X_{(k-1)\Delta}$. Also, let

$$\hat{\lambda}_{T,\Delta} = \frac{1}{T} \sum_{k=2}^n (Y_{k\Delta} - Y_{(k-1)\Delta})^3.$$

Here, $\hat{\lambda}_{T,\Delta}$ is the demeaned sample third moment. Consider the statistic

$$LTS_{T,\Delta} = \frac{\sqrt{T}}{\Delta} \hat{\lambda}_{T,\Delta}. \quad (2.10)$$

The asymptotic behavior of $LTS_{T,\Delta}$ can be analyzed under the following assumption.

Assumption A: (i) $\ln X_t$ is generated by equation (2.1) and V_t is defined in equations (2.2), (2.3), or (2.4). (ii) $\ln X_t$ is generated by equation (2.1) and V_t is defined in equation (2.5). For C a generic constant, (iii) $E(|V_t|^k) \leq C$, $k \geq 3$, (iv) N_t satisfies equations (2.6)-(2.8), and λ_t is either constant, or satisfies equation (2.9). (v) The jump size, Z_t , is independently and identically distributed, and $E(|Z_t|^k) \leq C$, for $k \geq 6$.

Corradi et al. (2014) show that under assumptions A(i) and A(iii)-(v), assuming that as $n \rightarrow \infty$, $T \rightarrow \infty$ and $\Delta \rightarrow 0$, then under $H_0 : LTS_{T,\Delta} \xrightarrow{d} N(0, \omega_0)$, with $\omega_0 = 15E(V_{k,\Delta}^3) + 4(E(V_{k,\Delta}))^3 - 12E(V_{k,\Delta})E(V_{k,\Delta}^2)$. Also, they prove that under $H_A^{(1)}$, there exists an $\varepsilon > 0$, such that: $\lim_{T \rightarrow \infty, \Delta \rightarrow 0} \Pr\left(\frac{\Delta}{\sqrt{T}} |LTS_{T,\Delta}| > \varepsilon\right) = 1$; and under $H_A^{(2)}$, there exists an $\varepsilon > 0$, such that: $\lim_{T \rightarrow \infty, \Delta \rightarrow 0} \Pr(\Delta |LTS_{T,\Delta}| > \varepsilon) = 1$.

It follows immediately that $LTS_{T,\Delta}$ converges to a normal random variable under the null hypothesis, diverges at rate $\frac{\sqrt{T}}{\Delta}$ under the alternative of asymmetric jumps, and diverges at the slower rate of $\frac{1}{\Delta}$ under the alternative of symmetric jumps.

Given that the variance is of a different order of magnitude under the null and under each alternative, the “standard” nonparametric bootstrap is not asymptotically valid. This issue arises because the variance of the bootstrap statistic mimics the sample variance. This implies that the bootstrap statistic is of order Δ^{-1} under the alternative. This is not be a problem under $H_A^{(1)}$, since the statistic is of order $\sqrt{T}\Delta^{-1}$, but is a problem under $H_A^{(2)}$, since the actual and bootstrap statistics would be of the same order. To ensure power against $H_A^{(2)}$, it suffices to ensure that the bootstrap statistic is of a smaller order than the actual statistic. This can be accomplished by resampling observations over a rougher grid, $\tilde{\Delta}$, using the same time span, T .

Set the new discrete interval to be $\tilde{\Delta}$, such that $\Delta/\tilde{\Delta} \rightarrow 0$, and resample, with replacement,

$$\left(Y_{k\tilde{\Delta}}^* - Y_{(k-1)\tilde{\Delta}}^*, \dots, Y_{\tilde{n}\tilde{\Delta}}^* - Y_{(\tilde{n}-1)\tilde{\Delta}}^*\right) \text{ from } \left(Y_{k\tilde{\Delta}} - Y_{(k-1)\tilde{\Delta}}, \dots, Y_{\tilde{n}\tilde{\Delta}} - Y_{(\tilde{n}-1)\tilde{\Delta}}\right), \text{ where } \tilde{n} =$$

$\frac{T}{\Delta}$. Now, let

$$\tilde{\lambda}_{T,\tilde{\Delta}} = \frac{1}{T} \sum_{k=2}^{\tilde{n}} \left(Y_{k\tilde{\Delta}} - Y_{(k-1)\tilde{\Delta}} \right)^3,$$

and

$$\tilde{\lambda}_{T,\tilde{\Delta}}^* = \frac{1}{T} \sum_{k=2}^{\tilde{n}} \left(Y_{k\tilde{\Delta}}^* - Y_{(k-1)\tilde{\Delta}}^* \right)^3.$$

Further, define the bootstrap statistic as

$$LTS_{T,\tilde{\Delta}}^* = \frac{\sqrt{T}}{\tilde{\Delta}} \left(\tilde{\lambda}_{T,\tilde{\Delta}}^* - \tilde{\lambda}_{T,\tilde{\Delta}} \right).$$

Finally, let $c_{\alpha,B,\Delta,\tilde{\Delta}}^*$ and $c_{(1-\alpha),B,\Delta,\tilde{\Delta}}^*$ be the $(\alpha/2)^{th}$ and $(1-\alpha/2)^{th}$ critical values of the empirical distribution of $LTS_{T,\tilde{\Delta}}^*$, constructed using B bootstrap replications. Corradi

et al. (2014) show that under assumptions A(i) and A(iii)-(v), and assuming that as $n \rightarrow \infty$, $B \rightarrow \infty$, $T \rightarrow \infty$, $\Delta \rightarrow 0$, $\tilde{\Delta} \rightarrow 0$ and $\Delta/\tilde{\Delta} \rightarrow 0$, then under H_0 :

$$\lim_{T,B \rightarrow \infty, \Delta, \tilde{\Delta} \rightarrow 0} \Pr \left(c_{\alpha/2,B,\Delta,\tilde{\Delta}}^* \leq LTS_{T,\Delta} \leq c_{(1-\alpha/2),B,\Delta,\tilde{\Delta}}^* \right) = 1 - \alpha;$$

and under $H_A^{(1)} \cup H_A^{(2)}$:

$$\lim_{T,B \rightarrow \infty, \Delta, \tilde{\Delta} \rightarrow 0} \Pr \left(c_{\alpha/2,B,\Delta,\tilde{\Delta}}^* \leq LTS_{T,\Delta} \leq c_{(1-\alpha/2),B,\Delta,\tilde{\Delta}}^* \right) = 0.$$

It is immediate to see that rejecting the null whenever $\frac{\sqrt{T}}{\Delta} \hat{\lambda}_{T,\Delta} < c_{\alpha/2,B,\Delta,\tilde{\Delta}}^*$ or $\frac{\sqrt{T}}{\Delta} \hat{\lambda}_{T,\Delta} > c_{(1-\alpha/2),B,\Delta,\tilde{\Delta}}^*$, and otherwise failing to reject, delivers a test with asymptotic size equal to α and asymptotic power equal to unity. Note that the bootstrap statistic is of P^* -probability order $\frac{1}{\Delta}$ under both $H_A^{(1)}$ and $H_A^{(2)}$, while the actual statistic is of P -probability order $\frac{\sqrt{T}}{\Delta}$ under $H_A^{(1)}$ and $\frac{1}{\Delta}$ under $H_A^{(2)}$. Hence, the condition that $\Delta/\tilde{\Delta} \rightarrow 0$ ensures unit asymptotic power under $H_A^{(2)}$.

The $LTS_{T,\Delta}$ test is not robust to leverage. In particular, the results presented above rely on the fact that under the null of no jumps, returns are symmetrically distributed. More precisely, all results are derived under the assumption that $E\left((Y_{k\Delta} - Y_{(k-1)\Delta})^3\right) = 0$, whenever there are no jumps. However, in the presence of leverage, (i.e. V_t is generated as in (2.5)), $E\left(\left(\int_{(k-1)\Delta}^{k\Delta} V_s^{1/2} dW_{1,s}\right)^3\right) \neq 0$, and is instead of order Δ^2 . For example, if V_t is generated by a square root process (i.e., $dV_t = \kappa(\theta - V_t)dt + \eta V_t^{1/2} dW_{2,t}$), then

$E\left((Y_{k\Delta} - Y_{(k-1)\Delta})^3\right) = \lambda E(Z_t - E(Z_t))^3 \Delta + \frac{\eta\theta\rho}{2\kappa} \Delta^2$ (see Aït-Sahalia et al. (2015)). Although, the contribution to the third moment of the asymmetric jump component is of a larger order than that of the leverage component, inference based on the comparison of $LTS_{T,\Delta}$ with the bootstrap critical values $c_{\alpha,B,\Delta,\tilde{\Delta}}^*$ and $c_{(1-\alpha),B,\Delta,\tilde{\Delta}}^*$ will lead to the rejection of the null of no jumps, even if the null is true. To avoid spurious rejection due to the presence of leverage, use the following modified statistic:

$$\widetilde{LTS}_{T,\Delta} = \frac{1}{T^{1/2+\varepsilon}} LTS_{T,\Delta}, \quad (2.11)$$

with $\varepsilon > 0$, arbitrarily small. For this test statistic, Corradi et al. (2014) show that under assumptions A(ii)-(v) hold, and assuming that as $n \rightarrow \infty$, $T \rightarrow \infty$, $\Delta \rightarrow 0$, $\tilde{\Delta} \rightarrow 0$, and $(T^{1/2+\varepsilon}\Delta)/\tilde{\Delta} \rightarrow 0$, then under H_0 :

$$\lim_{T,B \rightarrow \infty, \Delta, \tilde{\Delta} \rightarrow 0} \Pr\left(c_{\alpha/2,B,\Delta,\tilde{\Delta}}^* \leq \widetilde{LTS}_{T,\Delta} \leq c_{(1-\alpha/2),B,\Delta,\tilde{\Delta}}^*\right) = 1;$$

and under $H_A^{(1)} \cup H_A^{(2)}$:

$$\lim_{T,B \rightarrow \infty, \Delta, \tilde{\Delta} \rightarrow 0} \Pr\left(c_{\alpha/2,B,\Delta,\tilde{\Delta}}^* \leq \widetilde{LTS}_{T,\Delta} \leq c_{(1-\alpha/2),B,\Delta,\tilde{\Delta}}^*\right) = 0.$$

It follows that inference based on the comparison of $\widetilde{LTS}_{T,\Delta}$ with the bootstrap critical values $c_{\alpha,B,\Delta,\tilde{\Delta}}^*$ and $c_{(1-\alpha),B,\Delta,\tilde{\Delta}}^*$ delivers a test with zero asymptotic size and unit asymptotic power.

2.5 Fixed Time Span Realized Jump Tests

In this section, we briefly review three fixed time span realized jump tests that are evaluated in our Monte Carlo and empirical experiments. These tests are the *ASJ*, *BNS* and *PZ* tests discussed above.

2.5.1 Aït-Sahalia and Jacod (*ASJ*: 2009) Test

Aït-Sahalia and Jacod (2009) propose a jump test based on calculating the ratio between two realized higher order power variations with different sampling intervals Δ and $k\Delta$,

respectively, where k is an integer chosen prior to test construction. The p^{th} order realized higher order power variation is defined as follows,

$$\widehat{B}(p, \Delta) = \sum_{k=2}^{\lfloor 1/\Delta \rfloor} |\ln X_{k\Delta} - \ln X_{(k-1)\Delta}|^p \quad (2.12)$$

The ratio between two realized power variations with different sampling intervals is then,

$$\widehat{S}(p, k, \Delta) = \frac{\widehat{B}(p, k\Delta)}{\widehat{B}(p, \Delta)}. \quad (2.13)$$

The test statistic is defined as,

$$ASJ = \frac{k^{\frac{p}{2}-1} - \widehat{S}(p, k, \Delta)}{\sqrt{\widehat{V}_n^c}}, \quad (2.14)$$

where in the denominator, $\widehat{V}_n^c$, can be estimated either using a truncated estimator,

$$\widehat{V}_n^c = \Delta \frac{\widehat{A}(2p, \Delta)M(p, k)}{\widehat{A}(p, \Delta)^2}, \quad (2.15)$$

where $\widehat{A}(p, \Delta)$ is defined as follows,

$$\widehat{A}(p, \Delta) = \frac{\Delta^{1-p/2}}{\mu_p} \sum_{k=2}^{\lfloor 1/\Delta \rfloor} |\ln X_{k\Delta} - \ln X_{(k-1)\Delta}|^p \mathbf{1}_{\{|\ln X_{k\Delta} - \ln X_{(k-1)\Delta}| \leq \alpha \Delta^\varpi\}}, \quad (2.16)$$

or using a multipower variation estimator,

$$\widetilde{V}_n^c = \Delta \frac{M(p, k)\widetilde{A}(p/([p]+1), 2[p]+2, \Delta)}{\widetilde{A}(p/([p]+1), [p]+1, \Delta)^2}, \quad (2.17)$$

where

$$\widetilde{A}(r, q, \Delta) = \frac{\Delta^{1-qr/2}}{\mu_r^q} \sum_{k=q}^{\lfloor 1/\Delta \rfloor - q + 1} \prod_{j=0}^{q-1} |\ln X_{(k+j)\Delta} - \ln X_{(k+j-1)\Delta}|^r, \quad (2.18)$$

$$M(p, k) = \frac{1}{\mu_p^2} (k^{p-2}(1+k)\mu_{2p} + k^{p-2}(k-1)\mu_p^2 - 2k^{p/2-1}\mu_{k,p}),$$

and

$$\mu_r = \mathbb{E}(|U|^r) \quad \text{and} \quad \mu_{k,p} = \mathbb{E}(|U|^p |U + \sqrt{k-1}V|^p),$$

for $U, V \stackrel{\text{i.i.d}}{\sim} N(0, 1)$.

In practice, for any significance level α , if $ASJ > Z_\alpha$, where Z_α is the $(1 - \alpha)^{th}$ quantile of the standard normal distribution, one rejects the null of no jumps on the fixed interval $[0, 1]$.

2.5.2 Barndorff-Nielsen and Shephard (*BNS: 2006*) Test

The Barndorff-Nielsen and Shephard (2006) test compares the difference between two estimators of integrated volatility; one which is robust to jumps and the other which is not, to test for jumps in a particular sample path. Barndorff-Nielsen and Shephard (2004) introduce the realized bipower variation (*BPV*) which is a robust estimator of the integrated volatility. Namely, they consider

$$BPV = \frac{\pi}{2} \left(\frac{m}{m-1} \right) \sum_{k=2}^{[1/\Delta]} |\ln X_{(k+1)\Delta} - \ln X_{k\Delta}| |\ln X_{k\Delta} - \ln X_{(k-1)\Delta}|, \quad (2.19)$$

where $m = [1/\Delta]$. Realized *BPV* is a special case of the following realized multipower variation with $p = 2$,

$$MPV(p) = \mu_{\frac{2}{p}}^{-p} \left(\frac{m}{m-p+1} \right) \sum_{k=p}^{[1/\Delta]-p+1} \prod_{j=0}^{p-1} |\ln X_{(k+j)\Delta} - \ln X_{(k+j-1)\Delta}|^{\frac{2}{p}}. \quad (2.20)$$

In this paper, we analyze the following version of their test statistic:

$$BNS = \Delta^{-\frac{1}{2}} \frac{1 - \frac{BPV}{RV}}{\sqrt{((\frac{\pi}{2})^2 + \pi - 5) \max(1, \frac{TPV}{(BPV)^2})}}, \quad (2.21)$$

where *RV* is the realized volatility (i.e., the sum of squared high-frequency returns), and *TPV* is tripower variation (i.e., $MPV(3)$).

The authors prove that under the null, $BNS \xrightarrow{d} N(0, 1)$. As a result, one rejects the null of no jumps on some fixed interval $[0, 1]$, if the test statistic $BNS > Z_\alpha$.

2.5.3 Podolskij and Ziggel (*PZ: 2010*) Test

Podolskij and Ziggel (2010) modify the original truncated power variation statistic proposed in Mancini (2009) by introducing a sequence of positive *i.i.d.* random variables $\{\eta_i\}_{i \in [1, [1/\Delta]]}$, with expectation one and finite variance. Namely, they consider

$$T(\ln X, p) = \Delta^{\frac{1-p}{2}} \sum_{k=2}^{[1/\Delta_n]} |\ln X_{k\Delta} - \ln X_{(k-1)\Delta}|^p (1 - \eta_i \mathbf{1}_{\{|\ln X_{k\Delta} - \ln X_{(k-1)\Delta}| \leq \alpha \Delta_n^{\frac{p}{2}}\}}). \quad (2.22)$$

The test statistic that they propose has the following form,

$$PZ = \frac{T(\ln X, p)}{Var^*(\eta) \hat{A}(2p, \Delta)}, \quad (2.23)$$

where $\widehat{A}(2p, \Delta)$ is the original truncated power variation in equation (2.16). The authors prove that under the null of no jumps, $PZ \xrightarrow{d} N(0, 1)$, and explodes under the alternative. As a result, one rejects the null if $PZ > Z_\alpha$.

2.6 Monte Carlo Simulations

In this section, we report the results of Monte Carlo experiments used to analyze the finite sample properties of the tests introduced above. The simulations are designed to show: (i) the relevance of inconsistency of the fixed time span tests, when tested against non-zero jump intensity in the underlying DGP; (ii) the relevance (or lack thereof) of sequential testing bias when performing daily jump tests, sequentially, along sample paths with a long time span; (iii) the empirical size and power of fixed time span jump tests when applied directly to samples with long time spans; and (iv) the finite sample properties of the $LTS_{T,\Delta}$ and $\widetilde{LTS}_{T,\Delta}$ tests.

The DGP under the null hypothesis in all simulations is the following stochastic volatility model,

$$\begin{aligned} \ln X_t &= \ln X_0 + \int_0^t \bar{\mu} ds + \int_0^t \sigma_s dW_s, \\ \sigma_t^2 &= \sigma_0^2 + \kappa_\sigma \int_0^t (\bar{\sigma}^2 - \sigma_s^2) ds + \zeta \int_0^t \sqrt{\sigma_s^2} dB_s, \end{aligned} \tag{2.24}$$

where the stochastic volatility follows a square root process. Leverage effects are characterized by $\text{corr}(dW_s, dB_s) = \rho$, where $\rho = \{0, -0.5\}$. Under the alternative, we simulate jumps as a compound Poisson process. Namely, we add $\sum_{i=1}^{N_t} J_i$ to the log-price equation, where N_t is a Poisson process characterized by intensity parameter λ , which determines the frequency of jump arrivals, and the J_i are independently and identically drawn from either a normal distribution or an exponential distribution, jump magnitude parameter. All parameter values for the various DGP permutations considered are given in Table 2.1. Of note is that the parameter values used in our experiments are chosen to regions of the parameter space where the tests shift from having strong finite sample properties to having weaker finite sample properties. Thus, for example, while we broadly mimic the parameterizations used in the extant literature (see e.g., Huang and Tauchen (2005) and Aït-Sahalia and Jacod (2009)), in some cases, our parameters

are slightly smaller. For example, Huang and Tauchen (2005) have jump magnitude standard deviation parameters ranging from 0.5 to 2.0, while ours range from 0.25 to 1.25. The sampling frequency in our simulations is 5-minute (i.e., 78 observations per day). Using the Milstein discretization scheme, we simulate log-price sample paths over $T = 500$ days, so that there are 39000 observations in total, for each sample path. Simulation results are calculated based on 1000 replications, and tests are implemented using 0.05 and 0.10 significance levels.

Table 2.2 reports empirical size of daily fixed time span jump tests. In this table, however, the test is applied in two different ways. For entries under the “Jump Days” header, the empirical size of the daily tests are reported. One can think of these experiments as reporting rejection frequencies of 500,000 tests (since $T = 500$ and there are 1000 Monte Carlo replications). For entries under the “Sequential Testing Bias” header, sequences of T tests (corresponding the the length of our daily samples) are run, and overall rejection frequencies across all T tests are reported, where T ranges from 1 to 500 days. Thus, these entries indicate the accumulation of sequential testing bias associated with repeated application of the tests across multiple days. Turning first to the “Jump Days” empirical size results, it is evident that the *BNS* test is least favorably sized, as expected, while the *ASJ* test is very accurately sized, across all *DGPs*; and is not consistently undersized at 0.05 significance level, like the *PZ* test. Now, consider the “Sequential Testing Bias” results in the table. As expected, sequential testing bias leads to a 1.000 rejection rate as T increases beyond 50 days, and these rejections rates are achieved surprisingly quickly, as T increases, although it is interesting to note that the *PZ* test suffers from slightly less bias, for smaller values of T .

Table 2.3 reports empirical power of daily fixed time span jump tests, defined as the rejection rate of daily jump tests across each individual day in each sample path, averaged across all 1000 replications. As in Table 2.2, one can think of these experiments as reporting rejection frequencies of 500,000 tests. Interestingly, power is often small, even when $\lambda = 0.4$, which is a relatively large value, for finite-activity jumps. Among the three jump tests, the *ASJ* test has the lowest power, while *BNS* and *PZ* test are

somewhat better. In interpreting these results, note that, intuitively speaking, the empirical power of daily jump tests against non-zero jump intensity is largely determined by the magnitude of the jump intensity, since this parameter determines the frequency or probability of jump arrivals. Even if these tests have good power against jumps when they occur, for daily intervals without any jumps, it is not surprising to observe that these tests do not reject the null in favor of non-zero jump intensity. Therefore, as long as the intensity is finite, the probability of jumps not occurring on a particular fixed interval is positive, which in turn affects the empirical power of all fixed time span jump tests. However, the tests are also clearly impacted by jump size magnitude. For example, when σ increase from 0.25 to 1.25 (compare *DGPs* 3 and 4 with *DGPs* 5 and 6 - symmetric jumps, or compare *DGPs* 7 and 8 with *DGPs* 9 and 10 - asymmetric jumps), in which cases, empirical power increases by around 30% under symmetric jumps. The exception is the *BNS* and *PZ* tests, which show little power improvement, under the asymmetric jump case. However, there is still a trade-off between the three tests, as the *ASJ* test has overall less power for the case of symmetric jumps.

Tables 2.4-2.7 report findings from experiments where the “entire” sample of T days was used in a single application of the fixed time span tests. This testing strategy is of interest, because there is no reason that fixed time span tests need be implemented using only one day worth of data; and when they are implemented in this manner, they constitute a direct alternative to the use of our long time span tests. First, turn to Table 2.4, where empirical size is reported. Among the three tests, *ASJ* is the clear winner, with size remaining stable even when $T = 500$. This is an interesting and surprising result, suggesting the broad usefulness of the *ASJ* test. The *BNS* and *PZ* tests perform as expected, on the other hand. For example, the empirical size of the *PZ* test approached unity very quickly, and is already approximately 0.5, even for $T = 50$, the empirical size is other tests indicating the ability of this test to control size. As expected, and as can be seen upon inspection of Table 2.5-2.7, the empirical power of all three tests approaches unity quickly as T increases. For example, the empirical power of the *ASJ* test are over 0.9 for almost all *DGPs*, when $T = 50$. In summary, the *ASJ* test is well sized and has great power under long time span testing. This test, thus, is

a clear alternative to the long time span tests discussed in the sequel.

Tables 2.8-2.10 contain the results of our experiments run using the $LTS_{T,\Delta}$ and $\widetilde{LTS}_{T,\Delta}$ tests. As discussed above, the leverage-robust test (i.e., $\widetilde{LTS}_{T,\Delta}$) sacrifices power in order to ensure robustness against leverage effects. Specifically, in equation (2.11), due to the extra term, $\frac{1}{T^{1/2+\epsilon}}$, the leverage-robust test statistic diverges at rate $\frac{1}{T^\epsilon\Delta}$ and $\frac{1}{T^{1/2+\epsilon}\Delta}$, under the alternatives of asymmetric and symmetric jumps, respectively. In practice, the sampling interval, Δ , is usually small and fixed, and the test statistic under the alternatives shrinks with an increasing time span, particularly when jumps are symmetric, while the bootstrapped critical values are of order $1/\widetilde{\Delta}$. As a result, when constructing the leverage-robust test, we propose a rule-of-thumb called our “ T -varying” strategy, in order to choose the subsampling interval, $\widetilde{\Delta}$, used in bootstrapping (see notes to Table 2.8 for details). This rule-of-thumb results in improved power in our experiments. However, it is an ad-hoc data driven method, and further research into its properties remains to be done. Summarizing, we utilize coarser $\widetilde{\Delta}$, as T is grows. Since a coarser $\widetilde{\Delta}$ (i.e., a larger subsampling interval), diminishes the magnitude bootstrap critical values, this strategy (partially) offsets decreases in power that are due to the adjustment term being inversely proportional to T . Size trade-offs associated with using this method are found to be small, and hence the method is utilized in all of our leverage-robust testing experiments, and later in our empirical analysis.

Turning to the results of these tables, first consider empirical size (see Table 2.8). It is immediately apparent that the $LTS_{T,\Delta}$ test has good empirical size for DGP1 (i.e., the “no leverage” case). However, and as expected, size diverges when there is leverage. Again as expected, $\widetilde{LTS}_{T,\Delta}$ has zero empirical size, regardless of the presence of leverage, for values of T greater than 5. Interestingly, though when $T = 5$, the test is approximately correctly sized; thus indicating that our long time span test is an alternative to the short time span tests discussed earlier for small values of T . Of course, T should clearly not be equal to one for the application of the long time span tests. Finally, notice in Table 2.9 that the empirical power of the $LTS_{T,\Delta}$ is good across all cases, including the case where jumps are symmetric with small magnitudes

of $\sigma = 0.25$ (i.e., *DGPs* 3 and 4). Finally, turn to Table 2.10, where empirical power of the $\widetilde{LTS}_{T,\Delta}$ is reported. As expected, empirical power is sacrificed, particularly when jumps are symmetric and $\sigma = 0.25$ (i.e., *DGPs* 3 and 4). However, when $\sigma = 1.25$ (i.e., *DGPs* 5 and 6), this sacrifice is substantially reduced, and power is quite good in all cases, even when λ is small. Coupled with our earlier findings concerning size, we thus have strong evidence that the $\widetilde{LTS}_{T,\Delta}$ is an adequate test for evaluating the presence of jumps in long time spans.

