

NBER WORKING PAPER SERIES

APT OR “AIPT”? THE SURPRISING DOMINANCE OF LARGE FACTOR MODELS

Antoine Didisheim
Shikun (Barry) Ke
Bryan T. Kelly
Semyon Malamud

Working Paper 33012
<http://www.nber.org/papers/w33012>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
September 2024, Revised August 2025

A previously circulated version of this article was titled “Complexity in Factor Pricing Models.” Antoine Didisheim is at the University of Melbourne. Shikun Ke is at Yale School of Management. Bryan Kelly is at Yale School of Management, AQR Capital Management, and NBER; www.bryankellyacademic.org. Semyon Malamud is at Swiss Finance Institute, EPFL, and CEPR, and is a consultant to AQR. We are grateful for helpful comments from Federico Baldi-Lanfranchi, John Campbell, Mike Chernov, Xavier Gabaix, Martin Lettau, Stefan Nagel, Mohammad Pourmohammadi, Seth Pruitt, Mirela Sandulescu, Paul Schneider, Fabio Trojani, Neng Wang, and participants at many seminars and conferences. Semyon Malamud gratefully acknowledges the financial support of the Swiss Finance Institute and the Swiss National Science Foundation, Grant 100018 192692. AQR Capital Management is a global investment management firm that may or may not apply similar investment techniques or methods of analysis as described herein. The views expressed here are those of the authors and not necessarily those of AQR or the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2024 by Antoine Didisheim, Shikun (Barry) Ke, Bryan T. Kelly, and Semyon Malamud. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

APT or “AIPT”? The Surprising Dominance of Large Factor Models
Antoine Didisheim, Shikun (Barry) Ke, Bryan T. Kelly, and Semyon Malamud
NBER Working Paper No. 33012
September 2024, Revised August 2025
JEL No. C1, C40, C45, C55, G1, G11, G12, G14, G17

ABSTRACT

We introduce artificial intelligence pricing theory (AIPT). In contrast with the APT’s foundational assumption of a low dimensional factor structure in returns, the AIPT conjectures that returns are driven by a large number of factors. We first verify this conjecture empirically and show that complex nonlinear models with an exorbitant number of factors (many more than the number of training observations or base assets) are far more successful in describing the out-of-sample behavior of asset returns than simpler benchmark models. Economically, we show that complex models are more than twice as sensitive to long-run macroeconomic activity than simpler benchmarks. Finally, we theoretically characterize the behavior of large factor pricing models, from which we show that the AIPT’s “many factors” conjecture faithfully explains our empirical findings, while the APT’s “few factors” conjecture is contradicted by the data.

Antoine Didisheim
Department of Finance, University
of Lausanne
Department of Finance
antoine.didisheim@unimelb.edu.au

Shikun (Barry) Ke
Yale University
Yale School of Management
barry.ke@yale.edu

Bryan T. Kelly
Yale University
and NBER
bryan.kelly@yale.edu

Semyon Malamud
Swiss Finance Institute
semyon.malamud@epfl.ch

1 Introduction

Ross’s (1976) APT conjectures that a small number of common factors govern the joint variation of returns. This premise, combined with no-arbitrage arguments, delivers the prediction that assets’ expected returns are determined by exposures to a few common factors. Most empirical investigations of the risk-return tradeoff in the past fifty years have occurred within the confines of the APT’s assumption—small linear factor models.

In this paper, we consider a different conjecture: that asset pricing models with an exorbitant number of factors are better suited to describe the behavior of asset returns. Our conjecture is rooted in the burgeoning theory of artificial intelligence, which suggests that “complex” statistical models, those with many more parameters (P) than available training observations (T), tend to achieve better out-of-sample success than smaller models.¹ Given its origins, we refer to this conjecture as artificial intelligence pricing theory, or AIPT, to contrast it with the parsimony assumption at the heart of the APT.