2.7 Empirical Examination of Stock Market Data

2.7.1 Data

We analyze intraday TAQ stock price data sampled at a 5-minute frequency, for the period including observations from the beginning of 2006 through 2013. In particular, we examine: (i) twelve individual stocks including American Express Company (AXP), Bank of America Corporation (BAC), Cisco Systems, Inc. (CSCO), Citigroup Inc. (C), The Coca-Cola Company (KO), Intel Corporation (INTC), JPMorgan Chase & Co. (JPM), Merck & Co., Inc. (MRK), Microsoft Corporation (MSFT), The Procter & Gamble Company (PG), Pfizer Inc. (PFE) and Wal-Mart Stores, Inc. (WMT)); nine sector ETFs including Materials Select Sector SPDR ETF (XLB), Energy Select Sector SPDR ETF (XLE), Financial Select Sector SPDR ETF (XLF), Industrial Select Sector SPDR ETF (XLI), Technology Select Sector SPDR ETF (XLK), Consumer Staples Select Sector SPDR ETF (XLP), Utilities Select Sector SPDR ETF (XLU), Health Care Select Sector SPDR ETF (XLV) and Consumer Discretionary Select Sector SPDR ETF (XLY); and (iii) the SPDR S&P 500 ETF (SPY). Overnight returns are excluded from our dataset.

2.7.2 Empirical Findings

Turning our discussion first to Figures 2.1–2.2, note that the bar charts in these figures depict annual ratios of jump days for all of our stocks and ETFs, based on application of the *ASJ*, *BNS*, and *PZ* tests (see legend to Figure 2.1). For example, 0.2 indicates

that there were jumps found on 20% of the trading days in a given year. As expected, jumps are widely detected in asset prices and indexes over almost any year. Sometimes, the annual percentage of jump days even appears to be inconceivably large, at near 50%. Additionally, while the alternative tests often perform similarly (e.g. all three testing methods find jumps during around 40% of the days in 2006 for XLU and XLP), there are substantial differences for some stocks (e.g. in 2013 the *PZ* tests detects jumps twice as frequently as the other fixed time span tests).

As expected, the *ASJ* test is the most conservative among the three tests. In almost all cases, the *ASJ* test detects the fewest number of “jump days”. For instance, in 2008, *ASJ* test only finds 7.5% jump-days for XLK, while the *PZ* and *BNS* tests find jumps on 17.4% and 22.5% of days, respectively. For SPY, the *ASJ* test finds around 1/3 as many jumps as the other tests, in 2009. This finding is consistent with evidence from our Monte Carlo experiments (see Table 2.3). However, even with the most conservative test, we regularly detect over 15% jump days for many assets, including XLV for 2006 through 2010, XLB and XLY in 2006, and XLF and XLI for 2006 and 2007. Additionally, jump-day percentages are generally larger for our ETFs than for individual stocks, as should be expected. Still, it is also apparent, upon inspection of the figures, that the percentage of jumps detected in our ETFs is declining over time, on an annual basis. This pattern does not characterize individual stocks, however. We conjecture that a possible reason for this is that ETFs were not as frequently traded in the early years of our sample. For instance, typical daily trading volume for XLP or XLY was around 1 million, including pre-market trading and after-hours trading volumes, between 2006 and 2008. This volume is around 1% to 10% of the trading volume of BAC, and 0.15% to 2.5% of the trading volume of SPY, over the same period.

We now turn to a discussion of the results tabulated in Tables 2.11–2.16. In these tables, jump tests results based on the examination of long time spans are reported for the *ASJ*, $LTS_{T,\Delta}$ and $\widetilde{LTS}_{T,\Delta}$ tests. In these tables, tabulated entries are test statistics, and those entries with *, **, and *** indicate rejections of the “no-jump” null at 0.10, 0.05, and 0.01 significance levels, respectively. In these tables, the “long

span” considered is one year, corresponding to the period of time for which annual jump-day ratios were reported in Figures 2.1–2.2. Consider first the results of the *ASJ* test reported in Table 2.11 for our ETFs. Interestingly, there are various ETFs for which no jumps are found. For example, for XLE, no jumps are found in 2006, 2008, 2010, and 2011. In 2011, no jumps are found for 7 of 10 ETFs. Still, in 2007, jumps are found for all 9 ETFs, and in 2008, jumps are found for 7 ETFs. Thus, the evidence concerning jumps appears much more nuanced when the *ASJ* tests is utilized using long time spans. Of course, of discussion above concerning trading volume effects during the early years of our sample still applies, however. Thus, it is difficult to be sure whether the increase in the frequency of jumps found in earlier years for our ETFs is an indicator of the ensuing financial collapse of 2008, or whether this finding is simply an artifact of the data. We leave further investigation of this issue to future research.

Turning to Tables 2.12 and 2.13, which again report on ETFs, note that these tables include results for the $LTS_{T,\Delta}$ (Table 2.12) and $\widetilde{LTS}_{T,\Delta}$ (Table 2.13) tests. As expected, given our Monte Carlo findings, and assuming the presence of leverage, rejections based on the $LTS_{T,\Delta}$ test are not only frequent, but are actually more frequent than rejections based on the *ASJ* test. Indeed, given the presence of leverage, these results carry little weight. However, we know that the $\widetilde{LTS}_{T,\Delta}$ performs adequately, given the presence of leverage. It is perhaps not surprising, then, that the number of years for which jumps are found decreases substantially when $\widetilde{LTS}_{T,\Delta}$ is used, relative to when testing using $LTS_{T,\Delta}$. Indeed, in Table 2.13, note that there are many ETFs for which no jumps are found across multiple different years. Still, it should be stressed that while $\widetilde{LTS}_{T,\Delta}$ is robust to the presence of leverage, the cost of making it thus is a reduction in power, as discussed in the previous sections of this paper. Thus, our conjecture is that the “truth” likely lies somewhere between the results reported based on application of the *ASJ* and the $\widetilde{LTS}_{T,\Delta}$ tests. Still, either way, it is clear that application of long time span tests results in fewer findings of jumps. It is this feature of the tests that is most intriguing, given its implications on the specification and estimation of diffusion models.

Finally, Tables 2.14–2.16 contain results that are analogous to those reported in Tables 2.11–2.13, except that individual stocks are analyzed. Interestingly, the test

rejection patterns that appear upon inspection of the entire in these tables confirms our above discussion based on ETF analysis. Namely, there are various years for which no jumps are found based on application of the ASJ test, and this incidence of “non-rejections” increases when one utilizes the $\widetilde{LTS}_{T,\Delta}$ test.

In summary, we conclude that the usual “toolbox” used by financial econometricians might be usefully augmented by including in it long time span ASJ and $\widetilde{LTS}_{T,\Delta}$ tests. If application of the $\widetilde{LTS}_{T,\Delta}$ results in rejection of the no-jumps null hypothesis, then we have very strong evidence of jumps in the DGP. If application of the $\widetilde{LTS}_{T,\Delta}$ does not result in rejection, then it is advisable to check results based on application of the long time span ASJ test. If the ASJ test “rejects”, then one must consider whether the failure to reject based on the $\widetilde{LTS}_{T,\Delta}$ test is a “power” issue. However, if both tests “reject” then evidence of jumps is very strong.

2.8 Concluding Remarks

In this paper we carry out a Monte Carlo and empirical investigation of long time span jump tests designed to indicate whether the jump intensity in the underlying DGPs is identically zero. This approach differs from the fixed time span variety of jump test broadly used in empirical finance, in the sense that fixed time span tests are not consistent, and can result in identification failure if problems associated with sequential testing bias are not carefully addressed, for example. Our Monte Carlo findings indicate that some fixed time span tests are actually quite adequate for detecting jumps using long time spans of data. In particular, the Aït-Sahalia and Jacod (2009) (ASJ) test performs favorably when compared with Corradi et al. (2018) type long time span jump tests. In an empirical illustration, we show that both of these tests find less prevalence of jumps than when a variety of fixed time span jump tests are applied on a daily basis.

Table 2.1: Data Generating Processes in Monte Carlo Experiments

Panel A: Parameter Values

$\kappa_\sigma, \bar{\sigma}, \zeta = \{5, 0.12, 0.5\}$		
$\mu = 0.05$		
$\Delta_n = 1/78$		
$\rho = \{0, -0.5\}$		
$\lambda = \{0.1, 0.4, 0.8\}$		
$J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\mu_J, \sigma_J^2),$	$\{\mu_J, \sigma_J\} = \{0, 0.25\},$	$\{\mu_J, \sigma_J\} = \{0, 5 \times 0.25\}$
	$\{\mu_J, \sigma_J\} = \{\sqrt{0.5}, 0.25\},$	$\{\mu_J, \sigma_J\} = \{2.5 \times \sqrt{0.5}, 5 \times 0.25\}$

Panel B: Data Generating Processes (DGPs)

DGP 1: Eq. (2.24) with $\mu = 0.05, \rho = 0, \kappa_\sigma = 5, \bar{\sigma} = 0.12, \zeta = 0.5$
DGP 2: Eq. (2.24) with $\mu = 0.05, \rho = -0.5, \kappa_\sigma = 5, \bar{\sigma} = 0.12, \zeta = 0.5$
DGP 3: DGP 1 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 0.25^2)$
DGP 4: DGP 2 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 0.25^2)$
DGP 5: DGP 1 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, (5 \times 0.25)^2)$
DGP 6: DGP 2 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, (5 \times 0.25)^2)$
DGP 7: DGP 1 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\sqrt{0.5}, 0.25^2)$
DGP 8: DGP 2 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(\sqrt{0.5}, 0.25^2)$
DGP 9: DGP 1 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(2.5 \times \sqrt{0.5}, (5 \times 0.25)^2)$
DGP 10: DGP 2 + $J_i \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(2.5 \times \sqrt{0.5}, (5 \times 0.25)^2)$

*Notes: DGP 1 is continuous process without leverage effect and DGP 2 is continuous process with leverage effect. DGPs 3-10 are continuous processes with or without leverage effect plus jumps characterized by various jump size densities. See Section 2.6 for complete details.

Table 2.2: Empirical Size of Daily Fixed Time Span Jump Tests and Sequential Testing

Test	Subject	T = 1	T = 5	T = 50	T = 150	T = 300	T = 500
DGP 1							
<i>ASJ</i>	Jump Days	0.112	0.115	0.118	0.120	0.119	0.120
		0.058	0.058	0.059	0.059	0.059	0.060
	Sequential	0.112	0.458	0.999	1.000	1.000	1.000
	Testing Bias	0.058	0.256	0.953	1.000	1.000	1.000
DGP 2							
<i>ASJ</i>	Jump Days	0.100	0.116	0.121	0.120	0.120	0.120
		0.052	0.057	0.059	0.059	0.059	0.059
	Sequential	0.100	0.458	0.998	1.000	1.000	1.000
	Testing Bias	0.052	0.256	0.950	1.000	1.000	1.000
DGP 1							
<i>BNS</i>	Jump Days	0.122	0.150	0.155	0.152	0.153	0.153
		0.071	0.094	0.095	0.093	0.094	0.095
	Sequential	0.122	0.557	1.000	1.000	1.000	1.000
	Testing Bias	0.071	0.380	0.992	1.000	1.000	1.000
DGP 2							
<i>BNS</i>	Jump Days	0.141	0.153	0.156	0.154	0.154	0.154
		0.098	0.094	0.096	0.095	0.095	0.095
	Sequential	0.141	0.571	1.000	1.000	1.000	1.000
	Testing Bias	0.098	0.398	0.993	1.000	1.000	1.000
DGP 1							
<i>PZ</i>	Jump Days	0.120	0.081	0.123	0.113	0.110	0.108
		0.043	0.031	0.050	0.049	0.046	0.045
	Sequential	0.120	0.347	0.999	1.000	1.000	1.000
	Testing Bias	0.043	0.146	0.921	0.999	1.000	1.000
DGP 2							
<i>PZ</i>	Jump Days	0.105	0.075	0.122	0.113	0.110	0.108
		0.046	0.029	0.048	0.049	0.046	0.045
	Sequential	0.105	0.331	0.998	1.000	1.000	1.000
	Testing Bias	0.046	0.140	0.916	0.999	1.000	1.000

*Notes: Entries in this table denote rejection frequencies based on applications of *ASJ*, *BNS* and *PZ* daily fixed time span jump tests. Results for 0.1 (row 1) and 0.05 (row 2) significance levels are reported. T denotes the number of days for which daily fixed time span jump tests are applied. “Jump Days” shows the average percentage of detected jump days at 0.1 and 0.05 significance levels, respectively. “Sequential Testing Bias” shows probability of finding at least one jump at 0.1 and 0.05 significance levels, respectively. See Sections 2.5 and 2.6 for complete details.

Table 2.3: Empirical Power of Daily Fixed Time Span Jump Tests

Jump Intensity	DGP 3	DGP 4	DGP 5	DGP 6	DGP 7	DGP 8	DGP 9	DGP 10
<i>ASJ</i>								
$\lambda = 0.1$	0.127	0.124	0.143	0.135	0.160	0.160	0.186	0.185
	0.068	0.067	0.077	0.077	0.093	0.093	0.118	0.125
$\lambda = 0.4$	0.183	0.170	0.244	0.232	0.285	0.284	0.353	0.353
	0.104	0.103	0.150	0.143	0.167	0.172	0.262	0.262
$\lambda = 0.8$	0.232	0.242	0.345	0.356	0.415	0.431	0.548	0.533
	0.137	0.138	0.210	0.219	0.240	0.252	0.411	0.425
<i>BNS</i>								
$\lambda = 0.1$	0.201	0.185	0.222	0.208	0.247	0.230	0.250	0.232
	0.154	0.124	0.179	0.150	0.208	0.177	0.211	0.179
$\lambda = 0.4$	0.315	0.274	0.379	0.359	0.432	0.404	0.441	0.411
	0.264	0.225	0.339	0.313	0.403	0.371	0.413	0.378
$\lambda = 0.8$	0.427	0.392	0.560	0.534	0.625	0.611	0.633	0.615
	0.377	0.336	0.526	0.499	0.600	0.583	0.612	0.588
<i>PZ</i>								
$\lambda = 0.1$	0.191	0.182	0.218	0.211	0.239	0.234	0.241	0.235
	0.107	0.101	0.136	0.131	0.159	0.157	0.163	0.158
$\lambda = 0.4$	0.295	0.287	0.377	0.369	0.425	0.416	0.431	0.424
	0.217	0.210	0.311	0.298	0.366	0.352	0.372	0.360
$\lambda = 0.8$	0.424	0.397	0.560	0.548	0.634	0.624	0.640	0.628
	0.357	0.339	0.510	0.498	0.589	0.587	0.595	0.591

*Notes: See notes to Table 2.2. Rejection frequencies are given based on repeated daily applications of jump tests across $T = 500$ days, for each Monte Carlo replication. Thus, one can think of these experiments as reporting rejection frequencies of 500,000 tests (since $T = 500$ and there are 1000 Monte Carlo replications).

Table 2.4: Empirical Size of Fixed Time Span Jump Tests When Utilized Using Long-span Samples

Test	T = 5	T = 25	T = 50	T = 150	T = 300	T = 500
DGP 1						
<i>ASJ</i>	0.113	0.109	0.119	0.114	0.135	0.147
	0.051	0.057	0.067	0.066	0.072	0.072
DGP 2						
<i>ASJ</i>	0.106	0.131	0.148	0.136	0.132	0.145
	0.045	0.075	0.069	0.065	0.072	0.080
DGP 1						
<i>BNS</i>	0.132	0.136	0.142	0.150	0.194	0.215
	0.070	0.071	0.075	0.081	0.109	0.127
DGP 2						
<i>BNS</i>	0.127	0.122	0.142	0.162	0.192	0.227
	0.081	0.065	0.080	0.095	0.104	0.132
DGP 1						
<i>PZ</i>	0.116	0.288	0.504	0.849	0.962	0.994
	0.069	0.279	0.484	0.846	0.950	0.993
DGP 2						
<i>PZ</i>	0.119	0.290	0.504	0.869	0.974	0.995
	0.071	0.265	0.482	0.856	0.964	0.995

*Notes: See notes to Table 2.2. Entries are rejection frequencies based on a single application of the *ASJ*, *BNS* and *PZ* tests using long time span samples with T days, for each Monte Carlo replication. For all values of T, 1000 replications are run.

Table 2.5: Empirical Power of the *ASJ* Jump Test When Utilized Using Long-span Samples

Jump Intensity	DGP 3	DGP 4	DGP 5	DGP 6	DGP 7	DGP 8	DGP 9	DGP 10
T = 5								
$\lambda = 0.1$	0.214	0.227	0.341	0.352	0.409	0.408	0.436	0.447
	0.151	0.169	0.285	0.286	0.345	0.351	0.395	0.405
$\lambda = 0.4$	0.498	0.482	0.750	0.731	0.838	0.846	0.876	0.872
	0.401	0.400	0.678	0.668	0.769	0.784	0.856	0.854
$\lambda = 0.8$	0.670	0.653	0.901	0.899	0.941	0.936	0.947	0.939
	0.565	0.563	0.839	0.833	0.897	0.895	0.934	0.923
T = 25								
$\lambda = 0.1$	0.544	0.510	0.835	0.803	0.913	0.913	0.924	0.924
	0.480	0.454	0.812	0.779	0.907	0.900	0.921	0.918
$\lambda = 0.4$	0.908	0.898	0.993	0.996	0.992	0.991	0.986	0.984
	0.872	0.870	0.993	0.994	0.991	0.991	0.986	0.984
$\lambda = 0.8$	0.988	0.988	0.995	0.995	0.979	0.984	0.989	0.985
	0.985	0.977	0.995	0.994	0.978	0.980	0.986	0.978
T = 50								
$\lambda = 0.1$	0.711	0.714	0.960	0.956	0.989	0.986	0.991	0.990
	0.663	0.655	0.954	0.949	0.985	0.985	0.990	0.990
$\lambda = 0.4$	0.987	0.989	0.997	0.998	0.992	0.996	0.993	0.994
	0.979	0.983	0.997	0.997	0.992	0.994	0.993	0.992
$\lambda = 0.8$	0.995	0.997	0.996	0.996	0.990	0.990	0.997	0.996
	0.995	0.994	0.995	0.995	0.989	0.988	0.995	0.994
T = 150								
$\lambda = 0.1$	0.965	0.961	0.999	0.999	0.998	0.997	0.996	0.996
	0.954	0.947	0.999	0.999	0.997	0.997	0.996	0.996
$\lambda = 0.4$	0.998	1.000	1.000	1.000	0.998	0.999	1.000	1.000
	0.998	0.999	1.000	1.000	0.998	0.998	1.000	1.000
$\lambda = 0.8$	0.999	0.999	0.999	0.999	1.000	1.000	1.000	1.000
	0.999	0.999	0.999	0.999	1.000	1.000	1.000	1.000
T = 300								
$\lambda = 0.1$	0.996	0.996	0.999	0.999	0.998	0.998	0.999	1.000
	0.995	0.995	0.999	0.999	0.998	0.998	0.999	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 500								
$\lambda = 0.1$	0.998	0.999	0.999	0.999	0.999	0.999	1.000	1.000
	0.998	0.999	0.999	0.999	0.999	0.999	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

*Notes: See notes to Tables 2.3 and 2.4.

Table 2.6: Empirical Power of the *BNS* Jump Test When Utilized Using Long-span Samples

Jump Intensity	DGP 3	DGP 4	DGP 5	DGP 6	DGP 7	DGP 8	DGP 9	DGP 10
T = 5								
$\lambda = 0.1$	0.264	0.270	0.396	0.390	0.452	0.446	0.464	0.458
	0.197	0.207	0.340	0.344	0.408	0.410	0.423	0.423
$\lambda = 0.4$	0.588	0.599	0.802	0.798	0.885	0.883	0.889	0.887
	0.533	0.540	0.771	0.768	0.871	0.872	0.880	0.878
$\lambda = 0.8$	0.817	0.815	0.948	0.961	0.983	0.981	0.986	0.984
	0.768	0.776	0.942	0.953	0.979	0.978	0.984	0.981
T = 25								
$\lambda = 0.1$	0.537	0.550	0.824	0.833	0.915	0.917	0.928	0.930
	0.471	0.476	0.796	0.804	0.906	0.907	0.923	0.923
$\lambda = 0.4$	0.945	0.943	0.998	0.997	1.000	1.000	1.000	1.000
	0.923	0.928	0.998	0.997	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	0.999	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	0.999	1.000	1.000	1.000	1.000	1.000	1.000
T = 50								
$\lambda = 0.1$	0.707	0.720	0.950	0.958	0.987	0.988	0.993	0.995
	0.625	0.650	0.946	0.944	0.983	0.985	0.993	0.994
$\lambda = 0.4$	0.997	0.996	1.000	1.000	1.000	1.000	1.000	1.000
	0.994	0.995	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 150								
$\lambda = 0.1$	0.947	0.958	1.000	1.000	1.000	1.000	1.000	1.000
	0.917	0.934	0.999	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 300								
$\lambda = 0.1$	0.994	0.996	1.000	1.000	1.000	1.000	1.000	1.000
	0.993	0.994	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 500								
$\lambda = 0.1$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

*Notes: See notes to Tables 2.3 and 2.4.

Table 2.7: Empirical Power of the *PZ* Jump Test When Utilized Using Long-span Samples

Jump Intensity	DGP 3	DGP 4	DGP 5	DGP 6	DGP 7	DGP 8	DGP 9	DGP 10
T = 5								
$\lambda = 0.1$	0.316	0.317	0.418	0.413	0.465	0.456	0.471	0.464
	0.277	0.284	0.385	0.387	0.433	0.433	0.439	0.441
$\lambda = 0.4$	0.682	0.661	0.829	0.810	0.888	0.879	0.890	0.880
	0.656	0.645	0.818	0.801	0.881	0.874	0.883	0.875
$\lambda = 0.8$	0.874	0.865	0.962	0.965	0.981	0.984	0.982	0.985
	0.864	0.851	0.960	0.960	0.979	0.980	0.980	0.981
T = 25								
$\lambda = 0.1$	0.788	0.777	0.911	0.900	0.936	0.935	0.939	0.939
	0.780	0.766	0.907	0.896	0.934	0.933	0.937	0.937
$\lambda = 0.4$	0.993	0.992	1.000	1.000	1.000	1.000	1.000	1.000
	0.993	0.992	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 50								
$\lambda = 0.1$	0.960	0.954	0.991	0.987	0.997	0.993	0.998	0.994
	0.956	0.950	0.991	0.987	0.997	0.993	0.998	0.994
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 150								
$\lambda = 0.1$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 300								
$\lambda = 0.1$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
T = 500								
$\lambda = 0.1$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.4$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
$\lambda = 0.8$	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000
	1.000	1.000	1.000	1.000	1.000	1.000	1.000	1.000

*Notes: See notes to Tables 2.3 and 2.4.

Table 2.8: Empirical Size of $LTS_{T,\Delta}$ and $\widetilde{LTS}_{T,\Delta}$ Jump Tests

Test Statistic	Leverage	T = 5	T = 25	T = 50	T = 150	T = 300	T = 500
$LTS_{T,\Delta}$	$\emptyset$	0.202	0.151	0.129	0.123	0.106	0.112
		0.145	0.094	0.072	0.076	0.055	0.057
	$\sqrt{}$	0.264	0.316	0.461	0.858	0.987	0.999
		0.185	0.221	0.341	0.775	0.968	0.997
$\widetilde{LTS}_{T,\Delta}$	$\emptyset$	0.075	0.007	0.001	0.000	0.000	0.000
		0.048	0.003	0.000	0.000	0.000	0.000
	$\sqrt{}$	0.089	0.002	0.000	0.000	0.000	0.000
		0.053	0.001	0.000	0.000	0.000	0.000

*Notes: As in Tables 2.2-2.7, jump test rejection frequencies are reported. As discussed in Section 2.6, the subsampling interval, $\tilde{\Delta}$, used in constructing critical values for the tests has been selected using a simple rule. Namely, $\{T=25 \text{ and } 50, \tilde{\Delta}_n^1=1/3, \tilde{\Delta}_n^2=1/26\}$; $\{T=150, \tilde{\Delta}_n^1=1, \tilde{\Delta}_n^2=1/13\}$; $\{T=300, \tilde{\Delta}_n^1=2, \tilde{\Delta}_n^2=1/13\}$; $\{T=500, \tilde{\Delta}_n^1=3.2, \tilde{\Delta}_n^2=1/10\}$, where $\tilde{\Delta}_n^1$ is the subsampling interval used in bootstrapping critical values for $\widetilde{LTS}_{T,\Delta}$ and $\tilde{\Delta}_n^2$ is the subsampling interval used in bootstrapping critical values for $LTS_{T,\Delta}$.

Table 2.9: Empirical Power of $LTS_{T,\Delta}$ Jump Test

Jump Intensity	DGP 3	DGP 4	DGP 5	DGP 6	DGP 7	DGP 8	DGP 9	DGP 10
T = 5								
$\lambda = 0.1$	0.346	0.361	0.435	0.463	0.487	0.517	0.507	0.540
	0.274	0.291	0.374	0.396	0.436	0.463	0.461	49.1
$\lambda = 0.4$	0.572	0.585	0.717	0.712	0.895	0.883	0.901	0.895
	0.498	0.514	0.665	0.650	0.879	0.865	0.890	0.887
$\lambda = 0.8$	0.664	0.677	0.790	0.784	0.968	0.973	0.969	0.976
	0.600	0.615	0.752	0.737	0.962	0.964	0.964	0.970
T = 25								
$\lambda = 0.1$	0.472	0.535	0.713	0.733	0.913	0.921	0.926	0.941
	0.400	0.447	0.656	0.670	0.895	0.896	0.917	0.928
$\lambda = 0.4$	0.630	0.634	0.690	0.676	0.997	0.998	1.000	1.000
	0.560	0.537	0.604	0.598	0.996	0.997	0.998	1.000
$\lambda = 0.8$	0.650	0.645	0.647	0.650	1.000	1.000	1.000	1.000
	0.583	0.569	0.580	0.572	1.000	1.000	1.000	1.000
T = 50								
$\lambda = 0.1$	0.559	0.587	0.719	0.736	0.982	0.981	0.990	0.990
	0.459	0.492	0.644	0.657	0.977	0.975	0.989	0.988
$\lambda = 0.4$	0.641	0.656	0.642	0.628	1.000	1.000	0.999	0.999
	0.557	0.561	0.548	0.559	1.000	1.000	0.999	0.999
$\lambda = 0.8$	0.623	0.604	0.628	0.627	1.000	1.000	1.000	1.000
	0.534	0.527	0.531	0.539	1.000	1.000	1.000	1.000
T = 150								
$\lambda = 0.1$	0.763	0.786	0.847	0.828	1.000	1.000	1.000	1.000
	0.722	0.732	0.804	0.788	1.000	0.999	1.000	1.000
$\lambda = 0.4$	0.762	0.768	0.807	0.802	1.000	1.000	1.000	1.000
	0.716	0.717	0.759	0.764	1.000	1.000	1.000	1.000
$\lambda = 0.8$	0.761	0.750	0.805	0.795	1.000	1.000	1.000	1.000
	0.715	0.704	0.756	0.751	1.000	1.000	1.000	1.000
T = 300								
$\lambda = 0.1$	0.766	0.784	0.819	0.825	1.000	1.000	1.000	1.000
	0.709	0.738	0.784	0.786	1.000	1.000	1.000	1.000
$\lambda = 0.4$	0.789	0.768	0.780	0.791	1.000	1.000	1.000	1.000
	0.742	0.713	0.733	0.755	1.000	1.000	1.000	1.000
$\lambda = 0.8$	0.759	0.756	0.808	0.802	1.000	1.000	1.000	1.000
	0.711	0.702	0.764	0.757	1.000	1.000	1.000	1.000
T = 500								
$\lambda = 0.1$	0.791	0.819	0.865	0.855	1.000	1.000	1.000	1.000
	0.755	0.790	0.833	0.825	1.000	1.000	1.000	1.000
$\lambda = 0.4$	0.819	0.808	0.838	0.827	1.000	1.000	1.000	1.000
	0.771	0.774	0.804	0.795	1.000	1.000	1.000	1.000
$\lambda = 0.8$	0.786	0.776	0.823	0.836	1.000	1.000	1.000	1.000
	0.753	0.739	0.784	0.795	1.000	1.000	1.000	1.000

*Notes: See notes to Tables 2.3, 2.4 and 2.8.

Table 2.10: Empirical Power of $\widetilde{LTS}_{T,\Delta}$ Jump Test

Jump Intensity	DGP 3	DGP 4	DGP 5	DGP 6	DGP 7	DGP 8	DGP 9	DGP 10
T = 5								
$\lambda = 0.1$	0.181	0.221	0.300	0.337	0.374	0.401	0.406	0.438
	0.145	0.162	0.272	0.287	0.348	0.355	0.390	0.407
$\lambda = 0.4$	0.428	0.443	0.651	0.642	0.856	0.837	0.871	0.866
	0.371	0.388	0.597	0.585	0.835	0.819	0.865	0.863
$\lambda = 0.8$	0.596	0.586	0.761	0.738	0.950	0.951	0.951	0.952
	0.515	0.513	0.717	0.710	0.943	0.944	0.948	0.948
T = 25								
$\lambda = 0.1$	0.262	0.278	0.670	0.680	0.887	0.879	0.915	0.918
	0.228	0.233	0.644	0.656	0.875	0.863	0.915	0.915
$\lambda = 0.4$	0.567	0.565	0.824	0.811	0.998	0.999	0.998	0.998
	0.505	0.492	0.798	0.783	0.996	0.998	0.998	0.998
$\lambda = 0.8$	0.646	0.643	0.793	0.801	0.997	0.998	1.000	1.000
	0.602	0.579	0.766	0.766	0.997	0.996	1.000	1.000
T = 50								
$\lambda = 0.1$	0.169	0.192	0.694	0.687	0.946	0.959	0.983	0.983
	0.129	0.149	0.630	0.624	0.933	0.947	0.978	0.978
$\lambda = 0.4$	0.436	0.416	0.708	0.696	0.998	0.999	0.998	0.997
	0.349	0.342	0.632	0.634	0.998	0.997	0.995	0.996
$\lambda = 0.8$	0.489	0.469	0.654	0.635	0.996	0.993	0.999	0.999
	0.411	0.386	0.584	0.586	0.995	0.990	0.997	0.997
T = 150								
$\lambda = 0.1$	0.163	0.149	0.766	0.740	0.997	1.000	1.000	1.000
	0.120	0.099	0.726	0.695	0.996	0.998	1.000	1.000
$\lambda = 0.4$	0.320	0.307	0.743	0.742	1.000	1.000	1.000	1.000
	0.245	0.215	0.701	0.709	1.000	1.000	1.000	1.000
$\lambda = 0.8$	0.390	0.353	0.709	0.712	1.000	1.000	1.000	1.000
	0.311	0.286	0.659	0.657	1.000	1.000	1.000	1.000
T = 300								
$\lambda = 0.1$	0.090	0.093	0.769	0.753	1.000	1.000	1.000	1.000
	0.058	0.064	0.728	0.719	1.000	1.000	1.000	1.000
$\lambda = 0.4$	0.223	0.216	0.733	0.744	1.000	1.000	1.000	1.000
	0.158	0.154	0.696	0.708	1.000	1.000	1.000	1.000
$\lambda = 0.8$	0.274	0.248	0.702	0.708	1.000	1.000	1.000	1.000
	0.199	0.188	0.659	0.646	1.000	1.000	1.000	1.000
T = 500								
$\lambda = 0.1$	0.054	0.063	0.757	0.751	1.000	1.000	1.000	1.000
	0.026	0.038	0.710	0.700	1.000	1.000	1.000	1.000
$\lambda = 0.4$	0.153	0.158	0.717	0.700	1.000	1.000	1.000	1.000
	0.091	0.106	0.650	0.651	1.000	1.000	1.000	1.000
$\lambda = 0.8$	0.188	0.171	0.667	0.649	1.000	1.000	1.000	1.000
	0.126	0.129	0.607	0.585	1.000	1.000	1.000	1.000

*Notes: See notes to Tables 2.3, 2.4 and 2.8.

Table 2.11: *ASJ* Jump Test Results for ETFs

	2006	2007	2008	2009	2010	2011	2012	2013
SPY	3.264 (***)	1.694 (**)	0.002	2.579 (***)	5.213 (***)	0.745	0.874	1.745 (**)
XLB	5.011 (***)	2.528 (***)	1.941 (**)	3.207 (***)	4.161 (***)	0.870	2.682 (***)	0.986
XLE	0.952	4.862 (***)	0.019	7.023 (***)	1.061	0.058	5.665 (***)	1.772 (**)
XLF	3.207 (***)	4.128 (***)	1.825 (**)	1.774 (**)	0.843	0.822	1.286 (*)	1.663 (**)
XLI	4.233 (***)	5.903 (***)	1.625 (*)	1.827 (**)	7.367 (***)	1.486 (*)	0.951	1.813 (**)
XLK	7.909 (***)	4.214 (***)	0.991	1.766 (**)	1.384 (*)	0.759	0.922	2.492 (***)
XLP	3.180 (***)	8.996 (***)	7.979 (***)	4.212 (***)	0.373	1.493 (*)	6.078 (***)	1.565 (*)
XLU	7.066 (***)	2.730 (***)	1.481 (*)	10.000 (***)	0.528	0.642	2.769 (***)	3.324 (***)
XLV	7.030 (***)	6.084 (***)	1.828 (**)	2.386 (***)	2.352 (***)	1.729 (**)	0.866	2.530 (***)
XLY	3.368 (***)	1.845 (**)	3.279 (***)	4.399 (***)	0.457	0.735	0.022	2.721 (***)

*Notes: See notes to Tables 2.4 and 2.5. Entries are jump test statistics, and (***), (**), and (*) indicate rejections of the “no jump” null hypothesis at 0.01, 0.05 and 0.1 significance levels, respectively.

Table 2.12: *LTS_{T,Δ}* Jump Test Results for ETFs

	2006	2007	2008	2009	2010	2011	2012	2013
SPY	2.04E-07	-3.88E-06	4.87E-04 (***)	2.65E-05	1.65E-06	-1.19E-05	-3.15E-07	-2.90E-06 (**)
XLB	-1.40E-05	-8.45E-05 (***)	1.35E-03 (***)	-1.90E-04 (***)	-9.21E-05 (***)	9.39E-06	-1.33E-06	1.87E-06
XLE	2.38E-05 (***)	-2.43E-05 (**)	1.07E-03 (**)	-9.77E-05 (**)	4.74E-05 (***)	-4.76E-05	-6.35E-06 (**)	-4.45E-06
XLF	-1.11E-05 (***)	-2.58E-04 (***)	2.00E-03 (***)	-2.26E-05	-5.04E-05 (**)	-3.69E-05	3.76E-06	-6.79E-06 (**)
XLI	-2.03E-05 (***)	9.96E-05 (***)	7.23E-04 (***)	-9.39E-05 (**)	3.24E-04 (***)	9.48E-06	1.98E-06	-2.93E-06 (**)
XLK	-1.99E-04 (***)	-6.68E-05 (***)	1.86E-03 (***)	-2.68E-05	-1.45E-04 (***)	-9.44E-06	-2.63E-06	-3.39E-06 (***)
XLP	5.92E-06 (***)	1.83E-05 (***)	-2.02E-03 (***)	-5.91E-05 (***)	-2.19E-05 (**)	-4.69E-06	-4.30E-05 (***)	8.25E-07
XLU	-9.42E-05 (***)	1.67E-04 (***)	-8.17E-05	-1.04E-01 (***)	1.75E-04 (***)	-1.91E-05	7.89E-06 (***)	-7.47E-06
XLV	-1.17E-04 (***)	9.52E-05 (***)	4.08E-04 (***)	-4.11E-05 (***)	-1.71E-04 (***)	1.24E-05	-8.03E-07	-2.53E-06
XLY	-8.05E-05 (***)	-1.06E-04 (***)	-2.97E-03 (***)	-2.72E-04 (***)	1.62E-05	-1.56E-05	-5.33E-06 (**)	2.57E-07

*Notes: See notes to Tables 2.8 and 2.11.

Table 2.13: $\widetilde{LTS}_{T,\Delta}$ Jump Test Results for ETFs

	2006	2007	2008	2009	2010	2011	2012	2013
SPY	1.28E-08	-2.43E-07 (*)	3.05E-05	1.66E-06	1.03E-07	-7.46E-07	-1.98E-08 (*)	-1.82E-07
XLB	-8.78E-07 (*)	-5.31E-06 (**)	8.41E-05	-1.19E-05	-5.77E-06	5.88E-07	-8.38E-08 (*)	1.17E-07
XLE	1.49E-06	-1.53E-06	6.71E-05	-6.12E-06	2.97E-06	-2.98E-06	-3.99E-07	-2.79E-07
XLF	-6.97E-07	-1.62E-05 (**)	1.25E-04	-1.42E-06	-3.16E-06	-2.31E-06	2.36E-07 (*)	-4.26E-07
XLI	-1.27E-06	6.25E-06	4.52E-05	-5.88E-06	2.03E-05 (***)	5.94E-07	1.24E-07 (*)	-1.83E-07
XLK	-1.25E-05 (***)	-4.19E-06 (**)	1.16E-04	-1.68E-06	-9.06E-06 (**)	-5.91E-07	-1.66E-07	-2.12E-07
XLP	3.71E-07 (*)	1.15E-06	-1.26E-04 (***)	-3.70E-06 (**)	-1.37E-06 (*)	-2.94E-07	-2.70E-06 (***)	5.17E-08
XLU	-5.92E-06	1.05E-05 (***)	-5.11E-06 (*)	-6.49E-03 (***)	1.10E-05 (***)	-1.19E-06	4.96E-07	-4.68E-07
XLV	-7.37E-06 (***)	5.98E-06	2.55E-05	-2.57E-06	-1.07E-05 (***)	7.76E-07	-5.05E-08 (*)	-1.58E-07
XLV	-5.05E-06	-6.68E-06 (**)	-1.86E-04 (**)	-1.71E-05	1.01E-06	-9.79E-07	-3.35E-07 (*)	1.61E-08

*Notes: See notes to Tables 2.8 and 2.11.

Table 2.14: ASJ Jump Test Results for Individual Stocks

	2006	2007	2008	2009	2010	2011	2012	2013
American Express	3.115 (***)	4.248 (***)	2.044 (**)	2.672 (***)	0.996	2.343 (***)	3.555 (***)	2.464 (***)
Bank of America	2.914 (***)	3.014 (***)	1.529 (*)	3.576 (***)	0.819	3.171 (***)	1.784 (**)	0.803
Cisco	3.811 (***)	5.997 (***)	0.180	1.818 (**)	2.865 (***)	1.665 (**)	2.581 (***)	0.922
Citigroup	3.215 (***)	0.802	0.520	3.302 (***)	6.023 (***)	0.752	0.197	0.895
Coca-Cola	8.039 (***)	10.005 (***)	6.134 (***)	4.563 (***)	0.826	1.774 (**)	2.551 (***)	3.476 (***)
Intel	2.632 (***)	4.142 (***)	3.332 (***)	0.914	5.627 (***)	6.425 (***)	1.821 (**)	1.772 (**)
JPMorgan	3.446 (***)	2.532 (***)	3.232 (***)	0.962	1.733 (**)	3.461 (***)	0.987	0.983
Merck & Co.	5.700 (***)	8.016 (***)	1.909 (**)	3.184 (***)	0.051	1.559 (*)	2.648 (***)	0.246
Microsoft	2.982 (***)	6.997 (***)	0.909	0.811	0.652	3.434 (***)	4.579 (***)	0.370
Procter & Gamble	3.285 (***)	4.933 (***)	1.814 (**)	9.998 (***)	2.387 (***)	3.218 (***)	9.003 (***)	1.745 (**)
Pfizer	2.356 (***)	4.024 (***)	0.845	0.890	7.436 (***)	1.960 (**)	1.740 (**)	2.516 (***)
Wal-Mart	0.823	4.011 (***)	1.187	1.730 (**)	0.799	2.439 (***)	0.131	3.461 (***)

*Notes: See notes to Tables 2.8 and 2.11.

Table 2.15: $LTS_{T,\Delta}$ Jump Test Results for Individual Stocks

	2006	2007	2008	2009	2010	2011	2012	2013
American Express	1.45E-04 (***)	-8.03E-05 (***)	4.57E-03 (***)	1.30E-03 (**)	-7.80E-05	-1.34E-04 (**)	-6.61E-05 (***)	6.24E-05 (***)
Bank of America	-8.20E-04 (***)	-1.58E-04 (***)	4.16E-03 (***)	-2.26E-02 (**)	-1.29E-04	-1.39E-03 (***)	4.47E-05	-3.16E-05
Cisco	6.94E-05 (***)	-7.64E+02 (***)	9.32E-04 (***)	9.61E-05	-2.02E-04 (***)	1.31E-04 (***)	-1.85E-05	8.56E-06
Citigroup	-3.93E-05 (***)	-7.93E-05	-1.13E-02	-1.60E-02	-1.64E-03 (***)	-2.61E-04	6.52E-05 (**)	-2.26E-05
Coca-Cola	8.49E-05 (***)	5.56E-05 (***)	-2.05E-03 (***)	1.06E-04 (***)	-3.79E-05	-5.11E-06	-2.31E-05 (***)	-1.46E-05 (***)
Intel	5.55E-05 (***)	-1.49E-04 (***)	1.57E-03 (***)	2.10E-04 (**)	1.72E-04 (***)	-3.69E-05	-1.74E-05	4.64E-05 (***)
JPMorgan	4.15E-05	-1.83E-04 (***)	1.98E-03	-2.43E-04	3.88E-05	3.54E-04 (***)	7.16E-05	1.96E-05
Merck & Co.	1.39E-04 (***)	2.12E-04 (***)	-6.85E-03 (***)	-5.17E-04 (***)	3.81E-04 (***)	4.54E-06	2.80E-05 (***)	6.49E-06
Microsoft	9.93E-06 (**)	-7.46E+02 (***)	1.76E-04	6.35E-05	-4.09E-05	-1.23E-05	-6.93E-05 (***)	1.34E-05
Procter & Gamble	2.76E-05 (***)	8.05E-05 (***)	1.37E-04	-1.29E-03 (***)	2.33E-03 (***)	-2.88E-05 (***)	2.47E-05 (***)	5.72E-06
Pfizer	2.08E-04 (***)	-2.74E-03 (***)	1.11E-04	6.97E-05	3.96E-06	8.69E-05 (**)	9.01E-06	1.78E-05 (***)
Wal-Mart	6.75E-05	8.67E-05 (***)	8.60E-04 (***)	5.27E-05 (***)	-9.57E-06	3.19E-05 (**)	6.87E-06	-9.51E-06

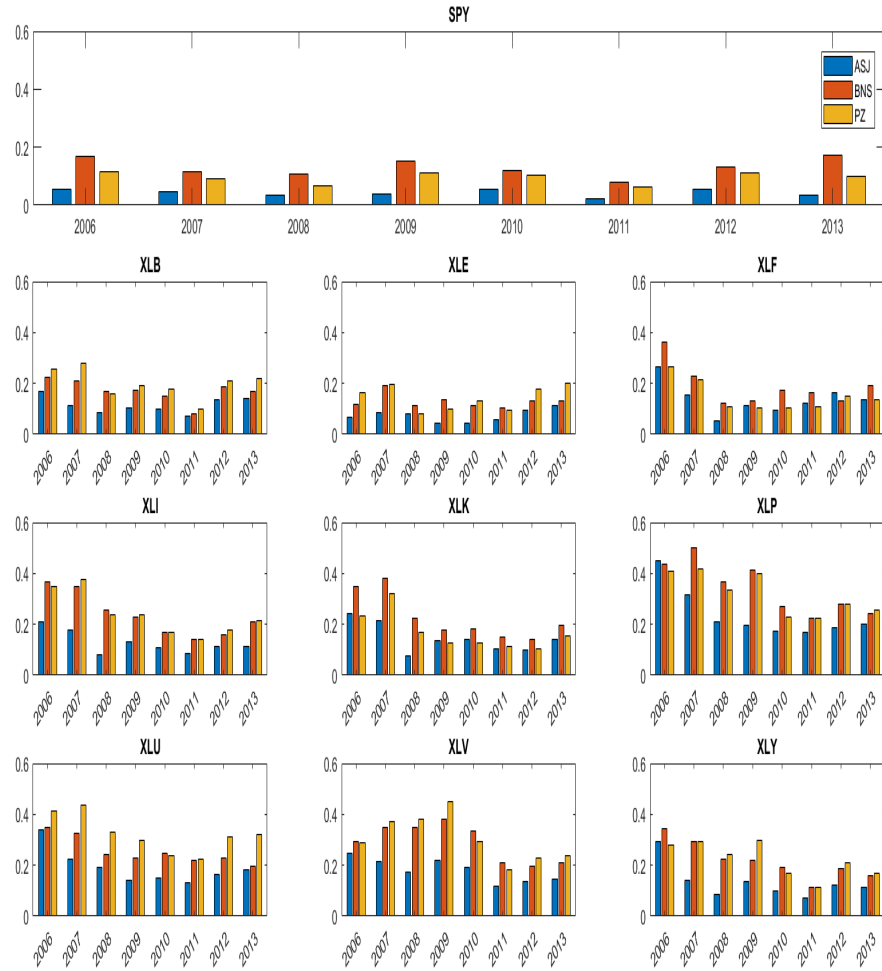
*Notes: See notes to Tables 2.8 and 2.11.

Table 2.16: $\widetilde{LTS}_{T,\Delta}$ Jump Test Results for Individual Stocks

	2006	2007	2008	2009	2010	2011	2012	2013
American Express	9.09E-06 (***)	-5.04E-06 (*)	2.86E-04	8.13E-05	-4.88E-06	-8.42E-06	-4.16E-06	3.91E-06 (**)
Bank of America	-5.15E-05 (***)	-9.94E-06 (**)	2.60E-04 (*)	-1.42E-03	-8.10E-06	-8.69E-05	2.81E-06 (**)	-1.98E-06
Cisco	4.36E-06 (*)	-4.79E+01 (***)	5.82E-05	6.02E-06	-1.27E-05 (*)	8.23E-06	-1.17E-06	5.36E-07
Citigroup	-2.47E-06	-4.98E-06 (*)	-7.08E-04 (**)	-1.00E-03	-1.03E-04 (***)	-1.64E-05	4.10E-06 (*)	-1.42E-06
Coca-Cola	5.33E-06 (***)	3.49E-06 (**)	-1.28E-04 (**)	6.64E-06	-2.37E-06 (**)	-3.20E-07	-1.45E-06 (**)	-9.15E-07
Intel	3.48E-06 (*)	-9.33E-06	9.83E-05 (**)	1.32E-05	1.08E-05 (*)	-2.31E-06	-1.10E-06	2.91E-06
JPMorgan	2.61E-06 (*)	-1.15E-05 (**)	1.24E-04	-1.52E-05	2.43E-06	2.22E-05 (*)	4.50E-06 (*)	1.23E-06
Merck & Co.	8.72E-06 (*)	1.33E-05 (***)	-4.28E-04 (***)	-3.24E-05	2.39E-05 (***)	2.84E-07	1.76E-06 (**)	4.07E-07
Microsoft	6.23E-07 (*)	-4.68E+01 (***)	1.10E-05 (*)	3.98E-06	-2.56E-06	-7.71E-07	-4.36E-06 (**)	8.42E-07
Procter & Gamble	1.73E-06 (*)	5.05E-06 (*)	8.58E-06	-8.06E-05 (***)	1.46E-04 (***)	-1.81E-06	1.55E-06 (**)	3.58E-07
Pfizer	1.30E-05 (**)	-1.72E-04 (***)	6.94E-06 (*)	4.37E-06	2.48E-07	5.44E-06	5.67E-07 (*)	1.11E-06
Wal-Mart	4.24E-06 (**)	5.44E-06	5.38E-05 (***)	3.30E-06	-6.00E-07	2.00E-06	4.32E-07 (*)	-5.96E-07

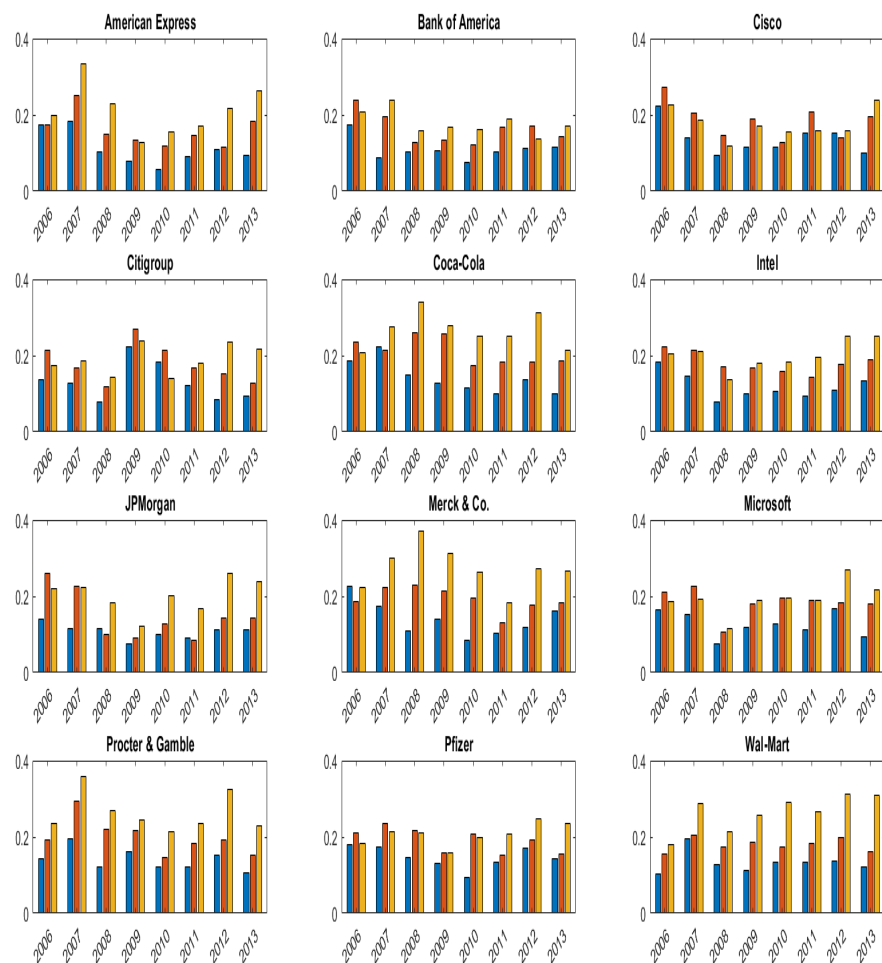
*Notes: See notes to Tables 2.8 and 2.11.

Figure 2.1: Annual Ratios of Jump Days for ETFs



*Notes: Entries in the above charts denote annual ratios of detected jump days, based on daily applications of *ASJ*, *BNS* and *PZ* fixed time span jump tests. See Sections 2.5 and 2.7 for complete details.

Figure 2.2: Annual Ratios of Jump Days for Individual Stocks



*Notes: See notes to Figure 2.1.