One reason for the APT’s position as a cornerstone of empirical asset pricing is the convenience of working with small linear factor models. The most obvious convenience stems from the tractability of estimating small models. More subtle is the convenience of *comparing* small models. The APT sets up conditions necessary for comparing models based on *in-sample* test statistics (namely, that the number of parameters is small compared to the number of training observations). Large factor models, on the other hand, are admittedly inconvenient. With so many parameters, they can be costly to train, and due to their high

¹See, among others, [Hastie et al. \(2019\)](#), [Bartlett et al. \(2020\)](#), and [Kelly et al. \(2024b\)](#). Following these papers, we define model complexity as the ratio of factor dimension to the number of training observations $c = P/T$. Our theoretical derivations rely crucially on the ratio c , hence our focus on this definition of complexity. The statistics literature proposes and analyzes other measures of model complexity as well, each capturing somewhat different qualities. As noted by [Bartlett et al. \(2020\)](#): “The classical perspective in statistical learning theory is that there should be a tradeoff between the fit to the training data and the complexity of the prediction rule. Whether complexity is measured in terms of the number of parameters, the number of nonzero parameters in a high-dimensional setting, the number of neighbors averaged in a nearest neighbor estimator, the scale of an estimate in a reproducing kernel Hilbert space, or the bandwidth of a kernel smoother, this tradeoff has been ubiquitous in statistical learning theory.”

statistical complexity, standard in-sample asset pricing tests are not applicable. Instead, they require out-of-sample model performance comparisons.

We show that, despite their inconveniences, large factor models are far better at pricing assets than the standard low-dimensional factor models in the literature. In the first half of this paper, we present extensive empirical evidence that challenges the APT parsimony conjecture and favors the complexity conjecture of AIPT. In the second half of the paper we develop a statistical theory that characterizes the behavior of complex factor models and rationalizes their surprising dominance in pricing assets.

1.1 Empirical Behavior of Large Factor Models

Our central empirical finding is that the out-of-sample performance of factor pricing models is increasing in the number of factors. We focus on two metrics of model performance—the Sharpe ratio of the tangency portfolio among factors and the magnitude of pricing errors among test assets.

Our analysis gradually increases the number of factors while holding the underlying information set fixed. The information set includes 130 well-known characteristics that predict the cross-section of US stocks. To vary the number of pricing factors, we use a neural network to generate new stock characteristics—up to hundreds of thousands—that are nonlinear transformations of the original 130 variables. For each nonlinear characteristic, we build the corresponding factor (using a standard characteristic-managed portfolio approach). Training the factor model amounts to building the stochastic discount factor (SDF) as the tangency portfolio of factors (using ridge shrinkage where applicable). We then track the complex SDF’s performance out-of-sample, presenting our main empirical results in the form of “VoC curves” (Kelly et al., 2024b, KMZ henceforth) that plot performance as a function of the number of factors. These curves document a “virtue of complexity” in factor pricing models, with out-of-sample Sharpe ratios increasing and pricing errors decreasing

as we expand the number of factors. The complex SDF performs well even when—in fact, especially when— P far exceeds T .

Large factor models outperform well-known, low-dimensional factor models from the literature by a wide margin. For example, the largest model we consider (with 360,000 factors, or a complexity ratio of 1,000 given our 360-month rolling training window) reduces pricing errors by 54.8% relative to the six-factor [Fama and French \(2015\)](#) model (including momentum) and by 50.4% relative to the five-factor [Hou et al. \(2021\)](#) model. Meanwhile, our largest factor model produces an out-of-sample tangency Sharpe ratio of 3.7, compared to 0.8 and 1.2 for the aforementioned benchmarks, respectively.²

Improvements in SDF measurement achieved by large factor models offer a window into the economic nature of the unobservable “true” SDF. We argue that the higher an SDF’s out-of-sample Sharpe ratio, the better it reflects the drivers of investors’ intertemporal marginal rate of substitution, and thus the better it should reflect expectations of future economic conditions. In other words, portfolio efficiency and informational efficiency go hand-in-hand.

We estimate a macroeconomic vector autoregression and find that the complex SDF is an unusually strong forecaster of macroeconomic conditions as far as three years into the future. For example, a positive one standard deviation annual return to the complex SDF predicts a roughly 3% rise in industrial production over three years and a reduction in macroeconomic uncertainty of about 7.5% over two years. A positive SDF return also predicts significant increases in employment, consumption, consumer sentiment, and housing over horizons of up to three years, and it predicts a decrease in credit spreads and inflation.

The asset pricing benefits of complexity and nonlinearity accrue even when the conditioning information set is relatively small. For example, we construct a “nonlinear Fama-French model” using a large set of nonlinear factors derived from *only* five characteristics: size, value, investment, profitability, and momentum, rather than the full set of 130 characteristics in our main analysis. The added complexity more than doubles the out-of-sample tangency Sharpe

²We study a number of additional benchmark models as well.

ratio relative to the baseline Fama-French model. The same result obtains for complex nonlinear versions of other benchmark models in the literature.