Chapter 3

Latent Common Volatility Factors: Capturing Elusive Predictive Accuracy Gains When Forecasting Volatility

3.1 Introduction

Accurate price volatility estimation and prediction is crucial to successful risk management and asset allocation. In light of this fact, many new estimators relevant for volatility analysis have recently been introduced, including but not limited to, realized variance (RV) (Andersen et al. (2001)), jump robust RV based on multi-power variation and truncation (Barndorff-Nielsen and Shephard (2004), Mancini (2009), Corsi et al. (2010), Podolskij and Ziggel (2010)), and multi-scale (Aït-Sahalia et al. (2011)) and pre-averaging (Jacod et al. (2009)) estimators, which are designed to eliminate microstructure effects. Making use of these sorts of integrated volatility estimators, heterogeneous autoregressive (HAR) type forecasting models have been studied extensively in the financial econometrics literature. For example, Corsi (2009) introduces a basic HAR-RV model, and Andersen et al. (2007a) and Corsi et al. (2010) analyze jump variation augmented HAR-RV models. Additionally, Duong and Swanson (2015) examine HAR model performance when so-called upside and downside jump variations are included. The authors also utilize q -th order variations of jump components, with $0.1 \leq q \leq 6$, and consider the usefulness of large jumps (i.e., jump size exceeds a given threshold) in HAR-RV type regressions. Patton and Sheppard (2015) study how the positive and negative price jumps affect the future volatility, respectively. Audrino and Hu (2016) exploit the prevalence of the leverage effect and investigate the characteristics of different components of continuous risks and jump risks on volatility persistence. Bollerslev et al. (2016) further improve volatility forecasting by allowing for the change of model coefficients according to the degree of measurement error. Other types of volatility forecasting model are also widely used, such as stochastic volatility (SV) models (Meddahi (2001), Andersen et al. (2004), Andersen et al. (2011)), (G)ARCH-type

models (Andersen et al. (2003), Hansen and Lunde (2005), Brandt and Jones (2006)), and Mixed Data Sampling (MIDAS) models (Ghysels et al. (2006), Ghysels and Sinko (2011)).

Although very parsimonious, the HAR-type models discussed in the above papers only utilize information on the target asset that is being predicted. A little explored question is whether there are sources of information other than the target asset itself can help improve the predictive accuracy in HAR regressions. In this paper, we attempt to answer this question by augmenting benchmark HAR models with estimates of latent integrated volatility (IV) factors extracted from latent common asset return factors, which are themselves extracted from a large dimensional and high-frequency asset returns dataset, and investigating whether these latent IV estimates are informative about the future volatility of selected target assets. As shall be discussed below, inclusion of latent IV factors substantially improve volatility forecasting performance for various assets at market, sector and individual-stock levels, with the notable exception of the financial sector.

The dimension reduction approach that we use in order to estimate factors combines several cutting-edge methods widely used in the literature. In particular, the motivation for our two-step dimension reduction procedure is based on new results on the use of principal component analysis (PCA) in the construction of latent factors using large dimensional and high-frequency asset return datasets that are developed in Aït-Sahalia and Xiu (2017a) and Aït-Sahalia and Xiu (2017b). However, in addition to focusing on PCA, our procedure attempts to take account of the fact that we are interested in targeted or individual market, sector or stock return prediction. Such targeted prediction is potentially inconsistent with the direct use of principal component analysis (PCA) for the extraction of common factors, since the common factors estimated using PCA that are used in prediction models are usually those associated with the largest eigenvalues in an eigenvalue-eigenvector decomposition of the correlation matrix of the dataset being examined (e.g., see Stock and Watson (2002a,b, 2006), Bai and Ng (2006a,b, 2008), and the references cited therein). Namely, the factors that account for the largest share of the variability of the covariance (correlation) matrix are assumed to be the best

candidate predictors for a given target variable. Clearly, this may not always be the case, as discussed in Bai and Ng (2008), Carrasco and Rossi (2016), and Swanson and Xiong (2017). To address this problem, we begin, in a first step, by selecting a subset of assets from the total asset pool. This is done by carrying out shrinkage of the set of all integrated volatility estimates for the asset return variables in our dataset. Shrinkage is done using the least absolute shrinkage operator (LASSO) or the elastic net. Then, in a second step, we estimate latent asset return factors by applying either PCA or sparse PCA (SPCA) to the selected subset of asset return variables corresponding to the integrated volatility variables selected in our first step. Finally, these latent asset return factors are used to construct latent integrated volatility factors, which are in turn used as explanatory variables in our HAR-type regression model prediction experiments.

One important aspect of our investigation is our novel use of SPCA. While PCA is well known, sparse principal component analysis (SPCA) is relatively new to the field, as discussed in Kim and Swanson (2017). Intuitively, SPCA can be viewed as a form of “double” shrinkage (see Zou et al. (2006) and Qi et al. (2013)). More specifically, while PCA can be interpreted as penalized regression with an L-2 penalty (akin to the penalty used in ridge regression), SPCA can be interpreted as penalized regression with either an L-1 norm penalty (i.e., a LASSO variant of PCA), or a combined L-1 and L-2 norm penalty (i.e., an elastic net variant of PCA). In both cases, sparseness is imposed on the factor loadings, with a regularization parameter controlling the degree of sparseness. In our setup, thus, sparseness is first imposed in our variable selection step (i.e., in our first step, where the lasso and elastic net are used to analyze integrated volatility variables), and then again imposed in our latent factor construction step (i.e., in our second step, where PCA and SPCA are used to analyze high frequency asset returns). Broadly speaking, the first step of our approach follows, and builds on, methods developed in in Bai and Ng (2008) in which “targeted predictors” are selected before the estimation of common factors. Again broadly speaking, our second step follows, and builds on, methods developed in Aït-Sahalia and Xiu (2017a) and Aït-Sahalia and Xiu (2017b), in which latent integrated volatility variables are constructed using PCA.

Our dataset consists of intra-day observations on 267 constituents of the S&P 500

index, 9 sector ETFs, and one market EFT (i.e., SPY, which is the SPDR S&P 500 ETF). Data were analyzed for the sample period from January 3, 2006 to December 31, 2010, and were collected from the TAQ database. We report the results based on prediction of SPY, 9 sector ETFs, and 11 individual stocks, for the period of July 1, 2009 to December 31, 2010. We also report the *in-sample* fit of various forecasting models, common factor estimators, and data aggregation permutations. Our key findings are summarized below, and explained in detail in a later section of the paper.

First, *in-sample* fit is surprisingly stable across different models, including our benchmark HAR model and our volatility-factor augmented models, across three different data frequencies, including 1-minute, 5-minute, and 10-minute frequencies. Thus, there is little to choose between data frequencies when comparing *in-sample* model fit. Moreover, *in-sample* model fit is surprisingly similar across different asset classes (i.e., market index, sector ETFs, and individual stocks), with most R^2 values ranging rather tightly between 0.35 and 0.55.

Second, our *in-sample* findings are highly mis-leading, when the objective of interest is *out-of-sample* volatility prediction. Namely, all of the above findings become irrelevant when ex ante prediction experiments are carried out. In particular, for forecasting, data frequency is crucial, and the “best” frequency varies across different assets and asset classes. However, we still recommend using the 5-minute frequency, as a general rule-of-thumb. This is because our factor augmented HAR models generally yield the “best” predictions (see below for further discussion) using 5-minute frequency data, when comparing results factor augmented model predictive accuracy across different frequencies. Intuitively, note that on one hand, using higher frequency data may result in a substantial amount of microstructure noise being absorbed by extracted factors, hence potentially deteriorating predictive performance. On the other hand, if the sampling frequency is relatively low, it is more difficult to eliminate individual jumps when estimating latent factors, leading to forecast deterioration.

The above argument is buttressed by our finding that models utilizing SPCA in factor construction generally forecast “better” than those utilizing PCA. Moreover, the

performance of SPCA, relative to PCA, is greatest when one moves from using 10-minute to 5-minute frequency data, as well as when one moves from using 1-minute to 5-minute frequency data.

Third, and perhaps most importantly, predictive accuracy improves appreciably when latent common volatility factors are included in benchmark HAR-type models. For example, for Johnson & Johnson (see Table 3.15), the benchmark model using 5-minute frequency data achieves an *out-of-sample* R^2 value of only 0.14. This is approximately one-third of the *out-of-sample* R^2 value associated with our “best” factor-augmented model. This pattern occurs for many firms and sectors; as well as for the market ETF. Interestingly, if only *in-sample* R^2 values were examined in order to assess the usefulness of common factors, then the story would change markedly. For example, again using Johnson & Johnson to illustrate our findings, the benchmark model using 5-minute frequency data (without a common factor) achieves an *in-sample* R^2 value of 0.39, while *in-sample* R^2 values for our factor-augmented models are all between 0.43 and 0.48. This small increase associated with utilizing common factors in an *in-sample* context characterizes all of our experiments. Indeed, substantial increases in performance only arise when using latent factors for ex ante prediction. This finding constitutes strong evidence of an important difference between findings based on in- and *out-of-sample* experiments.

A different way to interpret the above key finding is as follows. *In-sample* R^2 values are widely known to be substantively greater than out of sample R^2 values in financial forecasting applications. This feature has been extensively discussed in the literature, and reasons for it range from the presence of (smooth) structural breaks and state transitions, to the general inability of linear models to capture inherently nonlinear interactions among financial variables and markets (e.g., see Paye and Timmermann (2006), Aiolfi et al. (2009), and Ang and Timmermann (2012)). In our experiments, when comparing benchmark HAR models, *in-sample* R^2 values are indeed much greater than their *out-of-sample* benchmark HAR counterparts, as might be expected. For example, using IBM (see the 5-minute panel in Table 3.14) to illustrate our findings, the benchmark model (without a common factor) achieves an *in-sample* R^2 value of

0.61, as opposed to an *out-of-sample* R^2 value of 0.24. However, when the “best” factor augmented *in-sample* and out-of sample performances are compared in this example, the R^2 values are 0.65 and 0.38, respectively. Thus, the relative *out-of-sample* gains associated with utilizing latent volatility factors are greater than the *in-sample* gains. This feature characterizes our results at all market, sector, and individual-stock levels, although it is more starkly apparent at the individual stock level.

Fourth, there is an important wrinkle to the above story. Namely, for financial assets, out-of sample R^2 values are approximately 0 in some cases. A particularly interesting example of this is the financial sector ETF. For this ETF, *in-sample* R^2 values range from around 0.53 to 0.64, while *out-of-sample* R^2 range from around 0.08 to 0.30. At the individual stock level, the picture is even more stark. Consider Goldman Sachs (see Table 3.13). *In-sample* R^2 values are always around 0.40, while *out-of-sample* R^2 values are always less than 0. However, all is not lost. As discussed above, for many of our target variables, there is substantial predictable content. For example, *out-of-sample* R^2 values for Coca-Cola (see Table 3.17), Exxon Mobil (see Table 3.22) and IBM (see Table 3.14) range from 0.35 to 0.41, from 0.30 to 0.37, and from 0.23 to 0.38, respectively, when using common factors constructed via our two-step procedure, and based on IV estimators constructed using 5-minute frequency data.

Fifth, financial stocks are frequently selected in our first variable selection (or shrinkage) step. However, they are often assigned small weights in the second step (i.e., the latent factor estimation step), particularly when SPCA is used in this step. For instance, when we forecast the volatility of our energy sector ETF using 1-minute frequency data, over 33% of the most frequently selected stocks in the first step are in financial sector. However, the average weight assigned by PCA to, for instance, Goldman Sachs is only around 0.09, while the corresponding weight assigned to Texas Instruments is around double that (see Table 3.24). Even more starkly, the average weight assigned by SPCA to Goldman Sachs drops is only around 0.02. This is in part due to the fact that over 50% of weights assigned by SPCA are identically zero. On the contrary, the average weight on Texas Instruments Incorporated rises to 0.19. Therefore, we conjecture that the contribution of financial stocks to common volatility factors may be less than that

of stocks in other sectors, based on these rather surprising findings. Moreover, and as a result of the above findings, it is very likely that the marginal predictive content of common volatility factors is largely accounted for by information in sectors other than the financial sector, such as the industrial and technology sectors.

The rest of the paper is organized as follows. Section 3.2 outlines our setup and modeling assumptions, and includes a brief discussion of some of the realized measures that we construct. Section 3.3 discusses the forecasting framework used, and briefly introduces PCA, SPCA, LASSO and elastic net methods. Section 3.4 includes a discussion of the data used in our forecasting experiments, and summarizes our key empirical findings. Finally, Section 3.5 contains concluding remarks.

3.2 Setup

Denote by X the d -dimensional log-price process of d assets. Following the high-frequency literature, we assume that X follows an Itô-semimartingale defined on some filtered probability space $(\Omega, \mathbb{F}, (\mathcal{F}_t)_{t \geq 0}, \mathbb{P})$, and has the following representation:

$$\begin{aligned} X_t = X_0 &+ \int_0^t b_s ds + \int_0^t \sigma_s dW_s \\ &+ \int_0^t \int_{\{|x| \leq \epsilon\}} x(\mu - \nu)(ds, dx) + \int_0^t \int_{\{|x| \geq \epsilon\}} x\mu(ds, dx), \end{aligned} \quad (3.1)$$

where b_t is the instantaneous drift term, σ_t is the spot volatility, and both are adapted and càdlàg. Additionally, W_t is a multidimensional standard Brownian motion, μ is a Poisson random measure with compensator ν , and $\epsilon > 0$ is an arbitrary number. For more details on Itô-semimartingale and continuous-time asset price modeling, see Aït-Sahalia and Jacod (2014) and the references therein.

Since volatility is unobservable, realized measures are often employed to consistently estimate it on a fixed interval $[0, T]$, using high-frequency intraday data. For instance, one of the most widely known measures, realized volatility, is defined as follows:

$$\text{RV}_t = \sum_{i=1}^{\lfloor t/\Delta_n \rfloor} (\Delta_i^n X)^2, \quad \forall t \in [0, T], \quad (3.2)$$

where $\lfloor m \rfloor$ is the integer part of m and $\Delta_i^n X = X_{i\Delta_n} - X_{(i-1)\Delta_n}$, where Δ_n is the

equally-spaced sampling interval that shrinks to zero. It is well-known that when asset prices are continuous on a fixed interval $[0, T]$, we have that:

$$\sum_{i=1}^{\lfloor t/\Delta_n \rfloor} (\Delta_i^n X)^2 \xrightarrow{\mathbb{P}} \int_0^t \sigma_s^2 ds, \quad \forall t \in [0, T]. \quad (3.3)$$

However, when asset prices are discontinuous on $[0, T]$:

$$\sum_{i=1}^{\lfloor t/\Delta_n \rfloor} (\Delta_i^n X)^2 \xrightarrow{\mathbb{P}} \int_0^t \sigma_s^2 ds + \sum_{0 \leq s \leq t} (\Delta X_s)^2, \quad \forall t \in [0, T]. \quad (3.4)$$

where $\Delta X_s := X_s - X_{s-} \neq 0$, if and only if X jumps at time s .

To separate the integrated volatility from jump variation, one can use the threshold technique developed in Mancini (2001, 2009):

$$\sum_{i=1}^{\lfloor t/\Delta_n \rfloor} (\Delta_i^n X)^2 \mathbf{1}_{\{|\Delta_i^n X| \leq \alpha \Delta_n^\varpi\}} \xrightarrow{\mathbb{P}} \int_0^t \sigma_s^2 ds, \quad (3.5)$$

or use the multipower variation (MPV) estimator developed in Barndorff-Nielsen and Shephard (2004) and Barndorff-Nielsen et al. (2006):

$$\Delta_n^{1-p^+/2} \sum_{i=1}^{\lfloor t/\Delta_n \rfloor - k + 1} |\Delta_i^n X|^{p_1} \dots |\Delta_{i+k-1}^n X|^{p_k} \xrightarrow{\mathbb{P}} m_{p_1} \dots m_{p_k} \int_0^t |\sigma_s|^{p^+} ds \quad (3.6)$$

where $p_j \geq 0$, $p^+ = p_1 + \dots + p_k$ and $m_p = \mathbb{E}[|\mathcal{N}(0, 1)|^p]$. One can also combine these two methods and use a truncated multipower variation estimator. Apparently, different components of the quadratic variation can be analyzed or used separately in econometric analysis.

We also assume that the continuous part of asset log-prices follows an underlying continuous-time factor model on $[0, T]$. Namely, define:

$$Y_t = \Lambda_t F_t + Z_t \quad (3.7)$$

where $Y_t := X_0 + \int_0^t b_s ds + \int_0^t \sigma_s dW_s$ is the continuous part of X , F_t is an r -dimensional continuous factor ($r < d$), Z_t is an idiosyncratic component, and Λ_t is a d -by- r factor loading matrix, each element of which is adapted and has càdlàg paths almost surely. Here, we specifically call F_t the common price factor in order to distinguish it from the common volatility factor defined later. The common price factor F 's and the

idiosyncratic component Z 's are continuous Itô-semimartingales as well, with:

$$F_t = F_0 + \int_0^t h_s ds + \int_0^t \eta_s dB_s \quad (3.8)$$

and

$$Z_t = Z_0 + \int_0^t g_s ds + \int_0^t \gamma_s d\tilde{B}_s, \quad (3.9)$$

where B_s and $\tilde{B}_s$ are independent Brownian motions. All of the coefficient processes, h , η , g and γ are adapted to $(\mathcal{F}_t)_{t \geq 0}$ and have càdlàg paths, almost surely. The above factor models and general settings follow Aït-Sahalia and Xiu (2017b).

3.3 Dimension Reduction and Forecasting Methods

The original HAR model is given below.

$$\text{RM}_{t+h} = \beta_0 + \beta_1 \text{RM}_t + \beta_2 \text{RM}_{[t,t-4]} + \beta_3 \text{RM}_{[t,t-21]} + \epsilon_t, \quad (3.10)$$

where RM's are realized measures of integrated volatility, and $\text{RM}_{[t,t-p]}$ is the average of RM's over the most recent $p + 1$ days. For instance, if realized volatility is used in the model, then:

$$\text{RV}_{[t,t-p]} = \frac{1}{p+1} \sum_{i=0}^p \text{RV}_{t-i}. \quad (3.11)$$

To eliminate the jump variation from the total quadratic variation, we employ the truncated realized volatility in (3.5) to consistently estimate the integrated volatility¹. Therefore, the benchmark model that we consider in this paper is as follows:

$$\text{TRV}_{t+h} = \beta_0 + \beta_1 \text{TRV}_t + \beta_2 \text{TRV}_{[t,t-4]} + \beta_3 \text{TRV}_{[t,t-21]} + \epsilon_t, \quad (3.12)$$

where TRV stands for truncated realized volatility.

We propose using the following factor-augmented model in our forecasting experiments,

$$y_{t+h} = \beta_0 + \beta_\Psi^\top \Psi_t + \beta_w^\top w_t + \varepsilon_t, \quad (3.13)$$

¹We actually combine the two methods, i.e. (3.5) and (3.6), in the following way: we first use bipower variation to get an initial consistent estimate of the integrated volatility, and then use this to determine an initial choice for α . Then, we obtain a second estimate of the integrated volatility using the truncation method, and a second choice of α . We iterate this procedure until the estimated integrated volatility converges.

where y_{t+h} is the h -step-ahead forecast of daily integrated volatility. We focus on one-day-ahead forecasts (i.e., $h = 1$). Here, w_t is a vector consisting of truncated realized volatility on day t , the weekly average of truncated realized volatility from days $t - 4$ to t , and the monthly average of truncated realized volatility from days $t - 21$ to t (i.e., w_t contains all predictors in the benchmark model). Furthermore, Ψ_t consists of r -dimensional unobservable predictors. Based on the structure of factors assumed in (3.7), we define

$$\Psi_t := \int_0^t \text{diag}(\Lambda_s \eta_s \eta_s^\top \Lambda_s^\top) ds$$

and name it the common volatility factor. Note that we can not disentangle Λ from η unless imposing certain identification condition such as $\eta \eta^\top = I_r$. So we don't distinguish them and treat Ψ_t as the integrated volatility matrix of the r uncorrelated common factors.

Here, common price factors are extracted using PCA or SPCA applied to a high-frequency dataset, the constituent members of which are specified using LASSO or elastic net shrinkage on our 274 variable original dataset. Intuitively, common price factors in (3.7) can be interpreted as “composite stocks” (the name comes from the fact that they are linear combinations of all individual stocks in the dataset) that in general affect a majority of stocks in the market. Therefore, we first construct those “composite stocks”, next estimate the integrated volatility for each, and finally use the estimated integrated volatilities as predictors in (3.13) to forecast the integrated volatility of the target asset. Of note is that unlike many other applications of factor-augmented regressions, we do not directly use common factors $\Lambda_t F_t$ extracted from a large number of assets. Instead, what we actually use as predictors in forecasting models are the estimated integrated volatilities of these common factors, i.e. Ψ_t . As discussed above, we use PCA and SPCA when constructing “composite stocks” in this paper. These dimension reduction methods will be briefly discussed after we summarize the shrinkage methods utilized in the first step of our two step volatility factor extraction procedure.

3.3.1 LASSO and Elastic Net

Prior to construction of latent factors using PCA and SPCA, we first select targeted predictor assets. For this, we use two shrinkage or variable selection methods, including the LASSO (see Tibshirani (1996)) and the elastic net (see Zou and Hastie (2005)). Both techniques can be interpreted as regularized or penalized regression methods. Briefly, let RSS be the sum of squared residuals from a regression of y_{t+h} on w_t and χ_t , where χ_t is a vector of estimates of integrated volatility on day t for all assets in X_t . The LASSO estimator is the solution to:

$$\min_{\phi} \text{RSS} + \lambda \sum_j |\phi_j|, \quad (3.14)$$

where the ϕ 's are coefficients in the regression. Only assets with nonzero ϕ 's are retained in our final set of selected target predictor assets, say $\tilde{X}_t$, and the sparsity (number of variables) in $\tilde{X}_t$ only depends on λ . Therefore, instead of X_t , we actually apply PCA or SPCA to the variance-covariance matrix of $\tilde{X}_t$ when constructing estimates of latent asset return factors that are in turn used to construct latent volatility factors.

Similarly, the elastic net estimator is the solution to:

$$\min_{\phi} \text{RSS} + \lambda \sum_j \left(\frac{(1-\alpha)}{2} \phi_j^2 + \alpha |\phi_j| \right), \quad (3.15)$$

with $\alpha \in [0, 1]$. Of note is that when $\alpha = 1$, the elastic net is equivalent to LASSO. As α shrinks toward 0, the elastic net approaches ridge regression. In our experiments, we set $\alpha = 0.2$ and 0.6 for the elastic net. For both the LASSO and the elastic net, λ is selected using 10-fold cross validation in the training dataset used to calibrate prediction models.

Note that for any two different target assets, the information pool ($\tilde{X}_t$) from which we construct the $\hat{F}_t$'s and subsequently the $\hat{\Psi}_t$'s can be quite different (though the probability of them being equivalent is still positive). Additionally, note that after selecting $\tilde{X}_t$ via LASSO or elastic net shrinkage targeted to a specific asset, we construct (sparsely loaded) latent factors that are specifically related to the asset of interest. Therefore, it is reasonable to assume that their integrated volatilities (i.e., the $\hat{\Psi}_t$'s)

forecast y . In both the case of PCA and SPCA, the latent integrated volatility variables used in the HAR regressions discussed below are estimated using high frequency uncorrelated latent asset return factor data.

We conclude this subsection with two remarks. First, the above PCA procedure delivers the eigens (eigenvalues and eigenvectors) of the integrated volatility matrix. According to Aït-Sahalia and Xiu (2017a), these eigens are different from the integrated eigens of the spot volatility matrix, when t does not shrink to zero. However, in finite samples (e.g., in our empirical application), the time horizon t (one day) is small relative to Δ_n (1-minute, 5-minute, and 10-minute), and this difference is small compared with other sources of estimation error. Therefore, we do not address eigens of integrated volatility versus integrated eigens of spot volatility differences in our empirical application, following the approach taken by Aït-Sahalia and Xiu (2017b).

Second, it is well-known that eigens are nonlinear functions of the corresponding data matrix. Jacod and Rosenbaum (2013) show that various bias terms arise when estimating integrals of nonlinear functions of the spot volatility matrix, although only one bias term remains when local window sizes that are used are chosen to be relatively small. They further demonstrate that this remaining bias can be consistently estimated. Hence, it is possible to construct bias-corrected estimators. Moreover, according to Aït-Sahalia and Xiu (2017a), these bias terms are proportional to their associated eigens. Consequently, they share the same source of predictive power as eigens. In addition, analogous to our earlier arguments, the ratio t/Δ_n is small in our empirical application, making the bias term that can be treated using the methods of Jacod and Rosenbaum (2013), which is an integral over $[0, t]$, relatively small compared with other estimation errors. In view of these observations, we don't remove this bias term in our empirical application.

3.3.3 Sparse Principal Component Analysis

In general, PCA yields nonzero factor loadings for (almost) all variables, which exacerbates difficulty in interpretation. To avoid this drawback of PCA, and to generally induce parsimony, we also consider estimating factors using SPCA, as developed by Zou

et al. (2006) and Qi et al. (2013).

Let $\hat{\Sigma}$ be the same covariance matrix estimator defined in (3.16). The eigenvector $\hat{\xi}_1$ of the first sparse principal component is the solution to:

$$\max_{\|\xi_1\|_2=1} \frac{\xi_1^T \hat{\Sigma} \xi_1}{\|\xi_1\|_{\lambda_1}^2}, \quad (3.19)$$

where $\|\cdot\|_{\lambda_1}$ is a mixed norm defined as $\sqrt{(1-\lambda_1)\|\cdot\|_2^2 + \lambda_1\|\cdot\|_1^2}$, with $\lambda_1 \in [0, 1]$. Note that if $\lambda_1 = 0$, this mixed norm is equivalent to the L-2 norm, while it is equivalent to the L-1 norm if $\lambda_1 = 1$. With $\hat{\xi}_1$, one can sequentially obtain subsequent eigenvectors by solving the following optimization problems for $j = 2, 3, \dots, r$:

$$\max_{\|\xi_j\|_2=1, \xi_{j-1} \perp \xi_j} \frac{\xi_j^T \hat{\Sigma} \xi_j}{\|\xi_j\|_{\lambda_j}^2}, \quad (3.20)$$

where λ_k is the tuning parameter for ξ_k (which might be different for each k). In short, SPCA produces “sparse” factor loadings in the sense that many of them are identically zero, while factors are still constructed in the spirit of PCA, since explained data variances are maximized under constraints. Qi et al. (2013) show that their proposed algorithm for optimizing objective functions yields a stable limit which consistently estimates the eigenvectors under certain conditions. Our approach, as discussed above, is to first estimate high-frequency sparse principal components, and then construct and utilize realized volatilities from these estimated factors in (3.13) to forecast any given target of interest.

3.3.4 Forecasting Methods

The proposed one-step forecasting model is:

$$\widehat{\text{TRV}}_{t+1} = \beta_0 + \beta_1 \widehat{\text{TRV}}_t + \beta_2 \widehat{\text{TRV}}_{[t, t-4]} + \beta_3 \widehat{\text{TRV}}_{[t, t-21]} + \beta_\Psi^T \hat{\Psi}_t + \epsilon_t. \quad (3.21)$$

The estimated factors’ volatilities, $\hat{\Psi}_t$, are constructed by implementing the above mentioned two-step procedure. Recall that the first step involves using LASSO or elastic net shrinkage to select a subset of the asset dataset, as outlined in Section 3.3.1. The second step involves latent volatility factor construction, as discussed in Sections 3.3.2 and 3.3.3.

We choose the number of latent factors in our experiments by following an easy-to-implement, albeit ad-hoc rule. First, we sort all eigenvalues in descending order and select (additional) principal components based on their corresponding eigenvalues until their cumulative contribution exceeds (or is equal to) 90% of the total variation of the dataset. Next, we discard principal components with individual contributions that are less than 5% of total variation. For instance, if the first 5 principal components contribute 60%, 10%, 10%, 6%, 4%, respectively, we keep the first 4 principal components. The idea is very simple and natural: there is a trade-off between a more parsimonious model and a less informative one. Although the choice of cutoffs is somewhat arbitrary, our experiments suggest that the findings are robust to other cutoffs within a reasonable range of the above ones. Finally, we estimate daily integrated volatility of selected latent factors and use them as predictors in (3.21).

In summary, we consider six “permutations” of our two-step procedure in forecasting experiments, as follows:

I. EN1-PCA: First step - assets selected using elastic net (EN) shrinkage, with parameter $\alpha = 0.2$. Second step - latent integrated volatility factors constructed using PCA.

II. EN2-PCA: First step - assets selected using elastic net (EN) shrinkage, with parameter $\alpha = 0.6$. Second step - latent integrated volatility factors constructed using PCA.

III. LASSO-PCA: First step - assets selected using LASSO shrinkage, with parameter $\alpha = 0.2$. Second step - latent integrated volatility factors constructed using PCA.

IV. EN1-SPCA: First step - assets selected using elastic net (EN) shrinkage, with parameter $\alpha = 0.2$. Second step - latent integrated volatility factors constructed using SPCA.

V. EN2-SPCA: First step - assets selected using elastic net (EN) shrinkage, with parameter $\alpha = 0.6$. Second step - latent integrated volatility factors constructed using SPCA.

VI. LASSO-SPCA: First step - assets selected using LASSO shrinkage, with parameter $\alpha = 0.2$. Second step - latent integrated volatility factors constructed using

SPCA.

Model estimation and volatility prediction are carried out anew, each day, using a rolling-window estimation scheme. The length of rolling window (i.e. the *in-sample* period), is 630 days. For example, we first estimate models using data from December 28, 2006 to June 30, 2009 (630 trading days), and then construct one-day-ahead forecasts for July 1, 2009. Then, in order to forecast the volatility on July 2, 2009, we first estimate our models using data from December 29, 2006 to July 1, 2009 (630 trading days). We continue this procedure until we reach the end of our dataset. Finally, we obtain sequences of daily *out-of-sample* volatility forecasts for the sample period from July 1, 2009 to December 31, 2010, which constitutes 380 trading days.

Our benchmark HAR model is estimated using ordinary least squares. All factor-augmented regressions are estimated using constrained least squares, in order to guarantee that all parameters are nonnegative. By doing so, we avoid any potential negative forecasts of volatility.

To evaluate the forecasting performance of our factor-augmented models and compare them with the benchmark model, we consider three different criteria:

(a) *In-sample* R^2 .

(b) *Out-of-sample* R^2 (Campbell and Thompson (2008)), defined as:

$$R_{\text{OOS}}^2 = 1 - \frac{\sum_{t=1}^T (y_t - \hat{y}_t)^2}{\sum_{t=1}^T (y_t - \bar{y}_t)^2}, \quad (3.22)$$

where y_t is the ex-post value of volatility, $\bar{y}_t$ is the historical average of volatility, and $\hat{y}_t$ is our forecast.

(c) Heteroskedasticity adjusted root mean square error (HARMSE) (Corsi et al. (2010)), defined as:

$$\text{HARMSE} = \sqrt{\frac{1}{T} \sum_{t=1}^T \left(\frac{y_t - \hat{y}_t}{y_t} \right)^2} \quad (3.23)$$

The experimental setup discussed in this section is summarized in Table 3.1.

3.4 Empirical Results

3.4.1 Data

We collect intraday observations on 267 constituents of the S&P 500 index²; 9 sector ETFs, including; Materials (XLB), Energy (XLE), Financial (XLF), Industrial (XLI), Technology (XLK), Consumer Staples (XLP), Utilities (XLU), Health Care (XLV), and Consumer Discretionary (XLY); and the SPDR S&P 500 ETF (SPY). Our sample period is January 3, 2006 to December 31, 2010, and data are collected from the TAQ database.

In our forecasting experiments, target assets include SPY; the 9 sector ETFs listed above; and 11 individual stocks, including: Coca-Cola Company (KO), Exxon Mobil Corporation (XOM), General Electric Company (GE), Goldman Sachs (GS), International Business Machines (IBM), Johnson & Johnson (JNJ), JPMorgan Chase (JPM), McDonald’s (MCD), Merck (MRK), Microsoft (MSFT) and Wal-Mart (WMT).

It is worth mentioning that the original dataset we collected consists of 274 constituents of S&P 500 index. Of these, seven stocks, including AIG, C, F, GNW, HIG, LVL and STT, are deleted, leaving 267 stocks. The reason for this is that these stocks generate a small number of extreme integrated volatility values, even when data are filtered with using a judiciously chosen jump threshold. These stocks are thus viewed as “outliers” that contains strong microstructure noises and/or recording error, which are not informative about future volatilities, hence may consequently deteriorate forecasting performance of our models. As a robustness check, however, we did compare empirical results based on 267 constituents with those based on 274 constituents, although comparable results are only shown from the SPY case. Complete results based on the original dataset of 274 constituents are available upon request, although it is clear, upon comparison of our results in these two cases, that utilizing the 7 additional stocks result in a deterioration of the predictive performance of out latent volatility factors.

²Since the constituents of S&P 500 index change over time, we only collect those that are always present in the index between 2006 to 2010.

Finally, data cleaning, subsampling, etc., all follow standard procedures described in Aït-Sahalia and Jacod (2012). Overnight returns are excluded. Less frequently traded stocks are also excluded from the dataset since they do not generate high-frequency data.

3.4.2 Empirical Findings: Forecasting Performance

Tables 3.2–3.22 show the one-day ahead forecast performance of the benchmark HAR model and various factor-augmented HAR models, for the forecasting sample period from July 1, 2009 to December 31, 2010. All tables report *in-sample* and *out-of-sample* R^2 values, as well as HARMSE values. Table 3.2 (SPY) also compares the results with and without the aforementioned seven “outlier” stocks (first and second columns under each criterion). Moreover, to compare the performance across different sampling frequencies, we construct factors using 1-minute, 5-minute, and 10-minute frequency data, respectively. Finally, as discussed above, forecasting experiments are carried out using rolling windows to estimate all models, prior to ex ante forecast construction at each point in time. A number of clear-cut conclusions emerge upon inspection of the results contained in these tables.

First, *in-sample* fit is surprisingly stable across different models, including our benchmark HAR model and our volatility-factor augmented models, across three different data frequencies, including 1-minute, 5-minute, and 10-minute frequencies. Thus, there is little to choose between data frequencies when comparing *in-sample* model fit. Moreover, *in-sample* model fit is surprisingly similar across different asset classes (i.e., market index, sector ETFs, and individual stocks), with most R^2 values ranging rather tightly between 0.35 and 0.55. More specifically, most *in-sample* R^2 values for sector ETFs range rather tightly between approximately 0.50 and 0.65, regardless of whether our HAR specifications include a latent volatility factor or not. The exception to this appears to be XLP (Consumer Staples, Table 3.8), for which values range from 0.38 to 0.50. The market ETF (SPY, Table 3.2) delivers *in-sample* R^2 values between approximately 0.55 and 0.65. Finally, for individual stocks, the range is somewhat wider, including values from 0.35 to 0.65. Finally, *in-sample* fit changes little when volatility

factors are added to benchmark HAR models, regardless of asset class. Thus, based solely on *in-sample* diagnostics, there appears to be little gain to deploying volatility factors in HAR analysis. However, we shall see that this finding changes dramatically when *out-of-sample*, or true ex ante forecasting, is carried out.

Second, our *in-sample* findings are highly mis-leading, when the objective of interest is *out-of-sample* volatility prediction. Namely, all of the above findings become irrelevant when ex ante prediction experiments are carried out. For example, for forecasting, data frequency is crucial, and the “best” frequency varies across different assets and asset classes. However, we still recommend using the 5-minute frequency, as a general rule-of-thumb. This is because our factor augmented HAR models generally yield the “best” predictions (see below for further discussion) using 5-minute frequency data, when comparing results factor augmented model predictive accuracy across different frequencies. Intuitively, note that on one hand, using higher frequency data may result in a substantial amount of microstructure noise being absorbed by extracted factors, hence potentially deteriorating predictive performance. On the other hand, if the sampling frequency is relatively low, it is more difficult to eliminate individual jumps when estimating latent factors, leading to forecast deterioration.

Third, note that the the above findings are based on a comparison of predictions made using factor augmented HAR models. This is the correct comparison to make because predictive accuracy improves appreciably when latent common volatility factors are included in our benchmark HAR-type model. For example, for Johnson & Johnson (see Table 3.15), the benchmark model using 5-minute frequency data achieves an *out-of-sample* R^2 value of only 0.14. This is approximately one-third of the *out-of-sample* R^2 value associated with our “best” factor-augmented model. This pattern occurs for many firms and sectors; as well as for the market ETF. Interestingly, if only *in-sample* R^2 values were examined in order to assess the usefulness of common factors, then the story would change markedly. For example, again using Johnson & Johnson to illustrate our findings, the benchmark model using 5-minute frequency data (without a common factor) achieves an *in-sample* R^2 value of 0.39, while *in-sample* R^2 values for our factor-augmented models are all between 0.43 and 0.48. This small increase

associated with utilizing common factors in an *in-sample* context characterizes all of our experiments. Indeed, substantial increases in performance only arise when using latent factors for ex ante prediction. As discussed in the introduction to this paper, this finding constitutes strong evidence of an important difference between findings based on in- and *out-of-sample* experiments.

The above conclusion can perhaps best be understood by noting that *in-sample* R^2 values are widely known to be substantively greater than *out-of-sample* R^2 values in financial forecasting applications. This feature has been extensively discussed in the literature, and reasons for it range from the presence of (smooth) structural breaks and state transitions, to the general inability of linear models to capture inherently nonlinear interactions among financial variables and markets (e.g., see Ang and Timmermann (2012), Aiolfi et al. (2009), and Paye and Timmermann (2006)). Naturally, arguments centering around market efficiency may also play a role in explaining this phenomenon. Not surprisingly, then, when comparing benchmark HAR models, we find that *in-sample* R^2 values are indeed much greater than their *out-of-sample* benchmark HAR counterparts. For example, using IBM (see the 5-minute panel in Table 3.14) to illustrate our findings, the benchmark model (without a common factor) achieves an *in-sample* R^2 value of 0.61, as opposed to an *out-of-sample* R^2 value of 0.24. However, when the “best” factor augmented *in-sample* and out-of sample performances are compared in this example, the R^2 values are 0.65 and 0.38, respectively. Thus, the relative *out-of-sample* gains associated with utilizing latent volatility factors are greater than the *in-sample* gains, as the *out-of-sample* R^2 value increases from 0.24 to 0.38, which is more than a 50% gain. Indeed, analogous predictive accuracy gains exceed 50% for GE, JNJ, JPM, KO, MCD, MRK, WMT, and XOM (see Tables 3.12, 3.15, 3.16, 3.17, 3.18, 3.19, 3.21 and 3.22, respectively), with 5-minute frequency data. Lesser gains arise for only 2 of 11 stocks that we analyze. Broadly speaking, this feature also characterizes our results at all market and sector levels, although it is more starkly apparent at the individual stock level.

Fourth, models utilizing SPCA in factor construction generally forecast “better” than those utilizing PCA. Moreover, the gains to using SPCA, relative to PCA, are

greatest when one moves from using 10-minute to 5-minute frequency data, as well as when one moves from using 1-minute to 5-minute frequency data. This two-pronged finding is as expected, given that using high frequency data across many stocks, when constructing latent volatility factors, involves accounting for noisiness due not only to sampling frequency (i.e., microstructure noise), but also due to the large number of assets, a increasing number of which are transmitting noisy signals, as the cross sectional dimension of our dataset increases. This argument, parallels the argument outlined above, whereby using higher frequency data may result in more microstructure noise being absorbed by extracted factors, while when the sampling frequency is relatively low (or when the number of assets is relatively high), it may be more difficult to eliminate individual jumps when estimating latent factors.

Drilling down a bit further, the results in Table 3.12 indicate that at 1- and 5-minute frequencies, factor-augmented models with SPCA have a 25%–35% larger *out-of-sample* R^2 than those with PCA. Similar results can also be found in Tables 3.13, 3.14, 3.18, 3.20, 3.21 and 3.22. This pattern, however, becomes insignificant or even reversed at our lowest sampling frequency (i.e., the 10-minute frequency). Moreover, when forecasting individual stocks, as well as some ETFs, such as SPY, XLB, XLE, XLI, XLK and XLY (see Tables 3.2, 3.3, 3.4, 3.6, 3.7 and 3.11, respectively), factor-augmented models with SPCA yield much lower HARMSE, especially at when using higher frequency data. Again, this pattern becomes less significant at lower frequency. As discussed above, this finding likely due to the presence of microstructure noise in our data, given that SPCA assigns many identically zero weights on stocks, and consequently alleviates some of the effect of microstructure noise; particularly from stocks, which are non-informative about the volatility of the target asset. Therefore, we are not surprised that factor-augmented models using SPCA are more likely to perform better than those using PCA at higher frequencies. Of course, it is perhaps worth noting that due to aggregation, the impact of microstructure noise on our market index ETF and sector ETFs is much weaker. As a consequence, the difference among models utilizing SPCA and PCA when forecasting our ETFs is less pronounced, as mentioned above.

Fifth, there is an important wrinkle to the above story. Namely, for financial assets,

out-of sample R^2 values are approximately 0 in some cases. A particularly interesting example of this is the financial sector ETF. For this ETF, *in-sample* R^2 values range from around 0.53 to 0.64, while *out-of-sample* R^2 range from around 0.08 to 0.30 (see Table 3.5). At the individual stock level, the picture is even more stark. Consider Goldman Sachs (see Table 3.13). *In-sample* R^2 values are always around 0.40, while *out-of-sample* R^2 values are always less than 0. Evidently, integrated volatility of individual financial stocks is the most difficult to forecast. Unlike forecasting the financial sector as a whole, when it comes to individual financial stocks, HAR-type models performs very poorly. In Tables 3.13 and 3.16, entries in the column of *out-of-sample* R^2 for the benchmark model are almost all negative, HARMSE are in general much larger than those for other assets, and even *in-sample* R^2 values are much lower compared to other assets.

However, all is not lost. Incorporating common volatility factors extracted from a broad range of stocks into benchmark models sometimes helps in obtaining more precise forecasts for financial stocks, but only to a very limited extent. As discussed above, for many of our target variables, there is substantial predictable content. For example, *out-of-sample* R^2 values for Coca-Cola (see Table 3.17), Exxon Mobil (see Table 3.22), and IBM (see Table 3.14) range from 0.35 to 0.41, from 0.30 to 0.37, and from 0.23 to 0.38, respectively, when using common volatility factors constructed via our two-step procedure, and based on IV estimators constructed using 5-minute frequency data.

Sixth, financial stocks are frequently selected in our first variable selection (or shrinkage) step. However, they are often assigned small weights in the second step (i.e., the latent factor estimation step), particularly when SPCA is used in this step. For instance, when we forecast the volatility of our energy sector ETF using 1-minute frequency data, over 33% of the most frequently selected stocks in the first step are in financial sector. However, the average weight assigned by PCA to, for instance, Goldman Sachs is only around 0.09, while the corresponding weight assigned to Texas Instruments is around double that (see Table 3.24). Even more starkly, the average weight assigned by SPCA to Goldman Sachs drops is only around 0.02. This is in part due to the fact that over 50% of weights assigned by SPCA are identically zero. On the contrary, the average

weight on Texas Instruments Incorporated rises to 0.19. Therefore, we conjecture that the contribution of financial stocks to common volatility factors may be less than that of stocks in other sectors, based on these rather surprising findings. Moreover, and as a result of the above findings, it is very likely that the marginal predictive content of common volatility factors is largely accounted for by information in sectors other than the financial sector, such as the industrial and technology sectors.

3.4.3 Empirical Findings: Latent Factor Structures

Tables 3.23–3.25 contain factor structure details, for the case where we are interested in forecasting non-financial sector ETFs and individual stocks. A number of conclusions emerge when examining these results.

First, note that different shrinkage methods in the first step of our procedure select almost the same pool of stocks, for each sampling frequency. Thus, there appears to be little to choose between the LASSO and elastic net shrinkage. However, the pool of selected stocks changes with data frequency. For instance, consider the SPY ETF. Table 3.23 shows that at the 1-minute frequency, almost 32% of selected stocks belong to the financial sector. In contrast, at 5-minute and 10-minute frequencies, only around 15% to 20% of selected stocks are financials. Similar results can be seen upon inspection of Table 3.24 (sector ETF) and 3.25 (individual stock).

Second, an important feature of our volatility factors is that financial stocks tend to be selected frequently in the first step of our procedure, particularly when using higher frequency data. However, relatively little weight is placed on such stocks in the second step of our procedure, when utilizing PCA and SPCA to estimate asset return factors. For instance, in columns denoted “PCA” in these three tables, the average weight on HBAN (Huntington Bancshares) is only between 0.06 and 0.07, when using 1-minute frequency data. Similarly, BK (Bank of New York Mellon) has average weight around 0.06–0.09, when using 5-minute frequency data, and MMC (Marsh & McLennan Companies) in Table 3.23, GS (Goldman Sachs) in Table 3.24 and LM (Legg Mason) in Table 3.25 have average weights of around 0.1 or less, when using 10-minute frequency data. Furthermore, under “SPCA”, the average weights on financial stocks

are even smaller, and many are identically zero. For instance, Table 3.23 shows that at the 1-minute frequency, the average weight on PRU (Prudential Financial) decreases dramatically from 0.104 to 0.047 when factor estimation utilizes SPCA instead of PCA (under SPCA almost 28% of daily weights are zero). This finding is consistent with our above microstructure noise explanation of the superior performance of models that utilize SPCA, in conjunction with the use of higher frequency data.

Third, notice that stocks in the industrial and technology sectors usually have larger factor loadings (weights) under both PCA and SPCA. For instance, in Table 3.25, CSCO (Cisco), LLTC (Linear Technology) and SWKS (Skyworks Solutions) - in the technology sector, and MAS (Masco), UPS and UTX (United Technologies) - in the industrial sector, all have average weights greater than 0.15. Similarly, in Table 3.24, CERN (Cerner), NFLX (Netflix) and TXN (Texas Instruments) - in technology sector, and CSX, FAST (Fastenal) and HON (Honeywell) - in the industrial sector - have average weights larger than 0.15. Putting all of the above evidence together, we conclude that although financial stocks are frequently chosen in our first step shrinkage procedure, their contributions to common volatility factors appears to be less than that of industrial and technology stocks.

3.5 Concluding Remarks

This paper investigates whether latent common volatility factors extracted from a large-dimensional high-frequency intraday stock returns improve volatility forecasting. We propose a factor-augmented version of the widely studied HAR model. In our new model, factors are estimated using a two-step procedure involving variable selection using least absolute selection operator (LASSO) and elastic net shrinkage, followed by factor estimation using (sparse) principal components analysis (SPCA).

Our key findings are summarized as follows. First and foremost, we uncover substantial empirical evidence indicating that latent common volatility factors greatly improve the *out-of-sample* predictive accuracy of HAR models, as measured by both HARMSE

and *out-of-sample* R^2 . This improvement is seen across markets, sectors, and individual companies, with the greatest improvements noted at the individual company level. Second, *in-sample* performance is often irrelevant to *out-of-sample* performance. Indeed, if volatility modeling is viewed solely through the lens of *in-sample* fit, then little is gained by generalizing the HAR model using our procedure. Almost all gains are seen only when true ex ante prediction is carried out. Third, we recommend using high frequency datasets consisting of data sampled at 5-minute frequency, when constructing predictions of volatility using factor augmented regressions. This recommendation arises because of microstructure noise considerations, as well as because of the incidence of heterogeneous jumps associated with the large cross sectional dimension of our dataset. We also find that models utilizing SPCA perform better than those with PCA, when these methods are used to extract common volatility factors.

This chapter is meant as a starting point, as much remains to be done. For example, although substantial theoretical advances in the application of principal component analysis to high dimensional asset return datasets are made in Aït-Sahalia and Xiu (2017a) and Aït-Sahalia and Xiu (2017b), it remains to ascertain whether the results carry over to the use of SPCA. It also remains to theoretically analyze higher order latent (e.g., volatility) factors that are estimated based using first order latent factors constructed using observed (asset) data. From an empirical perspective, it will be of interest to further examine the robustness of the findings in this paper to the use of alternative sample periods for both *in-sample* estimation and out-of sample prediction. It will also be of interest to assess whether the findings in this paper can be translated into profitable investment strategies, in real-time trading contexts.

Table 3.1: Experimental Setup

Benchmark Model:	
$\widehat{\text{TRV}}_{t+1} = \beta_0 + \beta_1 \widehat{\text{TRV}}_t + \beta_2 \widehat{\text{TRV}}_{[t,t-4]} + \beta_3 \widehat{\text{TRV}}_{[t,t-21]} + \epsilon_t$	
Two-Step Procedure:	
Step 1: Shrinkage Methods (Variable Selection)	Step 2: Factor Estimation Methods
1. LASSO ($\alpha = 0$)	1. PCA
2. EN1 ($\alpha = 0.2$)	
3. EN2 ($\alpha = 0.6$)	2. SPCA
Sample Periods:	
<i>In-sample</i> period: January 3, 2006 – June 30, 2009	
<i>Out-of-sample</i> period: July 1, 2009 – December 31, 2010	
Regression Estimation Scheme:	
Rolling-window estimation.	
Window length: 630 days.	
Sampling Frequencies:	
1, 5, and 10 minutes.	
Factor Selection Rules:	
Contribution of all selected factors exceeds 90% of total variation.	
Contribution of every selected factor exceeds 5% of total variation.	
Evaluation Criteria:	
1. <i>In-sample</i> R^2	
2. <i>Out-of-sample</i> R^2	
3. Heteroskedasticity adjusted root mean square error (HARMSE)	

Table 3.2: SPDR S&P 500 ETF (SPY)

Frequency	Model	<i>In-Sample</i> R^2		<i>Out-of-Sample</i> R^2		HARMSE	
1-minute	Benchmark	0.5218	0.5218	0.2737	0.2737	1.2493	1.2493
	EN1-PCA	0.5302	0.5279	0.3030	0.3004	0.8443	1.0169
	EN2-PCA	0.5304	0.5279	0.3181	0.2823	0.8276	0.9794
	Lasso-PCA	0.5304	0.5280	0.3164	0.2985	0.8347	0.9981
	EN1-SPCA	0.5458	0.5408	0.3312	0.1822	0.6245	0.9497
	EN2-SPCA	0.5461	0.5408	0.3421	0.1626	0.6313	0.9911
	Lasso-SPCA	0.5461	0.5413	0.3197	0.1601	0.6350	0.9959
5-minute	Benchmark	0.6006	0.6006	0.3605	0.3605	1.2629	1.2629
	EN1-PCA	0.6071	0.6029	0.3897	0.3801	0.9828	1.1222
	EN2-PCA	0.6047	0.6031	0.3931	0.3780	1.0646	1.0984
	Lasso-PCA	0.6039	0.6030	0.3774	0.3759	1.0240	1.1122
	EN1-SPCA	0.6204	0.6088	0.4313	0.3995	0.7066	0.9393
	EN2-SPCA	0.6202	0.6088	0.4381	0.4000	0.7141	0.9156
	Lasso-SPCA	0.6193	0.6086	0.4233	0.4071	0.7012	0.9497
10-minute	Benchmark	0.5039	0.5039	0.2609	0.2609	1.6082	1.6082
	EN1-PCA	0.5445	0.5461	0.3829	0.3342	1.0496	1.0796
	EN2-PCA	0.5440	0.5373	0.3705	0.2729	1.0213	1.1176
	Lasso-PCA	0.5453	0.5363	0.3725	0.2810	1.0323	1.1523
	EN1-SPCA	0.5457	0.5428	0.3960	0.3239	1.0672	1.0790
	EN2-SPCA	0.5434	0.5362	0.3800	0.2816	1.0670	1.0963
	Lasso-SPCA	0.5449	0.5361	0.3833	0.2992	1.1066	1.1081

*Note: See Table 3.1. Entries are statistics that measure *in-sample* and *out-of-sample* volatility forecasting performance of the HAR model given in equation (3.21) of Section 3.3.4, for the target variable given in the title of the table (i.e., the SPY ETF). All models other than the benchmark (HAR) model, denoted as “Benchmark”, include latent volatility factors. EN1 and EN2 denote models for which elastic net shrinkage is used in initial variable selection, with $\alpha = 0.2$ and 0.6 , respectively. Lasso denotes use of the least absolute shrinkage operator in initial variable selection. After initial variable selection, either PCA or sparse PCA (i.e., SPCA) are utilized to obtain the latent volatility factor used in all models denoted as such. *In-sample* R^2 , *Out-of-sample* R^2 and HARMSE entries in this table consist of 2 columns each, the first of which corresponds to predictions made using 267 stocks in factor construction, and the second of which utilizes 274 stocks in the step of our analysis (see Section 3.4.1 for further details). All other tables report results based only on the analysis of 267 stocks. Complete details are given in Sections 3.3 and 3.4.