The insights of [Kozak et al. \(2020\)](#) suggest that a successful asset pricing model may not require many factors because the returns of most anomaly factors are adequately explained by a small number of their principal components. This begs the question: Can our complex pricing model be compressed to achieve similar performance with potentially many fewer parameters? Evidently, the answer is no. We consider replacing the large number of factors in our complex model with a smaller number of their principal components and find that dimension reduction significantly impairs out-of-sample performance relative to the full complex model.

Lastly, we conduct a variety of robustness analyses that consider different input characteristics, various subsets of the stock universe, different training windows, and so forth. Every robustness test reinforces our main findings. In summary, we establish the empirical fact that large factor models are better at pricing assets than existing parsimonious models.

1.2 Theoretical Findings

In order to understand our empirical findings, we develop a statistical theory of large factor pricing models. Consider the following interpretation of a large factor model from the asset pricing theory perspective of an SDF. A conditional SDF can be represented as a tradable portfolio of risky assets:

$$M_{t+1} = 1 - w(Z_t)'R_{t+1}. \tag{1}$$

R_{t+1} contains excess returns of risky assets, and $w(Z_t)$ are the SDF's weights on those assets. The weights are determined by variables Z_t that span the time t information set. Without further guidance about the functional form of w , a machine learning-based empirical approach would approximate the SDF with a nonparametric model such as a shallow neural

network: $w(Z_t) \approx \sum_{p=1}^P \lambda_p S_p(Z_t)$, where each $S_p(Z_t)$ is some nonlinear basis function of Z_t . Thus, a neural network SDF delivers a linear factor pricing model,

$$M_{t+1} \approx 1 - \sum_p \lambda_p F_{p,t+1}, \quad (2)$$

where each “factor” $F_{p,t+1}$ is a characteristic-managed portfolio that uses the transformed “characteristics” $S_p(Z_t)$ as portfolio weights, and $\lambda = (\lambda_1, \dots, \lambda_P)$ are parameters to be estimated.

This formulation captures the tension in machine learning asset pricing models. On the one hand, the more terms in the nonparametric representation of w (and thus the larger the number of factors), the better the model can approximate the true SDF. But as we improve the approximation, we simultaneously increase the number of parameters that must be estimated, so when P becomes large, the estimated model may be unstable and suffer out-of-sample. Our theory makes it possible to understand (and quantitatively characterize) this cost/benefit tradeoff in large factor models. This is non-trivial and requires our novel random matrix theory (RMT) results that allow us to tackle the “three infinities” problem (when the number of assets, observations, and parameters are simultaneously large).³

On the cost side of the tradeoff is a phenomenon we refer to as “limits to learning.” In low complexity settings—with many more observations than parameters to estimate—the law of large numbers kicks in, and appropriate estimators can recover the true model.⁴ Traditional econometric theory deals with estimator behavior in this data-rich environment. By contrast, in high-complexity settings, the number of factors is large relative to the number of observations. Thus, the law of large numbers breaks down. Even correctly specified estimators fail to converge on the true model because there is not enough data to go around. The first key aspect of our theory is an explicit characterization of the limits to the learning

³Established RMT results (e.g., those from [Hastie et al., 2019](#)) are insufficient in our context.

⁴I.e., if the model is correctly specified. If the model is mis-specified, estimators recover the nearest “pseudo-true” parameters (e.g. [White, 1996](#)).

effect. We prove that data limitations impose an unavoidable implicit shrinkage of the model, which depresses the model’s expressive power. Complexity induces bias through its implicit shrinkage. But on the benefit side of the tradeoff, complexity also *reduces* specification bias. More factors lead to a more accurate approximation of the true SDF. Our main theoretical result shows that which effect dominates depends on the eigenvalue distribution of the factors. When there is a concentrated eigenvalue distribution and thus a few dominant factors, asset pricing model performance tops out after adding only a few factors, and there is no virtue of complexity—a counterfactual prediction in light of our empirical findings.

However, when the underlying factor structure is not too concentrated—that is, when the data-generating process aligns with the APT conjecture—a virtue of complexity emerges that closely mimics the patterns we find in the data. Our theory rationalizes our surprising empirical facts, demonstrating that even if arbitrage is absent and an SDF exists, it is possible to continually find new empirical “risk” factors that are unpriced by others and that adding these factors to the pricing model continually improves its out-of-sample performance. The theory also helps rationalize the prominence of “anomaly” portfolios in empirical asset pricing. The abundance of anomalies (the so-called “factor zoo”) is not a puzzle to be solved or evidence of a corrupt research process.⁵ Instead, it is the theoretically expected outcome in a complex asset pricing environment because data limitations hamper our ability to learn the true nature of markets. In fact, our theory argues that the extant factor zoo is *too small* and that an SDF model can be beneficially expanded to incorporate a teeming Noah’s ark of factors by transforming raw asset characteristics into a wide variety of nonlinear signals (buttressed by appropriate shrinkage). Such a large factor set improves the out-of-sample SDF Sharpe ratio and reduces out-of-sample pricing errors.