Table 3.3: Materials Sector ETF (XLB)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5598	0.3204	0.8258
	EN1-PCA	0.5598	0.3204	0.8258
	EN2-PCA	0.5598	0.3204	0.8258
	Lasso-PCA	0.5598	0.3204	0.8258
	EN1-SPCA	0.5678	0.3616	0.6902
	EN2-SPCA	0.5673	0.3647	0.6904
	Lasso-SPCA	0.5673	0.3686	0.6910
5-minute	Benchmark	0.6234	0.2853	1.0050
	EN1-PCA	0.6274	0.3107	0.9057
	EN2-PCA	0.6271	0.3047	0.9316
	Lasso-PCA	0.6269	0.3053	0.9303
	EN1-SPCA	0.6341	0.3322	0.7841
	EN2-SPCA	0.6345	0.3351	0.7887
	Lasso-SPCA	0.6348	0.3445	0.7970
10-minute	Benchmark	0.5497	0.1131	1.2993
	EN1-PCA	0.5712	0.1699	1.0226
	EN2-PCA	0.5717	0.1684	1.0258
	Lasso-PCA	0.5709	0.1682	1.0329
	EN1-SPCA	0.5702	0.1833	1.0078
	EN2-SPCA	0.5694	0.1815	1.0187
	Lasso-SPCA	0.5690	0.1735	1.0100

* Notes: See notes to Table 3.2.

Table 3.4: Energy Sector ETF (XLE)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5221	0.1932	1.1601
	EN1-PCA	0.5239	0.2910	0.9712
	EN2-PCA	0.5236	0.2925	0.9773
	Lasso-PCA	0.5236	0.2910	0.9764
	EN1-SPCA	0.5445	0.3592	0.6126
	EN2-SPCA	0.5451	0.3664	0.6065
	Lasso-SPCA	0.5462	0.3637	0.6018
5-minute	Benchmark	0.6203	0.3153	1.1597
	EN1-PCA	0.6240	0.3750	1.0010
	EN2-PCA	0.6242	0.3577	0.9668
	Lasso-PCA	0.6231	0.3608	0.9830
	EN1-SPCA	0.6286	0.4192	0.7402
	EN2-SPCA	0.6308	0.4277	0.7335
	Lasso-SPCA	0.6298	0.3993	0.7473
10-minute	Benchmark	0.5374	0.1878	1.4904
	EN1-PCA	0.5667	0.3442	0.9174
	EN2-PCA	0.5656	0.3575	0.8816
	Lasso-PCA	0.5652	0.3547	0.9139
	EN1-SPCA	0.5601	0.3315	0.8931
	EN2-SPCA	0.5595	0.3793	0.8741
	Lasso-SPCA	0.5591	0.3820	0.8863

* Notes: See notes to Table 3.2.

Table 3.5: Financial Sector ETF (XLF)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5423	0.3026	0.7466
	EN1-PCA	0.5441	0.2695	0.7847
	EN2-PCA	0.5445	0.2433	0.7978
	Lasso-PCA	0.5450	0.2106	0.8075
	EN1-SPCA	0.6258	0.3585	0.6668
	EN2-SPCA	0.6299	0.3085	0.7173
	Lasso-SPCA	0.6335	0.1985	0.7334
5-minute	Benchmark	0.5823	0.2853	1.3230
	EN1-PCA	0.6085	0.2565	1.2094
	EN2-PCA	0.6028	0.2896	1.2651
	Lasso-PCA	0.5972	0.2508	1.3114
	EN1-SPCA	0.6145	0.2612	1.2482
	EN2-SPCA	0.6150	0.2648	1.3372
	Lasso-SPCA	0.6149	0.2652	1.3766
10-minute	Benchmark	0.4950	0.1276	1.7122
	EN1-PCA	0.5386	0.0960	1.8255
	EN2-PCA	0.5393	0.1032	1.8221
	Lasso-PCA	0.5391	0.1053	1.8167
	EN1-SPCA	0.5427	0.0852	1.8423
	EN2-SPCA	0.5400	0.0881	1.8621
	Lasso-SPCA	0.5402	0.0984	1.8578

*Notes: See notes to Table 3.2.

Table 3.6: Industrial Sector ETF (XLI)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5573	0.3389	0.8208
	EN1-PCA	0.5668	0.3589	0.6438
	EN2-PCA	0.5659	0.3607	0.6585
	Lasso-PCA	0.5658	0.3491	0.6513
	EN1-SPCA	0.5848	0.3724	0.5040
	EN2-SPCA	0.5887	0.3681	0.4973
	Lasso-SPCA	0.5890	0.3771	0.4896
5-minute	Benchmark	0.6219	0.3217	1.6667
	EN1-PCA	0.6398	0.3211	1.0840
	EN2-PCA	0.6389	0.3155	1.3193
	Lasso-PCA	0.6380	0.3094	1.2992
	EN1-SPCA	0.6547	0.3363	0.9299
	EN2-SPCA	0.6534	0.3324	1.0008
	Lasso-SPCA	0.6528	0.3364	0.9307
10-minute	Benchmark	0.5309	0.0715	1.5196
	EN1-PCA	0.5566	0.0982	1.0240
	EN2-PCA	0.5571	0.1125	1.0148
	Lasso-PCA	0.5575	0.1172	1.0048
	EN1-SPCA	0.5529	0.1136	1.0563
	EN2-SPCA	0.5540	0.1211	1.0552
	Lasso-SPCA	0.5538	0.1244	1.0401

*Notes: See notes to Table 3.2.

Table 3.7: Technology Sector ETF (XLK)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5311	0.2340	0.6839
	EN1-PCA	0.5370	0.2848	0.5778
	EN2-PCA	0.5372	0.2800	0.5765
	Lasso-PCA	0.5372	0.2801	0.5757
	EN1-SPCA	0.5458	0.3103	0.4815
	EN2-SPCA	0.5458	0.3046	0.4806
	Lasso-SPCA	0.5456	0.2981	0.4914
5-minute	Benchmark	0.6171	0.2849	0.9884
	EN1-PCA	0.6207	0.3009	0.9021
	EN2-PCA	0.6192	0.2892	0.9350
	Lasso-PCA	0.6198	0.3031	0.9043
	EN1-SPCA	0.6302	0.3183	0.7123
	EN2-SPCA	0.6309	0.3077	0.7020
	Lasso-SPCA	0.6311	0.3093	0.7011
10-minute	Benchmark	0.5118	0.0451	1.3730
	EN1-PCA	0.5363	0.0929	0.9936
	EN2-PCA	0.5362	0.1002	0.9901
	Lasso-PCA	0.5364	0.0937	0.9833
	EN1-SPCA	0.5342	0.1050	0.9818
	EN2-SPCA	0.5341	0.1051	0.9791
	Lasso-SPCA	0.5344	0.1028	0.9476

*Notes: See notes to Table 3.2.

Table 3.8: Consumer Staples Sector ETF (XLP)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3840	0.1047	0.7164
	EN1-PCA	0.4066	0.1905	0.4557
	EN2-PCA	0.4066	0.1931	0.4510
	Lasso-PCA	0.4067	0.1988	0.4554
	EN1-SPCA	0.4405	0.1830	0.3920
	EN2-SPCA	0.4353	0.1194	0.4026
	Lasso-SPCA	0.4342	0.1703	0.3983
5-minute	Benchmark	0.4578	0.2790	0.9885
	EN1-PCA	0.4929	0.4753	0.5346
	EN2-PCA	0.4908	0.4162	0.5487
	Lasso-PCA	0.4910	0.4236	0.5675
	EN1-SPCA	0.5165	0.4089	0.5666
	EN2-SPCA	0.5188	0.3479	0.5794
	Lasso-SPCA	0.5180	0.4234	0.5744
10-minute	Benchmark	0.4001	0.1796	1.3737
	EN1-PCA	0.4870	0.2990	1.0413
	EN2-PCA	0.4887	0.2953	1.0044
	Lasso-PCA	0.4893	0.2789	1.0221
	EN1-SPCA	0.4840	0.2621	1.0628
	EN2-SPCA	0.4930	0.2432	1.0490
	Lasso-SPCA	0.4942	0.2278	1.0707

*Notes: See notes to Table 3.2.

Table 3.9: Utilities Sector ETF (XLU)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5242	0.1309	0.8652
	EN1-PCA	0.5331	0.1515	0.5590
	EN2-PCA	0.5333	0.1359	0.5620
	Lasso-PCA	0.5332	0.1549	0.5587
	EN1-SPCA	0.5380	0.1869	0.5176
	EN2-SPCA	0.5374	0.1747	0.5161
	Lasso-SPCA	0.5376	0.1806	0.5265
5-minute	Benchmark	0.5683	0.1887	1.1073
	EN1-PCA	0.5906	0.2594	0.6719
	EN2-PCA	0.5870	0.2559	0.6629
	Lasso-PCA	0.5845	0.2746	0.6580
	EN1-SPCA	0.6160	0.2751	0.8010
	EN2-SPCA	0.6118	0.2618	0.7718
	Lasso-SPCA	0.6108	0.2655	0.7661
10-minute	Benchmark	0.4960	0.1882	1.3382
	EN1-PCA	0.5320	0.3671	0.8231
	EN2-PCA	0.5307	0.3879	0.8198
	Lasso-PCA	0.5304	0.3563	0.8261
	EN1-SPCA	0.5307	0.3662	0.8257
	EN2-SPCA	0.5292	0.3896	0.8234
	Lasso-SPCA	0.5292	0.3577	0.8382

*Notes: See notes to Table 3.2.

Table 3.10: Health Care Sector ETF (XLV)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5053	0.2481	0.6576
	EN1-PCA	0.5185	0.2678	0.4433
	EN2-PCA	0.5186	0.2706	0.4399
	Lasso-PCA	0.5186	0.2610	0.4389
	EN1-SPCA	0.5257	0.2629	0.4309
	EN2-SPCA	0.5269	0.2796	0.4163
	Lasso-SPCA	0.5276	0.2395	0.4230
5-minute	Benchmark	0.4735	0.2067	1.0695
	EN1-PCA	0.5027	0.3332	0.6355
	EN2-PCA	0.5019	0.3218	0.6613
	Lasso-PCA	0.5014	0.3263	0.6656
	EN1-SPCA	0.5400	0.3070	0.6025
	EN2-SPCA	0.5404	0.3218	0.6022
	Lasso-SPCA	0.5423	0.2934	0.6048
10-minute	Benchmark	0.4566	0.2016	1.2785
	EN1-PCA	0.5087	0.3486	0.7555
	EN2-PCA	0.5102	0.3498	0.7428
	Lasso-PCA	0.5098	0.3648	0.7550
	EN1-SPCA	0.5056	0.3654	0.7431
	EN2-SPCA	0.5077	0.3533	0.7678
	Lasso-SPCA	0.5071	0.3733	0.7317

*Notes: See notes to Table 3.2.

Table 3.11: Consumer Discretionary Sector ETF (XLY)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5255	0.3513	0.8867
	EN1-PCA	0.5370	0.4082	0.7845
	EN2-PCA	0.5367	0.3972	0.7855
	Lasso-PCA	0.5366	0.4126	0.7861
	EN1-SPCA	0.5557	0.4236	0.5854
	EN2-SPCA	0.5550	0.4018	0.5686
	Lasso-SPCA	0.5544	0.4171	0.5993
5-minute	Benchmark	0.5724	0.3408	1.3435
	EN1-PCA	0.5832	0.3581	1.2122
	EN2-PCA	0.5841	0.3691	1.2269
	Lasso-PCA	0.5849	0.3724	1.2084
	EN1-SPCA	0.6120	0.4058	0.9128
	EN2-SPCA	0.6119	0.4056	0.9064
	Lasso-SPCA	0.6117	0.4103	0.9152
10-minute	Benchmark	0.4784	0.1197	1.5265
	EN1-PCA	0.5177	0.1966	1.0291
	EN2-PCA	0.5195	0.1848	1.0086
	Lasso-PCA	0.5196	0.1853	1.0053
	EN1-SPCA	0.5133	0.1880	1.0273
	EN2-SPCA	0.5157	0.1854	0.9921
	Lasso-SPCA	0.5161	0.1904	0.9789

*Notes: See notes to Table 3.2.

Table 3.12: General Electric Company (GE)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5211	0.2898	0.8151
	EN1-PCA	0.5243	0.3123	0.7342
	EN2-PCA	0.5240	0.3132	0.7350
	Lasso-PCA	0.5240	0.3130	0.7367
	EN1-SPCA	0.5792	0.3823	0.5655
	EN2-SPCA	0.5816	0.4005	0.5602
	Lasso-SPCA	0.5800	0.3954	0.5665
5-minute	Benchmark	0.5189	0.1576	1.2367
	EN1-PCA	0.5554	0.1710	0.9424
	EN2-PCA	0.5435	0.2130	1.0303
	Lasso-PCA	0.5522	0.1848	0.9533
	EN1-SPCA	0.5821	0.2586	0.8025
	EN2-SPCA	0.5823	0.2576	0.8256
	Lasso-SPCA	0.5830	0.2665	0.8301
10-minute	Benchmark	0.4825	0.0555	1.5839
	EN1-PCA	0.4893	0.0943	1.3869
	EN2-PCA	0.4900	0.1012	1.3659
	Lasso-PCA	0.4908	0.1041	1.3368
	EN1-SPCA	0.4901	0.0997	1.3638
	EN2-SPCA	0.4905	0.0980	1.3590
	Lasso-SPCA	0.4912	0.1020	1.3359

*Notes: See notes to Table 3.2.

Table 3.13: The Goldman Sachs Group, Inc. (GS)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3939	-0.2130	1.6083
	EN1-PCA	0.3920	-0.2120	1.6196
	EN2-PCA	0.3920	-0.2120	1.6196
	Lasso-PCA	0.3920	-0.2120	1.6196
	EN1-SPCA	0.4078	0.0106	0.9982
	EN2-SPCA	0.4180	-0.0706	1.0676
	Lasso-SPCA	0.4317	-0.0856	1.0974
5-minute	Benchmark	0.4206	-0.2341	2.0352
	EN1-PCA	0.4187	-0.2100	2.0106
	EN2-PCA	0.4187	-0.2228	2.0288
	Lasso-PCA	0.4187	-0.2213	2.0235
	EN1-SPCA	0.4258	-0.1169	1.7640
	EN2-SPCA	0.4226	-0.1392	1.7918
	Lasso-SPCA	0.4402	-0.0366	1.3739
10-minute	Benchmark	0.3758	-0.2761	2.5707
	EN1-PCA	0.3957	-0.0406	1.5839
	EN2-PCA	0.4006	-0.0435	1.6176
	Lasso-PCA	0.3990	-0.0497	1.6556
	EN1-SPCA	0.3971	-0.0670	1.7213
	EN2-SPCA	0.4016	-0.0807	1.7189
	Lasso-SPCA	0.4007	-0.0789	1.7528

*Notes: See notes to Table 3.2.

Table 3.14: International Business Machines Corporation (IBM)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5380	0.1149	1.1117
	EN1-PCA	0.5423	0.1855	0.9083
	EN2-PCA	0.5424	0.1707	0.9089
	Lasso-PCA	0.5424	0.1693	0.9146
	EN1-SPCA	0.5623	0.1478	0.6466
	EN2-SPCA	0.5621	0.2774	0.6214
	Lasso-SPCA	0.5632	0.2785	0.6124
5-minute	Benchmark	0.6140	0.2374	1.0384
	EN1-PCA	0.6194	0.3004	0.9106
	EN2-PCA	0.6281	0.3167	0.8908
	Lasso-PCA	0.6319	0.3215	0.8284
	EN1-SPCA	0.6463	0.3826	0.7098
	EN2-SPCA	0.6518	0.3709	0.7436
	Lasso-SPCA	0.6538	0.3578	0.7521
10-minute	Benchmark	0.5936	0.1696	1.1609
	EN1-PCA	0.5993	0.2111	1.0156
	EN2-PCA	0.5993	0.2138	1.0000
	Lasso-PCA	0.5995	0.2177	0.9866
	EN1-SPCA	0.5989	0.2306	0.9837
	EN2-SPCA	0.5987	0.2295	0.9792
	Lasso-SPCA	0.5986	0.2283	0.9772

*Notes: See notes to Table 3.2.

Table 3.15: Johnson & Johnson (JNJ)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3611	0.1158	0.8426
	EN1-PCA	0.4040	0.2748	0.4300
	EN2-PCA	0.4039	0.2580	0.4331
	Lasso-PCA	0.4041	0.2588	0.4411
	EN1-SPCA	0.4415	0.2210	0.4721
	EN2-SPCA	0.4407	0.2319	0.4500
	Lasso-SPCA	0.4380	0.2300	0.4568
5-minute	Benchmark	0.3882	0.1398	1.0297
	EN1-PCA	0.4356	0.3251	0.5415
	EN2-PCA	0.4316	0.3355	0.5533
	Lasso-PCA	0.4310	0.3738	0.5387
	EN1-SPCA	0.4814	0.2968	0.5557
	EN2-SPCA	0.4816	0.3496	0.5746
	Lasso-SPCA	0.4815	0.3118	0.5709
10-minute	Benchmark	0.3740	0.1091	1.3244
	EN1-PCA	0.4395	0.2914	0.8430
	EN2-PCA	0.4432	0.3202	0.8348
	Lasso-PCA	0.4391	0.3136	0.8480
	EN1-SPCA	0.4450	0.3015	0.8406
	EN2-SPCA	0.4489	0.2860	0.8483
	Lasso-SPCA	0.4444	0.3513	0.8371

*Notes: See notes to Table 3.2.

Table 3.16: JPMorgan Chase & Co. (JPM)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5122	-0.0657	1.1058
	EN1-PCA	0.5122	-0.0659	1.1060
	EN2-PCA	0.5122	-0.0657	1.1058
	Lasso-PCA	0.5122	-0.0657	1.1058
	EN1-SPCA	0.5730	-0.0419	0.9726
	EN2-SPCA	0.5742	-0.0790	1.0505
	Lasso-SPCA	0.5711	-0.0569	1.0658
5-minute	Benchmark	0.5543	0.0369	1.3160
	EN1-PCA	0.5640	0.0699	1.3026
	EN2-PCA	0.5592	0.0438	1.2999
	Lasso-PCA	0.5574	0.0552	1.3058
	EN1-SPCA	0.5745	0.0542	1.2892
	EN2-SPCA	0.5703	0.0713	1.2437
	Lasso-SPCA	0.5709	0.0877	1.2393
10-minute	Benchmark	0.4516	-0.2183	1.8278
	EN1-PCA	0.4682	-0.2899	1.8026
	EN2-PCA	0.4763	-0.2392	1.7838
	Lasso-PCA	0.4787	-0.2438	1.7681
	EN1-SPCA	0.4810	-0.2596	1.7216
	EN2-SPCA	0.4899	-0.2128	1.7214
	Lasso-SPCA	0.4911	-0.2275	1.7314

*Notes: See notes to Table 3.2.

Table 3.17: The Coca-Cola Company (KO)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.4082	0.1379	0.9190
	EN1-PCA	0.4384	0.2487	0.5931
	EN2-PCA	0.4386	0.2568	0.5920
	Lasso-PCA	0.4385	0.2417	0.5880
	EN1-SPCA	0.4504	0.2621	0.4465
	EN2-SPCA	0.4500	0.2681	0.4571
	Lasso-SPCA	0.4501	0.2478	0.4592
5-minute	Benchmark	0.5598	0.2292	1.1106
	EN1-PCA	0.6039	0.3952	0.7784
	EN2-PCA	0.5996	0.3626	0.7949
	Lasso-PCA	0.5998	0.3954	0.7813
	EN1-SPCA	0.6194	0.3500	0.7107
	EN2-SPCA	0.6166	0.4174	0.6635
	Lasso-SPCA	0.6170	0.3807	0.6651
10-minute	Benchmark	0.5006	0.1572	1.4081
	EN1-PCA	0.5687	0.2966	1.0139
	EN2-PCA	0.5702	0.2943	0.9600
	Lasso-PCA	0.5679	0.2459	1.0471
	EN1-SPCA	0.5707	0.2637	0.9683
	EN2-SPCA	0.5699	0.2831	0.9421
	Lasso-SPCA	0.5671	0.2428	0.9493

*Notes: See notes to Table 3.2.

Table 3.18: McDonald's Corporation (MCD)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3738	-0.1118	0.9516
	EN1-PCA	0.3921	0.1359	0.6769
	EN2-PCA	0.3920	0.1485	0.6735
	Lasso-PCA	0.3920	0.1700	0.6680
	EN1-SPCA	0.4606	0.2318	0.5474
	EN2-SPCA	0.4576	0.1939	0.5311
	Lasso-SPCA	0.4604	0.0285	0.5411
5-minute	Benchmark	0.3785	-0.1553	1.3427
	EN1-PCA	0.4219	0.2491	0.8218
	EN2-PCA	0.4129	0.2093	0.8621
	Lasso-PCA	0.4152	0.2078	0.8535
	EN1-SPCA	0.4709	0.2464	0.7252
	EN2-SPCA	0.4711	0.2126	0.7607
	Lasso-SPCA	0.4721	0.1929	0.7566
10-minute	Benchmark	0.3299	-0.1506	1.9629
	EN1-PCA	0.4212	-0.0134	1.5636
	EN2-PCA	0.4174	0.0291	1.5366
	Lasso-PCA	0.4192	0.0759	1.5963
	EN1-SPCA	0.4241	0.0333	1.2639
	EN2-SPCA	0.4226	0.1089	1.3563
	Lasso-SPCA	0.4240	0.1185	1.3150

*Notes: See notes to Table 3.2.

Table 3.19: Merck & Co., Inc. (MRK)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3563	0.1561	0.7700
	EN1-PCA	0.3910	0.2377	0.5873
	EN2-PCA	0.3916	0.2401	0.5876
	Lasso-PCA	0.3917	0.2400	0.5863
	EN1-SPCA	0.4508	0.2370	0.3857
	EN2-SPCA	0.4525	0.2424	0.3835
	Lasso-SPCA	0.4523	0.2381	0.3654
5-minute	Benchmark	0.4129	0.1914	0.9668
	EN1-PCA	0.4721	0.2825	0.6938
	EN2-PCA	0.4686	0.2804	0.7012
	Lasso-PCA	0.4687	0.2847	0.7420
	EN1-SPCA	0.5452	0.2966	0.4908
	EN2-SPCA	0.5444	0.2929	0.4956
	Lasso-SPCA	0.5454	0.2965	0.4892
10-minute	Benchmark	0.4723	0.2843	1.1830
	EN1-PCA	0.5028	0.3981	1.0175
	EN2-PCA	0.5020	0.3873	1.0178
	Lasso-PCA	0.5020	0.4072	1.0136
	EN1-SPCA	0.5010	0.4063	0.9486
	EN2-SPCA	0.5005	0.4143	0.9526
	Lasso-SPCA	0.5006	0.4117	0.9304

*Notes: See notes to Table 3.2.

Table 3.20: Microsoft Corporation (MSFT)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.5622	0.2316	0.7706
	EN1-PCA	0.5699	0.2751	0.7567
	EN2-PCA	0.5701	0.2639	0.7563
	Lasso-PCA	0.5702	0.2681	0.7555
	EN1-SPCA	0.5944	0.3269	0.5756
	EN2-SPCA	0.5955	0.3086	0.5804
	Lasso-SPCA	0.5952	0.3499	0.5781
5-minute	Benchmark	0.6116	0.2394	1.0603
	EN1-PCA	0.6173	0.2816	1.0187
	EN2-PCA	0.6172	0.2784	1.0136
	Lasso-PCA	0.6171	0.2777	1.0110
	EN1-SPCA	0.6329	0.3305	0.8306
	EN2-SPCA	0.6324	0.3169	0.7993
	Lasso-SPCA	0.6327	0.3141	0.8153
10-minute	Benchmark	0.4991	0.1190	2.2867
	EN1-PCA	0.5126	0.1930	2.0942
	EN2-PCA	0.5120	0.1798	2.1108
	Lasso-PCA	0.5114	0.1817	2.1030
	EN1-SPCA	0.5158	0.2166	2.0068
	EN2-SPCA	0.5146	0.1975	1.8044
	Lasso-SPCA	0.5136	0.2134	1.9075

*Notes: See notes to Table 3.2.

Table 3.21: Wal-Mart Stores, Inc. (WMT)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3552	-0.3547	1.0014
	EN1-PCA	0.3693	-0.2423	0.8893
	EN2-PCA	0.3694	-0.2523	0.8926
	Lasso-PCA	0.3695	-0.2480	0.8926
	EN1-SPCA	0.3991	0.0710	0.6201
	EN2-SPCA	0.3977	0.1636	0.6056
	Lasso-SPCA	0.3976	0.1615	0.5976
5-minute	Benchmark	0.4571	-0.2414	1.1750
	EN1-PCA	0.4822	0.0186	0.8979
	EN2-PCA	0.4822	0.0918	0.9150
	Lasso-PCA	0.4849	0.1143	0.8939
	EN1-SPCA	0.5376	0.1146	0.8349
	EN2-SPCA	0.5274	0.0704	0.8387
	Lasso-SPCA	0.5263	0.0663	0.8486
10-minute	Benchmark	0.4012	-0.2658	1.5235
	EN1-PCA	0.4723	0.0772	1.0864
	EN2-PCA	0.4758	0.0955	1.0586
	Lasso-PCA	0.4759	0.0843	1.0716
	EN1-SPCA	0.4659	0.0665	1.1765
	EN2-SPCA	0.4684	0.0546	1.1446
	Lasso-SPCA	0.4690	0.0501	1.1425

*Notes: See notes to Table 3.2.

Table 3.22: Exxon Mobil Corporation (XOM)

Frequency	Model	<i>In-Sample</i> R^2	<i>Out-of-Sample</i> R^2	HARMSE
1-minute	Benchmark	0.3620	-0.0409	1.2201
	EN1-PCA	0.3676	0.2644	0.7383
	EN2-PCA	0.3668	0.2590	0.7495
	Lasso-PCA	0.3667	0.2555	0.7482
	EN1-SPCA	0.3998	0.3272	0.4526
	EN2-SPCA	0.3995	0.3065	0.4518
	Lasso-SPCA	0.4001	0.3000	0.4684
5-minute	Benchmark	0.4823	0.0819	1.2181
	EN1-PCA	0.5032	0.3501	0.7590
	EN2-PCA	0.5011	0.3082	0.7848
	Lasso-PCA	0.4989	0.3053	0.7824
	EN1-SPCA	0.5444	0.3630	0.7415
	EN2-SPCA	0.5408	0.3450	0.7232
	Lasso-SPCA	0.5387	0.3744	0.7138
10-minute	Benchmark	0.4102	0.0355	1.5327
	EN1-PCA	0.4801	0.2805	1.0449
	EN2-PCA	0.4791	0.2996	1.0267
	Lasso-PCA	0.4812	0.2546	1.0546
	EN1-SPCA	0.4786	0.2386	1.0775
	EN2-SPCA	0.4781	0.2673	1.0312
	Lasso-SPCA	0.4809	0.1981	1.0852

*Notes: See notes to Table 3.2.

Table 3.23: Factor Structure (SPY)

Sampling Frequency: 1 Minute																	
Stock		EN-1				EN-2				Stock		Lasso					
Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker		
AFL	F	1.000	0.129	0.105	0.955	AFL	F	1.000	0.131	0.106	0.961	AFL	F	1.000	0.131	0.105	0.958
AMO	CS	1.000	0.083	0.024	0.397	AMO	CS	1.000	0.085	0.025	0.392	AMO	CS	1.000	0.086	0.025	0.389
AMT	T	1.000	0.107	0.057	0.808	AMT	T	1.000	0.109	0.057	0.808	AMT	T	1.000	0.109	0.057	0.808
ABC	H	1.000	0.126	0.112	0.787	ABC	H	1.000	0.128	0.114	0.779	ABC	H	1.000	0.129	0.112	0.789
AMGN	F	1.000	0.126	0.101	0.832	AMGN	F	1.000	0.129	0.106	0.829	AMGN	F	1.000	0.129	0.105	0.832
BK	F	1.000	0.087	0.019	0.463	BK	F	1.000	0.088	0.019	0.463	BK	F	1.000	0.088	0.019	0.468
BBY	CD	1.000	0.137	0.126	0.963	BBY	CD	1.000	0.139	0.128	0.963	BBY	CD	1.000	0.140	0.128	0.963
CPB	CS	1.000	0.103	0.055	0.668	CPB	CS	1.000	0.105	0.055	0.661	CPB	CS	1.000	0.105	0.056	0.639
CAH	H	1.000	0.135	0.126	0.892	CAH	H	1.000	0.138	0.130	0.882	CAH	H	1.000	0.138	0.128	0.884
CI	H	1.000	0.109	0.059	0.805	CI	H	1.000	0.111	0.060	0.813	CI	H	1.000	0.112	0.060	0.797
CAG	CS	1.000	0.110	0.071	0.629	CAG	CS	1.000	0.111	0.073	0.645	CAG	CS	1.000	0.112	0.073	0.639
HSY	CS	1.000	0.105	0.054	0.632	HSY	CS	1.000	0.107	0.057	0.653	HSY	CS	1.000	0.107	0.057	0.642
HBAN	F	1.000	0.064	0.011	0.168	HBAN	F	1.000	0.066	0.011	0.184	HBAN	F	1.000	0.065	0.011	0.168
ILMN	CS	1.000	0.105	0.055	0.629	ILMN	CS	1.000	0.106	0.055	0.642	ILMN	CS	1.000	0.108	0.056	0.634
KR	CS	1.000	0.090	0.034	0.487	KR	CS	1.000	0.092	0.033	0.487	KR	CS	1.000	0.092	0.033	0.484
LM	F	1.000	0.095	0.034	0.624	LM	F	1.000	0.096	0.034	0.618	LM	F	1.000	0.096	0.033	0.618
LNC	F	1.000	0.117	0.074	0.861	LNC	F	1.000	0.118	0.074	0.861	LNC	F	1.000	0.119	0.074	0.863
MMC	F	1.000	0.092	0.033	0.600	MMC	F	1.000	0.094	0.034	0.595	MMC	F	1.000	0.094	0.035	0.600
PEP	CS	1.000	0.114	0.072	0.771	PEP	CS	1.000	0.116	0.076	0.768	PEP	CS	1.000	0.116	0.074	0.766
PRU	F	1.000	0.104	0.047	0.718	PRU	F	1.000	0.106	0.048	0.721	PRU	F	1.000	0.106	0.047	0.734
SWN	E	1.000	0.128	0.082	0.900	SWN	E	1.000	0.129	0.083	0.907	SWN	E	1.000	0.119	0.081	0.834
SBUX	CD	1.000	0.151	0.158	0.987	SBUX	CD	1.000	0.154	0.160	0.987	SBUX	CD	1.000	0.154	0.158	0.987
SYF	CS	1.000	0.108	0.063	0.729	TROW	F	1.000	0.115	0.068	0.889	TROW	F	1.000	0.116	0.068	0.895
TROW	F	1.000	0.114	0.067	0.887	TIF	CD	1.000	0.133	0.112	0.968	TIF	CD	1.000	0.133	0.110	0.968
UNP	CD	1.000	0.130	0.110	0.971	UNP	CD	1.000	0.142	0.132	0.966	UNP	CD	1.000	0.142	0.132	0.963
UNM	I	1.000	0.139	0.129	0.966	UNAM	F	1.000	0.125	0.092	0.887	UNM	F	1.000	0.125	0.092	0.884
UNX	F	1.000	0.123	0.092	0.882	WYNN	CD	1.000	0.143	0.133	0.987	WYNN	CD	1.000	0.143	0.132	0.982
CNX	U	0.997	0.133	0.114	0.955	MET	F	1.000	0.114	0.068	0.853	MET	F	1.000	0.114	0.068	0.853
WYNN	CD	0.997	0.141	0.131	0.984	CNX	U	0.992	0.135	0.115	0.960	SYK	H	0.992	0.144	0.140	0.947
CERN	T	0.992	0.147	0.150	0.947	CERN	T	0.992	0.150	0.142	0.960	CYH	F	0.987	0.136	0.114	0.960
SYK	H	0.989	0.142	0.138	0.939	SYK	H	0.982	0.144	0.142	0.944	WYNN	CD	0.987	0.146	0.136	0.960
MRK	I	0.971	0.123	0.094	0.894	SYF	CS	0.976	0.109	0.062	0.733	CERN	T	0.976	0.151	0.154	0.962
GE	F	0.968	0.134	0.120	0.927	AXP	F	0.961	0.124	0.090	0.945	SYF	CS	0.974	0.110	0.064	0.743

Sampling Frequency: 5 Minute																	
Stock		EN-1				EN-2				Stock		Lasso					
Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker		
BK	F	1.000	0.077	0.020	0.418	BK	F	1.000	0.075	0.018	0.439	BK	F	1.000	0.075	0.018	0.434
BBY	CD	1.000	0.125	0.108	0.924	BBY	CD	1.000	0.123	0.105	0.913	BBY	CD	1.000	0.121	0.103	0.929
BSX	T	1.000	0.063	0.003	0.542	BSX	T	1.000	0.063	0.003	0.539	BSX	T	1.000	0.062	0.006	0.532
CNX	I	1.000	0.125	0.103	0.924	CNX	I	1.000	0.122	0.103	0.929	CNX	I	1.000	0.121	0.098	0.937
CSX	U	1.000	0.130	0.115	0.932	CSX	U	1.000	0.127	0.111	0.934	CSX	U	1.000	0.127	0.112	0.926
CAH	H	1.000	0.128	0.117	0.861	CAH	H	1.000	0.124	0.112	0.845	CAH	H	1.000	0.124	0.114	0.826
CERN	T	1.000	0.146	0.149	0.916	CERN	T	1.000	0.142	0.146	0.924	CERN	T	1.000	0.141	0.145	0.916
CI	H	1.000	0.098	0.053	0.689	CI	H	1.000	0.095	0.050	0.676	CI	H	1.000	0.094	0.050	0.668
CCI	CD	1.000	0.085	0.040	0.584	CCI	CD	1.000	0.084	0.039	0.584	CCI	CD	1.000	0.083	0.039	0.587
DHI	CD	1.000	0.099	0.056	0.763	DHI	CD	1.000	0.098	0.056	0.750	DHI	CD	1.000	0.098	0.057	0.758
FITB	F	1.000	0.080	0.028	0.497	FITB	F	1.000	0.079	0.026	0.492	FITB	F	1.000	0.079	0.026	0.492
HSY	CS	1.000	0.094	0.051	0.674	HSY	CS	1.000	0.091	0.048	0.658	HSY	CS	1.000	0.090	0.048	0.649
ILMN	H	1.000	0.096	0.059	0.632	ILMN	H	1.000	0.097	0.061	0.629	ILMN	H	1.000	0.092	0.056	0.629
LNC	F	1.000	0.103	0.063	0.805	LNC	F	1.000	0.100	0.059	0.784	LNC	F	1.000	0.100	0.060	0.784
MUR	E	1.000	0.130	0.114	0.934	MUR	E	1.000	0.126	0.109	0.934	MUR	E	1.000	0.125	0.108	0.947
NWSA	CD	1.000	0.143	0.148	0.947	NWSA	CD	1.000	0.141	0.144	0.950	NWSA	CD	1.000	0.141	0.147	0.942
PRU	F	1.000	0.092	0.041	0.716	PRU	F	1.000	0.090	0.040	0.713	PRU	F	1.000	0.090	0.039	0.711
DCX	H	1.000	0.114	0.089	0.771	DCX	H	1.000	0.110	0.087	0.742	DCX	H	1.000	0.107	0.076	0.758
HOT	CD	1.000	0.133	0.125	0.911	HOT	CD	1.000	0.130	0.122	0.913	HOT	CD	1.000	0.131	0.125	0.918
THC	F	1.000	0.104	0.061	0.826	TROW	F	1.000	0.101	0.057	0.818	TROW	F	1.000	0.102	0.059	0.832
TIF	CD	1.000	0.076	0.034	0.487	THC	H	1.000	0.076	0.034	0.497	THC	H	1.000	0.073	0.030	0.476
UTX	I	1.000	0.143	0.148	0.966	UTX	I	1.000	0.138	0.141	0.942	UTX	I	1.000	0.140	0.145	0.963
WYNN	CD	1.000	0.136	0.132	0.966	UTX	I	1.000	0.133	0.128	0.968	UTX	I	1.000	0.133	0.130	0.955
PEP	CS	0.992	0.100	0.062	0.687	WYNN	CD	1.000	0.127	0.111	0.926	WYNN	CD	1.000	0.127	0.113	0.929
WMB	E	0.989	0.126	0.106	0.904	TYG	CD	1.000	0.103	0.069	0.784	PEP	CS	1.000	0.096	0.058	0.689
MCD	CD	0.987	0.093	0.047	0.699	PEP	CS	0.997	0.098	0.061	0.691	MCD	CD	1.000	0.096	0.058	0.689
TYC	I	0.984	0.106	0.072	0.794	MCD	CD	0.992	0.091	0.046	0.706	TYC	I	0.997	0.103	0.070	0.773
SYK	H	0.979	0.117	0.110	0.871	SMB	CD	0.997	0.091	0.046	0.706	TGT	CD	0.984	0.101	0.058	0.797
SWN	E	0.971	0.112	0.078	0.818	FDO	CD	0.955	0.073	0.021	0.408	WMB	E	0.958	0.123	0.104	0.915
FDO	CD	0.966	0.075	0.024	0.420	SYK	H	0.939	0.124	0.110	0.866	FDO	CD	0.939	0.073	0.021	0.434
TGT	E	0.963	0.101	0.059	0.803	TGT	E	0.937	0.098	0.055	0.795	MRK	H	0.924	0.101	0.070	0.758
						SWN	E	0.926	0.110	0.074	0.807	LLTC	T	0.921	0.161	0.193	0.989

Sampling Frequency: 10 Minute																	
Stock		EN-1				EN-2				Stock		Lasso					
Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker		
AFL	F	1.000	0.115	0.083	0.837	AFL	F	1.000	0.115	0.083	0.834	AFL	F	1.000	0.117	0.084	0.834
BBBY	CD	1.000	0.153	0.145	0.929	BBBY	CD	1.000	0.151	0.142	0.924	BBBY	CD	1.000	0.153	0.143	0.924
BBY	CD	1.000	0.138	0.120	0.876	BBY	CD	1.000	0.136	0.119	0.900	BBY	CD	1.000	0.139	0.121	0.892
CERN	T	1.000	0.157	0.153	0.957	CERN	T	1.000	0.153	0.147	0.958	CERN	T	1.000	0.157	0.158	0.958
CI	H	1.000	0.104	0.065	0.747	CI	H	1.000	0.100	0.060	0.7.						

Table 3.24: Factor Structure (XLE)

Sampling Frequency: 1 Minute																	
EN-1						EN-2						Lasso					
Ticker	Stock	Sector	Freq.	PCA	SPCA	Ticker	Stock	Sector	Freq.	PCA	SPCA	Ticker	Stock	Sector	Freq.	PCA	SPCA
AFL	F	1.000	0.133	0.109	0.961	AFL	F	1.000	0.130	0.105	0.945	AFL	F	1.000	0.130	0.104	0.953
ALL	F	1.000	0.105	0.048	0.682	ALL	F	1.000	0.102	0.045	0.684	ALL	F	1.000	0.102	0.045	0.695
ABC	H	1.000	0.130	0.113	0.771	ABC	H	1.000	0.127	0.109	0.763	ABC	H	1.000	0.129	0.109	0.776
BBY	CD	1.000	0.140	0.128	0.950	BBY	CD	1.000	0.137	0.127	0.961	BBY	CD	1.000	0.139	0.127	0.955
CX	CAH	1.000	0.130	0.120	0.958	CAH	H	1.000	0.126	0.106	0.936	CAH	H	1.000	0.137	0.123	0.945
CAH	H	1.000	0.140	0.130	0.879	CI	H	1.000	0.110	0.057	0.779	CI	H	1.000	0.110	0.057	0.795
CI	H	1.000	0.113	0.061	0.797	CAG	CS	1.000	0.111	0.071	0.618	CAG	CS	1.000	0.108	0.066	0.600
CAG	CS	1.000	0.111	0.069	0.611	FTTB	CS	1.000	0.088	0.023	0.471	FTTB	CS	1.000	0.088	0.024	0.471
FTTB	F	1.000	0.090	0.024	0.474	GS	F	1.000	0.090	0.022	0.479	GS	F	1.000	0.090	0.022	0.476
GS	F	1.000	0.130	0.123	0.958	HSY	CS	1.000	0.104	0.023	0.479	HSY	CS	1.000	0.107	0.023	0.485
HSY	CS	1.000	0.092	0.022	0.487	ILMN	H	1.000	0.104	0.052	0.637	ILMN	H	1.000	0.105	0.056	0.642
HBAN	F	1.000	0.106	0.052	0.629	KR	CS	1.000	0.092	0.034	0.474	KR	CS	1.000	0.091	0.036	0.479
ILMN	CR	1.000	0.067	0.012	0.161	LM	I	1.000	0.095	0.032	0.637	LM	I	1.000	0.095	0.032	0.621
KR	CS	1.000	0.104	0.053	0.611	LNC	F	1.000	0.117	0.073	0.874	LNC	F	1.000	0.117	0.073	0.871
LM	I	1.000	0.092	0.032	0.611	LLTC	F	1.000	0.178	0.073	0.907	LLTC	F	1.000	0.177	0.073	0.905
LNC	F	1.000	0.097	0.033	0.621	MMC	F	1.000	0.092	0.032	0.584	MMC	F	1.000	0.092	0.033	0.559
LLTC	F	1.000	0.121	0.077	0.874	MDT	H	1.000	0.140	0.129	0.918	MDT	H	1.000	0.139	0.127	0.900
LLTC	F	1.000	0.182	0.228	1.000	MET	F	1.000	0.113	0.066	0.853	MET	F	1.000	0.113	0.065	0.858
MON	M	1.000	0.090	0.034	0.918	PEP	CS	1.000	0.105	0.016	0.726	PEP	CS	1.000	0.105	0.017	0.738
MDT	H	1.000	0.143	0.133	0.918	SWN	E	1.000	0.118	0.080	0.834	SWN	E	1.000	0.118	0.079	0.826
MET	F	1.000	0.116	0.070	0.837	SYK	H	1.000	0.142	0.136	0.947	TROW	F	1.000	0.114	0.066	0.897
PEP	CS	1.000	0.116	0.071	0.750	TROW	F	1.000	0.114	0.066	0.892	TIF	CD	1.000	0.130	0.107	0.968
PRU	F	1.000	0.108	0.048	0.711	TIF	CD	1.000	0.130	0.106	0.961	UNM	F	1.000	0.123	0.090	0.889
SWN	E	1.000	0.128	0.108	0.947	AMT	U	1.000	0.123	0.089	0.859	GRMN	CD	1.000	0.114	0.073	0.718
SYK	CS	1.000	0.146	0.142	0.947	AMT	U	1.000	0.108	0.056	0.803	GRMN	CD	1.000	0.114	0.073	0.718
SYK	CS	1.000	0.146	0.142	0.947	CNX	U	0.989	0.135	0.114	0.963	SYK	H	0.995	0.142	0.134	0.939
TROW	F	1.000	0.117	0.069	0.903	PEP	CS	0.987	0.113	0.068	0.755	CA	T	0.995	0.138	0.127	0.944
TIF	CD	1.000	0.133	0.109	0.958	GRMN	CD	0.982	0.115	0.077	0.724	CNX	T	0.992	0.134	0.112	0.958
UNM	F	1.000	0.120	0.084	0.879	HSY	CS	0.975	0.105	0.064	0.755	PRU	F	0.997	0.107	0.082	0.867
VLO	E	1.000	0.120	0.081	0.847	MON	F	0.971	0.099	0.041	0.604	MON	M	0.963	0.099	0.040	0.590
CSX	I	0.997	0.150	0.150	0.976	CA	T	0.966	0.138	0.129	0.948	AXP	F	0.961	0.124	0.091	0.953
MON	M	0.997	0.102	0.042	0.612	AXP	F	0.961	0.124	0.090	0.942	HBAN	F	0.958	0.066	0.011	0.162
CSX	I	0.995	0.110	0.057	0.794	EXP	T	0.939	0.166	0.193	0.997	TXN	T	0.945	0.167	0.194	0.992
SBUX	CD	1.000	0.135	0.129	0.959	UTX	I	0.928	0.139	0.128	0.938	TXN	T	0.945	0.167	0.194	0.992
CERN	T	0.950	0.150	0.152	0.972	UTX	I	0.916	0.142	0.139	0.994	ENDP	H	0.929	0.114	0.082	0.717

Sampling Frequency: 5 Minute																	
EN-1						EN-2						Lasso					
Ticker	Stock	Sector	Freq.	PCA	SPCA	Ticker	Stock	Sector	Freq.	PCA	SPCA	Ticker	Stock	Sector	Freq.	PCA	SPCA
BK	F	1.000	0.069	0.017	0.437	BK	F	1.000	0.065	0.015	0.432	BK	F	1.000	0.066	0.016	0.429
BSX	H	1.000	0.080	0.048	0.511	BSX	H	1.000	0.077	0.046	0.521	BSX	H	1.000	0.078	0.048	0.518
CX	U	1.000	0.111	0.089	0.908	CX	U	1.000	0.107	0.085	0.911	CX	U	1.000	0.107	0.085	0.918
CX	U	1.000	0.116	0.100	0.918	CX	U	1.000	0.112	0.097	0.926	CX	U	1.000	0.112	0.098	0.918
CAH	H	1.000	0.100	0.095	0.829	CAH	H	1.000	0.105	0.102	0.918	CAH	H	1.000	0.107	0.093	0.918
CERN	T	1.000	0.129	0.130	0.908	CERN	T	1.000	0.123	0.122	0.921	CERN	T	1.000	0.121	0.119	0.921
CHK	H	1.000	0.107	0.081	0.868	CHK	H	1.000	0.101	0.076	0.874	CI	H	1.000	0.081	0.039	0.661
CI	H	1.000	0.083	0.040	0.658	CI	H	1.000	0.080	0.039	0.661	CCI	T	1.000	0.074	0.035	0.550
CCI	T	1.000	0.078	0.037	0.579	CCI	T	1.000	0.074	0.034	0.571	DHI	CD	1.000	0.097	0.090	0.734
DHI	CD	1.000	0.096	0.052	0.727	DHI	CD	1.000	0.088	0.046	0.685	FDO	CD	1.000	0.065	0.020	0.418
FDO	CD	1.000	0.068	0.021	0.418	FDO	CD	1.000	0.065	0.019	0.424	FTTB	F	1.000	0.070	0.023	0.461
FTTB	F	1.000	0.073	0.024	0.474	FTTB	F	1.000	0.070	0.023	0.471	HD	CD	1.000	0.100	0.074	0.868
HD	CD	1.000	0.104	0.078	0.897	HD	CD	1.000	0.099	0.073	0.879	ILMN	H	1.000	0.080	0.046	0.605
ILMN	H	1.000	0.080	0.050	0.652	ILMN	H	1.000	0.085	0.052	0.705	MUR	E	1.000	0.102	0.085	0.739
ILMN	H	1.000	0.092	0.055	0.787	LNC	F	1.000	0.089	0.053	0.784	MUR	E	1.000	0.109	0.090	0.929
LLTC	F	1.000	0.152	0.179	0.989	LLTC	F	1.000	0.149	0.178	0.992	PRU	F	1.000	0.080	0.035	0.679
MUR	E	1.000	0.115	0.098	0.921	MUR	E	1.000	0.110	0.093	0.926	SWN	E	1.000	0.096	0.066	0.826
PRU	F	1.000	0.082	0.036	0.703	PRU	F	1.000	0.079	0.034	0.700	SWKS	T	1.000	0.128	0.133	0.921
PRU	F	1.000	0.106	0.070	0.811	SWN	E	1.000	0.096	0.067	0.831	TIF	I	1.000	0.115	0.108	0.911
SWKS	T	1.000	0.131	0.134	0.932	SWKS	T	1.000	0.128	0.133	0.942	TROW	F	1.000	0.090	0.051	0.816
HOT	CD	1.000	0.121	0.113	0.929	SYK	H	1.000	0.107	0.092	0.858	TIF	CD	1.000	0.106	0.085	0.942
SYK	H	1.000	0.110	0.094	0.850	TROW	F	1.000	0.090	0.052	0.826	UTX	I	1.000	0.115	0.108	0.961
TROW	F	1.000	0.093	0.053	0.826	TIF	CD	1.000	0.106	0.086	0.937	WYNN	CD	1.000	0.113	0.089	0.921
TIF	CD	1.000	0.110	0.089	0.933	UTX	I	1.000	0.116	0.100	0.915	WYNN	CD	1.000	0.112	0.078	0.888
UTX	I	1.000	0.120	0.113	0.958	WYNN	CD	1.000	0.113	0.100	0.932	BBY	CD	1.000	0.107	0.091	0.911
WYNN	CD	1.000	0.118	0.103	0.929	KEY	F	1.000	0.072	0.032	0.476	CHK	E	0.997	0.101	0.076	0.868
CSX	CD	0.992	0.122	0.118	0.944	CSX	CD	0.992	0.119	0.116	0.944	LNC	F	0.997	0.090	0.053	0.802
CSX	CD	0.989	0.080	0.039	0.526	HOT	CD	0.986	0.106	0.091	0.904	TXN	T	0.995	0.167	0.195	0.995
MCD	CD	0.989	0.083	0.040	0.670	HSY	CS	0.989	0.076	0.037	0.513	TXN	T	0.992	0.134	0.146	0.976
LMT	I	0.987	0.093	0.061	0.744	BBY	CD	0.989	0.106	0.091	0.904	HOT	CD	0.992	0.117	0.110	0.926
NFLX	F	0.979	0.120	0.116	0.820	TXN	T	0.976	0.133	0.145	0.976	HSY	CS	0.987	0.076	0.036	0.504
KEY	F	0.966	0.074	0.033	0.488	NFLX	T	0.968	0.115	0.112	0.821	NFLX	T	0.976	0.113	0.108	0.814
NFLX	F	0.947	0.128	0.128	0.937	CSX	CD	0.967	0.066	0.033	0.468	CSX	CD	0.953	0.086	0.037	0.504
MAS	H	0.945	0.132	0.132	0.916	LMT	I	0.921	0.090	0.060	0.749	TYC	I	0.937	0.088	0.057	0.784
BHI	E	0.921	0.131	0.131	0.934	MCD	CD	0.928	0.077	0.037	0.690	ALL	F	0.903	0.075	0.034	0.653

Sampling Frequency: 10 Minute																	
EN-1						EN-2						Lasso					
Ticker	Stock	Sector	Freq.	PCA	SPCA	Ticker	Stock	Sector	Freq.	PCA	SPCA	Ticker	Stock	Sector	Freq.	PCA	SPCA
AEP	U	1.000	0.105	0.068	0.639	AEP	U	1.000	0.108	0.071	0.674	AEP	U	1.000	0.110	0.075	0.679
BBBY	CD	1.000	0.166	0.160	0.937	BBBY	CD	1.000	0.170	0.165	0.937	BBBY	CD	1.000	0.167	0.162	0.934
CBS	CD	1.000	0.162	0.156	0.916	CBS	CD	1.000	0.165	0.158	0.913	CBS	CD	1.000	0.164	0.157	0.908
CX	U	1.000	0.140	0.127	0.887	CX	U	1.000	0.142	0.134	0.871	CX	U	1.000	0.143	0.137	0.871
CERN	T	1.000	0.166	0.160	0.882</												

*Notes: See notes to Table 3.23.