⁵Jensen et al. (2023) reach a similar conclusion based on the rationale that the risk-return trade-off is difficult to measure and complexity manifests as an inability to find a single silver-bullet characteristic that pins down expected returns. Instead, researchers gradually expand and refine the set of noisy signals and conclude “a more positive take on the factor zoo is not as a collective exercise in data mining and false discovery, but rather as a natural outcome of a decentralized effort in which researchers make contributions that are correlated with, but incrementally improve on, the body of knowledge.”

1.3 Literature Review

This paper is related to an emergent literature that shows machine learning asset pricing models achieve higher out-of-sample Sharpe ratios and smaller pricing errors than their parsimonious predecessors. Examples include [Kozak et al. \(2020\)](#), [Gu et al. \(2020a\)](#), [Chen et al. \(2023\)](#), [Bryzgalova et al. \(2020\)](#), [Cong et al. \(2022\)](#), [Fan et al. \(2022\)](#) and [Preite et al. \(2022\)](#), among others. It also relates to machine learning methods for analyzing factor models, including [Connor et al. \(2012\)](#), [Fan et al. \(2016\)](#), [Kelly et al. \(2020\)](#), [Lettau and Pelger \(2020\)](#), [Giglio and Xiu \(2021\)](#), [Giglio et al. \(2022\)](#), and [He et al. \(2023\)](#) (see [Kelly and Xiu \(2023\)](#) and [Rapach and Zhou \(2020\)](#) for more complete survey treatments of these topics). We extend this literature with detailed documentation of the counterintuitive phenomenon that out-of-sample asset pricing model performance appears to improve with ever richer model specifications.

An understanding of this surprising empirical phenomenon is only beginning to take shape. [KMZ](#) theoretically analyze machine learning models for return prediction. We extend this work in a number of ways, including reorienting the statistical objective around minimizing pricing errors and maximizing SDF Sharpe ratios, and by moving from a single asset to a cross-section of assets.⁶ Together with [Martin and Nagel \(2021\)](#), [Da et al. \(2022\)](#), [Fan et al. \(2022\)](#), and [KMZ](#), our paper belongs to an emergent literature analyzing “limits to learning” in high-dimensional asset pricing models. Our paper also relates to the theoretical

⁶From a technical standpoint, we overcome a number of new theoretical hurdles relative to [KMZ](#). In time series regressions of [KMZ](#), the random matrix behavior of time series signal covariances dictates the market timing strategies. In the panel problem, behavior is determined not just by time series covariances but also by the covariance of signals across assets. While standard random matrix theory in [KMZ](#) requires dealing with double limits (number of observations and number of parameters), our analysis requires the development of a non-trivial theoretical extension to deal with a third limiting dimension—the number of base assets. This contrasts with recent results of [Martin and Nagel \(2021\)](#), who lack the time dimension and assume that all stock characteristics are constant. Such an assumption implies that factors (characteristics-managed portfolios) are, in fact, static portfolios of stocks. In particular, the number of factors in their setting is always smaller than the number of stocks. Importantly, we also remove the equal ex-ante predictive power assumption of [KMZ](#) and allow for a generic distribution of risk premia across factors.

machine learning literature on topics surrounding “double descent” and “benign overfit” (e.g. [Spigler et al., 2019](#); [Hastie et al., 2019](#); [Belkin et al., 2020](#); [Bartlett et al., 2020](#)).

A related empirical literature analyzes the association between asset returns and future macroeconomic conditions. [Parker and Julliard \(2005\)](#) and [Croce et al. \(2023\)](#) show that stock returns covary with future consumption growth.⁷ [Bryzgalova et al. \(Forthcoming\)](#) and [Dew-Becker and Giglio \(2016\)](#) use returns to identify the role of long-lived consumption risks for pricing assets. [Backus et al. \(2007\)](#) show that equity returns lead GDP growth by a few quarters. We add to this literature in a number of ways. First, our complex SDF forecasts future not just consumption and output, but a wide range of macroeconomic activity and conditions. The macro predictability arising from our model is larger in magnitude than that of benchmark factor models in the literature. A particularly important new finding is our complex SDF’s strong predictability for macroeconomic uncertainty, giving credence to popular consumption-based asset pricing models with stochastic macroeconomic volatility (e.g. [Bansal and Yaron, 2004](#); [Campbell et al., 2018](#)). In contrast, none of the standard benchmarks demonstrate any ability predict macroeconomic uncertainty.⁸