Table 3.25: Factor Structure (IBM)

Sampling Frequency: 1 Minute																																		
Stock					EN-1					Stock					EN-2					Stock					Lasso					SPCA				
Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA					
ADM	CS	1.000	0.103	0.040	0.587	AFL	F	1.000	0.140	0.112	0.961	AFL	F	1.000	0.141	0.113	0.961	AFL	F	1.000	0.141	0.113	0.961	AFL	F	1.000	0.141	0.113	0.961					
ALL	F	1.000	0.137	0.110	0.966	ALL	F	1.000	0.111	0.051	0.726	ALL	F	1.000	0.111	0.051	0.726	ALL	F	1.000	0.111	0.051	0.726	ALL	F	1.000	0.111	0.051	0.726					
AMT	T	1.000	0.108	0.049	0.721	AMT	T	1.000	0.117	0.062	0.824	AMT	T	1.000	0.118	0.063	0.824	AMT	T	1.000	0.118	0.063	0.824	AMT	T	1.000	0.118	0.063	0.824					
ABC	H	1.000	0.114	0.062	0.829	ABC	H	1.000	0.137	0.118	0.755	ABC	H	1.000	0.138	0.117	0.758	ABC	H	1.000	0.138	0.117	0.758	ABC	H	1.000	0.138	0.117	0.758					
ABN	H	1.000	0.133	0.115	0.761	ABN	H	1.000	0.137	0.107	0.834	ABN	H	1.000	0.138	0.109	0.826	ABN	H	1.000	0.138	0.109	0.826	ABN	H	1.000	0.138	0.109	0.826					
BBY	H	1.000	0.130	0.107	0.829	BBY	H	1.000	0.150	0.137	0.955	BBY	H	1.000	0.151	0.139	0.958	BBY	H	1.000	0.151	0.139	0.958	BBY	H	1.000	0.151	0.139	0.958					
CPB	CS	1.000	0.146	0.137	0.963	CPB	CS	1.000	0.111	0.058	0.666	CPB	CS	1.000	0.112	0.058	0.653	CPB	CS	1.000	0.112	0.058	0.653	CPB	CS	1.000	0.112	0.058	0.653					
CAH	H	1.000	0.109	0.058	0.674	CAH	H	1.000	0.146	0.131	0.861	CAH	H	1.000	0.148	0.132	0.855	CAH	H	1.000	0.148	0.132	0.855	CAH	H	1.000	0.148	0.132	0.855					
CI	H	1.000	0.142	0.129	0.874	CI	H	1.000	0.118	0.063	0.803	CI	H	1.000	0.118	0.062	0.800	CI	H	1.000	0.118	0.062	0.800	CI	H	1.000	0.118	0.062	0.800					
CL	H	1.000	0.115	0.061	0.824	CAG	CS	1.000	0.120	0.080	0.632	CAG	CS	1.000	0.121	0.081	0.634	CAG	CS	1.000	0.121	0.081	0.634	CAG	CS	1.000	0.121	0.081	0.634					
CIG	CS	1.000	0.117	0.077	0.624	GE	I	1.000	0.146	0.129	0.937	GE	I	1.000	0.146	0.130	0.929	GE	I	1.000	0.146	0.130	0.929	GE	I	1.000	0.146	0.130	0.929					
GE	I	1.000	0.142	0.125	0.934	GRMN	CD	1.000	0.126	0.085	0.750	GRMN	CD	1.000	0.125	0.085	0.774	GRMN	CD	1.000	0.125	0.085	0.774	GRMN	CD	1.000	0.125	0.085	0.774					
GRMN	CD	1.000	0.122	0.082	0.742	HBAN	F	1.000	0.070	0.012	0.189	HBAN	F	1.000	0.071	0.013	0.189	HBAN	F	1.000	0.071	0.013	0.189	HBAN	F	1.000	0.071	0.013	0.189					
HBAN	F	1.000	0.070	0.013	0.176	ILMN	H	1.000	0.115	0.063	0.629	ILMN	H	1.000	0.116	0.064	0.642	ILMN	H	1.000	0.116	0.064	0.642	ILMN	H	1.000	0.116	0.064	0.642					
ILMN	H	1.000	0.113	0.063	0.653	KR	CS	1.000	0.099	0.040	0.495	KR	CS	1.000	0.101	0.042	0.497	KR	CS	1.000	0.101	0.042	0.497	KR	CS	1.000	0.101	0.042	0.497					
KR	CS	1.000	0.097	0.039	0.518	LM	F	1.000	0.103	0.036	0.639	LM	F	1.000	0.104	0.037	0.663	LM	F	1.000	0.104	0.037	0.663	LM	F	1.000	0.104	0.037	0.663					
LM	F	1.000	0.101	0.037	0.684	LNC	F	1.000	0.127	0.081	0.879	LNC	F	1.000	0.128	0.081	0.892	LNC	F	1.000	0.128	0.081	0.892	LNC	F	1.000	0.128	0.081	0.892					
LNC	F	1.000	0.124	0.078	0.879	MMC	F	1.000	0.100	0.037	0.629	MMC	F	1.000	0.101	0.038	0.613	MMC	F	1.000	0.101	0.038	0.613	MMC	F	1.000	0.101	0.038	0.613					
MMC	F	1.000	0.098	0.036	0.621	MON	M	1.000	0.107	0.045	0.616	MON	M	1.000	0.108	0.045	0.626	MON	M	1.000	0.108	0.045	0.626	MON	M	1.000	0.108	0.045	0.626					
MON	M	1.000	0.104	0.045	0.618	PEP	CS	1.000	0.122	0.075	0.771	PEP	CS	1.000	0.124	0.077	0.779	PEP	CS	1.000	0.124	0.077	0.779	PEP	CS	1.000	0.124	0.077	0.779					
PEP	CS	1.000	0.120	0.076	0.789	PRU	F	1.000	0.113	0.052	0.784	PRU	F	1.000	0.114	0.052	0.782	PRU	F	1.000	0.114	0.052	0.782	PRU	F	1.000	0.114	0.052	0.782					
PRU	F	1.000	0.111	0.051	0.795	SWN	E	1.000	0.129	0.091	0.855	SWN	E	1.000	0.130	0.091	0.858	SWN	E	1.000	0.130	0.091	0.858	SWN	E	1.000	0.130	0.091	0.858					
SWN	E	1.000	0.126	0.089	0.855	SBUX	CD	1.000	0.165	0.168	0.989	SBUX	CD	1.000	0.167	0.169	0.989	SBUX	CD	1.000	0.167	0.169	0.989	SBUX	CD	1.000	0.167	0.169	0.989					
SBUX	CD	1.000	0.161	0.165	0.992	TIF	CD	1.000	0.142	0.117	0.974	TIF	CD	1.000	0.143	0.117	0.976	TIF	CD	1.000	0.143	0.117	0.976	TIF	CD	1.000	0.143	0.117	0.976					
TIF	CD	1.000	0.138	0.116	0.971	UNP	I	1.000	0.150	0.135	0.968	UNP	I	1.000	0.151	0.137	0.963	UNP	I	1.000	0.151	0.137	0.963	UNP	I	1.000	0.151	0.137	0.963					
UNP	I	1.000	0.147	0.135	0.963	UNM	F	1.000	0.133	0.098	0.903	UNM	F	1.000	0.134	0.099	0.900	UNM	F	1.000	0.134	0.099	0.900	UNM	F	1.000	0.134	0.099	0.900					
UNM	F	1.000	0.130	0.096	0.897	LLTC	T	1.000	0.196	0.245	0.995	LLTC	T	1.000	0.198	0.246	0.995	LLTC	T	1.000	0.198	0.246	0.995	LLTC	T	1.000	0.198	0.246	0.995					
LLTC	T	0.997	0.191	0.239	0.997	EXPE	CD	1.000	0.180	0.204	0.989	EXPE	CD	1.000	0.181	0.206	0.992	EXPE	CD	1.000	0.181	0.206	0.992	EXPE	CD	1.000	0.181	0.206	0.992					
CNX	F	0.989	0.143	0.123	0.968	CNX	U	0.984	0.147	0.127	0.965	CNX	U	0.987	0.147	0.126	0.974	CNX	U	0.987	0.147	0.126	0.974	CNX	U	0.987	0.147	0.126	0.974					
SCHW	F	0.987	0.095	0.031	0.496	ADM	CS	0.945	0.106	0.042	0.604	TXN	T	0.916	0.186	0.220	0.994	TXN	T	0.916	0.186	0.220	0.994	TXN	T	0.916	0.186	0.220	0.994					
GS	F	0.966	0.096	0.026	0.529	TXN	T	0.926	0.184	0.220	0.994	AXP	F	0.871	0.138	0.105	0.964	AXP	F	0.871	0.138	0.105	0.964	AXP	F	0.871	0.138	0.105	0.964					
CD	0.945	0.154	0.149	0.964	AXP	F	0.861	0.136	0.103	0.969	GS	F	0.853	0.094	0.020	0.485	GS	F	0.853	0.094	0.020	0.485	GS	F	0.853	0.094	0.020	0.485						
EXPE	CD	0.937	0.175	0.201	0.994	GS	F	0.845	0.096	0.023	0.526	ENDP	H	0.845	0.129	0.099	0.929	ENDP	H	0.845	0.129	0.099	0.929	ENDP	H	0.845	0.129	0.099	0.929					
TXN	T	0.911	0.178	0.208	0.997	HOT	CD	0.834	0.159	0.152	0.962	HOT	CD	0.808	0.161	0.154	0.967	HOT	CD	0.808	0.161	0.154	0.967	HOT	CD	0.808	0.161	0.154	0.967					

Sampling Frequency: 5 Minute																																		
Stock					EN-1					Stock					EN-2					Stock					Lasso					SPCA				
Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA	Ticker	Sector	Freq.	PCA	SPCA					
ALL	F	1.000	0.105	0.050	0.684	ALL	F	1.000	0.106	0.050	0.708	ALL	F	1.000	0.106	0.051	0.700	ALL	F	1.000	0.106	0.051	0.700	ALL	F	1.000	0.106	0.051	0.700					
BK	F	1.000	0.090	0.025	0.508	BK	F	1.000	0.091	0.025	0.492	BK	F	1.000	0.091	0.024	0.487	BK	F	1.000	0.091	0.024	0.487	BK	F	1.000	0.091	0.024	0.487					
BBY	CD	1.000	0.147	0.130	0.937	BBY	CD	1.000	0.149	0.132	0.939	BBY	CD	1.000	0.150	0.132	0.945	BBY	CD	1.000	0.150	0.132	0.945	BBY	CD	1.000	0.150	0.132	0.945					
CNX	U	1.000	0.141	0.114	0.908	CNX	U	1.000	0.142	0.113	0.905	CNX	U	1.000	0.142	0.113	0.892	CNX	U	1.000	0.142	0.113	0.892	CNX	U	1.000	0.142	0.113	0.892					
CSX	I	1.000	0.149	0.132	0.942	CSX	I	1.000	0.151	0.133	0.945	CSX	I	1.000	0.150	0.131	0.942	CSX	I	1.000	0.150	0.131	0.942	CSX	I	1.000	0.150	0.131	0.942					
CPB	CS	1.000	0.110	0.066	0.618	CAH	H	1.000	0.147	0.130	0.855	CAH	H	1.000	0.147	0.131	0.842	CAH	H	1.000	0.147	0.131	0.842	CAH	H	1.000	0.147	0.131	0.842					
CAH	H	1.000	0.148	0.135	0.847	CERN	T	1.000	0.169	0.173	0.924	CERN	T	1.000	0.169	0.172	0.913	CERN	T	1.000	0.169													

Bibliography

- AIOLFI, M., RODRIGUEZ, M., AND TIMMERMAN, A. 2009. Understanding analysts earnings expectations: Biases, nonlinearities, and predictability. *Journal of Financial Econometrics* 8:305–334.
- AÏT-SAHALIA, Y. 2002a. Maximum likelihood estimation of discretely sampled diffusions: A closed-form approximation approach. *Econometrica* 70:223–262.
- AÏT-SAHALIA, Y. 2002b. Telling from discrete data whether the underlying continuous-time model is a diffusion. *The Journal of Finance* 57:2075–2112.
- AÏT-SAHALIA, Y., CACHO-DIAZ, J., AND LAEVEN, R. J. 2015. Modeling financial contagion using mutually exciting jump processes. *Journal of Financial Economics* 117:585–606.
- AÏT-SAHALIA, Y., FAN, J., LAEVEN, R. J., WANG, C. D., AND YANG, X. 2016. Estimation of the continuous and discontinuous leverage effects. *Journal of the American Statistical Association* .
- AÏT-SAHALIA, Y. AND JACOD, J. 2009. Testing for jumps in a discretely observed process. *The Annals of Statistics* 37:184–222.
- AÏT-SAHALIA, Y. AND JACOD, J. 2011. Testing whether jumps have finite or infinite activity. *the Annals of Statistics* 39:1689–1719.
- AÏT-SAHALIA, Y. AND JACOD, J. 2012. Analyzing the spectrum of asset returns: Jump and volatility components in high frequency data. *Journal of Economic Literature* 50:1007–1050.
- AÏT-SAHALIA, Y. AND JACOD, J. 2014. High-frequency financial econometrics. Princeton University Press: Princeton, NJ, USA.
- AÏT-SAHALIA, Y., JACOD, J., AND LI, J. 2012. Testing for jumps in noisy high frequency data. *Journal of Econometrics* 168:207–222.

- AÏT-SAHALIA, Y., MYKLAND, P. A., AND ZHANG, L. 2011. Ultra high frequency volatility estimation with dependent microstructure noise. *Journal of Econometrics* 160:160–175.
- AÏT-SAHALIA, Y. AND XIU, D. 2017a. Principal component analysis of high frequency data. *Journal of the American Statistical Association*, *forthcoming* .
- AÏT-SAHALIA, Y. AND XIU, D. 2017b. Using principal component analysis to estimate a high dimensional factor model with high-frequency data. *Journal of Econometrics* 201:384–399.
- ANDERSEN, T., BOLLERSLEV, T., DIEBOLD, F. X., AND LABYS, P. 2001. The distribution of realized exchange rate volatility. *Journal of the American Statistical Association* 96:42–55.
- ANDERSEN, T. G., BOLLERSLEV, T., AND DIEBOLD, F. X. 2007a. Roughing it up: Including jump components in the measurement, modeling, and forecasting of return volatility. *The review of economics and statistics* 89:701–720.
- ANDERSEN, T. G., BOLLERSLEV, T., DIEBOLD, F. X., AND LABYS, P. 2003. Modeling and forecasting realized volatility. *Econometrica* 71:579–625.
- ANDERSEN, T. G., BOLLERSLEV, T., AND DOBREV, D. 2007b. No-arbitrage semimartingale restrictions for continuous-time volatility models subject to leverage effects, jumps and iid noise: Theory and testable distributional implications. *Journal of Econometrics* 138:125–180.
- ANDERSEN, T. G., BOLLERSLEV, T., AND MEDDAHI, N. 2004. Analytical evaluation of volatility forecasts. *International Economic Review* 45:1079–1110.
- ANDERSEN, T. G., BOLLERSLEV, T., AND MEDDAHI, N. 2011. Realized volatility forecasting and market microstructure noise. *Journal of Econometrics* 160:220–234.
- ANDREWS, D. W. AND CHENG, X. 2012. Estimation and inference with weak, semi-strong, and strong identification. *Econometrica* 80:2153–2211.

- ANG, A. AND TIMMERMANN, A. 2012. Regime changes and financial markets. *Annu. Rev. Financ. Econ.* 4:313–337.
- AUDRINO, F. AND HU, Y. 2016. Volatility forecasting: Downside risk, jumps and leverage effect. *Econometrics* 4:1–24.
- BAI, J. AND NG, S. 2006a. Confidence intervals for diffusion index forecasts and inference for factor-augmented regressions. *Econometrica* 74:1133–1150.
- BAI, J. AND NG, S. 2006b. Evaluating latent and observed factors in macroeconomics and finance. *Journal of Econometrics* 131:507–537.
- BAI, J. AND NG, S. 2008. Forecasting economic time series using targeted predictors. *Journal of Econometrics* 146:304–317.
- BARNDORFF-NEILSEN, O. E. AND SHEPHARD, N. 2003. Realised power variation and stochastic volatility variance. *Bernoulli* 9:243–265.
- BARNDORFF-NIELSEN, O. E. 2002. Econometric analysis of realized volatility and its use in estimating stochastic volatility models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 64:253–280.
- BARNDORFF-NIELSEN, O. E., GRAVERSEN, S. E., JACOD, J., AND SHEPHARD, N. 2006. Limit theorems for bipower variation in financial econometrics. *Econometric Theory* 22:677–719.
- BARNDORFF-NIELSEN, O. E. AND SHEPHARD, N. 2004. Power and bipower variation with stochastic volatility and jumps. *Journal of financial econometrics* 2:1–37.
- BARNDORFF-NIELSEN, O. E. AND SHEPHARD, N. 2006. Econometrics of testing for jumps in financial economics using bipower variation. *Journal of financial Econometrics* 4:1–30.
- BENJAMINI, Y. AND HOCHBERG, Y. 1995. Controlling the false discovery rate: a practical and powerful approach to multiple testing. *Journal of the royal statistical society. Series B (Methodological)* 57:289–300.

- BOLLERSLEV, T., PATTON, A. J., AND QUAEDVLIEG, R. 2016. Exploiting the errors: A simple approach for improved volatility forecasting. *Journal of Econometrics* 192:1–18.
- BOWSER, C. G. 2007. Modelling security market events in continuous time: Intensity based, multivariate point process models. *Journal of Econometrics* 141:876–912.
- BRANDT, M. W. AND JONES, C. S. 2006. Volatility forecasting with range-based egarch models. *Journal of Business & Economic Statistics* 24:470–486.
- CAMPBELL, J. Y. AND THOMPSON, S. B. 2008. Predicting excess stock returns out of sample: Can anything beat the historical average? *Review of Financial Studies* 21:1509–1531.
- CARRASCO, M. AND ROSSI, B. 2016. In-sample inference and forecasting in misspecified factor models. *Journal of Business and Economic Statistics* 34:313–338.
- CHERNOV, M., GALLANT, A. R., GHYSELS, E., AND TAUCHEN, G. 2003. Alternative models for stock price dynamics. *Journal of Econometrics* 116:225–257.
- CORRADI, V., SILVAPULLE, M. J., AND SWANSON, N. R. 2014. Consistent pretesting for jumps. *Working Paper, University of Surrey, Monash University and Rutgers University*.
- CORRADI, V., SILVAPULLE, M. J., AND SWANSON, N. R. 2018. Testing for jumps and jump intensity path dependence. *Journal of Econometrics, forthcoming*.
- CORSI, F. 2009. A simple approximate long-memory model of realized volatility. *Journal of Financial Econometrics* 7:174–196.
- CORSI, F., PIRINO, D., AND RENO, R. 2009. Volatility forecasting: The jumps do matter. *Working Paper*.
- CORSI, F., PIRINO, D., AND RENO, R. 2010. Threshold bipower variation and the impact of jumps on volatility forecasting. *Journal of Econometrics* 159:276–288.

- DUMITRU, A.-M. AND URGAS, G. 2012. Identifying jumps in financial assets: A comparison between nonparametric jump tests. *Journal of Business & Economic Statistics* 30:242–255.
- DUONG, D. AND SWANSON, N. R. 2015. Empirical evidence on the importance of aggregation, asymmetry, and jumps for volatility prediction. *Journal of Econometrics* 187:606–621.
- GHYSELS, E., SANTA-CLARA, P., AND VALKANOV, R. 2006. Predicting volatility: getting the most out of return data sampled at different frequencies. *Journal of Econometrics* 131:59–95.
- GHYSELS, E. AND SINKO, A. 2011. Volatility forecasting and microstructure noise. *Journal of Econometrics* 160:257–271.
- HANSEN, P. R. AND LUNDE, A. 2005. A forecast comparison of volatility models: does anything beat a garch (1, 1)? *Journal of applied econometrics* 20:873–889.
- HOLM, S. 1979. A simple sequentially rejective multiple test procedure. *Scandinavian journal of statistics* 6:65–70.
- HUANG, X. AND TAUCHEN, G. 2005. The relative contribution of jumps to total price variance. *Journal of financial econometrics* 3:456–499.
- JACOD, J., LI, Y., MYKLAND, P. A., PODOLSKIJ, M., AND VETTER, M. 2009. Microstructure noise in the continuous case: the pre-averaging approach. *Stochastic processes and their applications* 119:2249–2276.
- JACOD, J. AND PROTTER, P. 2011. Discretization of processes, volume 67. Springer Science & Business Media.
- JACOD, J. AND ROSENBAUM, M. 2013. Quarticity and other functionals of volatility: Efficient estimation. *The Annals of Statistics* 41:1462–1484.
- KALNINA, I. AND XIU, D. 2017. Nonparametric estimation of the leverage effect: A trade-off between robustness and efficiency. *Journal of the American Statistical Association* 112:384–396.

- KIM, H. H. AND SWANSON, N. R. 2017. Mining big data using parsimonious factor, machine learning, variable selection, and shrinkage methods. *International Journal of Forecasting*, forthcoming .
- LEE, S. S. AND MYKLAND, P. A. 2008. Jumps in financial markets: A new nonparametric test and jump dynamics. *Review of Financial studies* 21:2535–2563.
- MANCINI, C. 2001. Disentangling the jumps of the diffusion in a geometric jumping brownian motion. *Giornale dell'Istituto Italiano degli Attuari* LXIV:19–47.
- MANCINI, C. 2009. Non-parametric threshold estimation for models with stochastic diffusion coefficient and jumps. *Scandinavian Journal of Statistics* 36:270–296.
- MEDDAHI, N. 2001. An eigenfunction approach for volatility modeling. CIRANO.
- PATTON, A. J. AND SHEPPARD, K. 2015. Good volatility, bad volatility: Signed jumps and the persistence of volatility. *Review of Economics and Statistics* 97:683–697.
- PAYE, B. S. AND TIMMERMANN, A. 2006. Instability of return prediction models. *Journal of Empirical Finance* 13:274–315.
- PODOLSKIJ, M. AND VETTER, M. 2009a. Bipower-type estimation in a noisy diffusion setting. *Stochastic processes and their applications* 119:2803–2831.
- PODOLSKIJ, M. AND VETTER, M. 2009b. Estimation of volatility functionals in the simultaneous presence of microstructure noise and jumps. *Bernoulli* 15:634–658.
- PODOLSKIJ, M. AND ZIGGEL, D. 2010. New tests for jumps in semimartingale models. *Statistical inference for stochastic processes* 13:15–41.
- QI, X., LUO, R., AND ZHAO, H. 2013. Sparse principal component analysis by choice of norm. *Journal of Multivariate Analysis* 114:127–160.
- ROMANO, J. P. AND WOLF, M. 2005. Stepwise multiple testing as formalized data snooping. *Econometrica* 73:1237–1282.

- STOCK, J. H. AND WATSON, M. W. 2002a. Forecasting using principal components from a large number of predictors. *Journal of the American Statistical Association* 97:1167–1179.
- STOCK, J. H. AND WATSON, M. W. 2002b. Macroeconomic forecasting using diffusion indexes. *Journal of Business & Economic Statistics* 20:147–162.
- STOCK, J. H. AND WATSON, M. W. 2006. Forecasting with many predictors. *Handbook of Economic Forecasting* 1:515–554.
- STOREY, J. D. 2003. The positive false discovery rate: a bayesian interpretation and the q-value. *The Annals of Statistics* 31:2013–2035.
- SWANSON, N. R. AND XIONG, W. 2017. Big data analytics in economics: what have we learned so far, and where should we go from here? *Working Paper, Rutgers University*.
- THEODOSIOU, M. AND ZIKES, F. 2009. A comprehensive comparison of alternative tests for jumps in asset prices. *Working Paper, Imperial College London*.
- TIBSHIRANI, R. 1996. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society. Series B (Methodological)* pp. 267–288.
- TODOROV, V. 2015. Jump activity estimation for pure-jump semimartingales via self-normalized statistics. *The Annals of Statistics* 43:1831–1864.
- TODOROV, V. AND TAUCHEN, G. 2011. Volatility jumps. *Journal of Business & Economic Statistics* 29:356–371.
- WHITE, H. 2000. A reality check for data snooping. *Econometrica* 68:1097–1126.
- ZOU, H. AND HASTIE, T. 2005. Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 67:301–320.
- ZOU, H., HASTIE, T., AND TIBSHIRANI, R. 2006. Sparse principal component analysis. *Journal of Computational and Graphical Statistics* 15:265–286.