Through its coupling with the literature on factor pricing models, our work also relates to the empirical literature surrounding the “factor zoo” and factor replicability, including [Harvey et al. \(2016\)](#), [McLean and Pontiff \(2016\)](#), [Hou et al. \(2020\)](#), [Feng et al. \(2020\)](#), [Jensen et al. \(2023\)](#), and [Chen and Zimmermann \(2021\)](#). Our theory helps rationalize the continued discovery of factors that are unspanned by simpler precedent models. Relatedly, the demonstrated success of machine learning models in predicting the cross-section of returns such as [Chinco et al. \(2019\)](#), [Han et al. \(2019\)](#), [Freyberger et al. \(2020\)](#), [Rapach and Zhou \(2020\)](#), [Gu et al. \(2020b\)](#), [Avramov et al. \(2023\)](#), and [Guijarro-Ordóñez et al. \(2021\)](#), is additional evidence of the virtue of complexity in financial markets research.

The remainder of the paper proceeds by introducing our empirical design in Section 2

⁷[Chernov et al. \(2023\)](#) demonstrate that currency returns are also predictive of future consumption.

⁸Relatedly, [Bryzgalova et al. \(Forthcoming\)](#) point out that consumption volatility shocks are conspicuously absent from their estimation of the risk-return tradeoff.

and our empirical findings in Section 3. Section 4 presents our theoretical analysis of large factor models, and Section 5 concludes.

2 Empirical Framework

We now present the details of our empirical framework for large factor pricing models and their associated SDF representations.

2.1 Background: Factor Models and the SDF

[Hansen and Richard \(1987\)](#) show that a true SDF, if one exists, is representable as a tradable portfolio of risky assets as introduced in Equation (1).

A factor pricing model is equivalent to an SDF that is linear in the factors. We can, therefore, analyze and compare asset pricing factor models as competing SDF specifications. Motivated by the APT and the principle of parsimony,⁹ the literature has primarily investigated (1) with tightly constrained factor specifications of the SDF weight function $w(Z_t)$. A leading example is the [Fama and French \(1993\)](#) model, which restricts M_{t+1} to be a three-parameter model:

$$w_i^{\text{FF}} = c_1 \text{MKT}_i + c_2 \text{SIZE}_i + c_3 \text{VALUE}_i,$$

where the weight of asset i in the Fama-French SDF (w_i^{FF}) depends only on the stock's weight in the market portfolio (MKT_i) and on researcher-dictated functions of the stock's size and book-to-market ratio (SIZE_i and VALUE_i).

The Fama-French SDF portfolio can equivalently (and more familiarly) be viewed as a

⁹See, e.g., [Tukey \(1961\)](#) and [Box and Jenkins \(1970\)](#). Economic theory also motivates functional restrictions on the SDF (e.g. [Hansen and Singleton, 1982](#)). However, restrictions derived from economic theory have had limited success to date in pricing cross-sections of assets such as stocks, bonds, and derivatives. The Fama-French SDF and many other factor models in the literature are motivated by empirical “anomalies” vis-a-vis the CAPM and not by a particular economic theory.

linear combination of three factors. The factors are portfolios whose individual stock weights are given by MKT_i , SIZE_i , and VALUE_i , respectively. The SDF combines these factors with three parameters, c_1 , c_2 , and c_3 , corresponding to the Sharpe ratio maximizing combination of the factors. These are the only three parameters that must be estimated in the Fama-French SDF, resulting in an extraordinarily parsimonious asset pricing model. Thus, the Fama-French model is a special case of (1) where

$$w(Z_t^{FF}) = Z_t^{FF} \lambda^{FF}.$$

The conditioning set is Z_t^{FF} , which has $D = 3$ columns containing stocks' Fama-French characteristics, and the weight function is linear in parameters $\lambda^{FF} = (c_1, c_2, c_3)'$. In this case, the SDF is

$$M_t^{FF} = 1 - \lambda^{FF'} Z_t^{FF'} R_{t+1},$$

where we recognize that $Z_t^{FF'} R_{t+1}$ is a vector of time t returns on the three Fama-French factors and $1 - M_t^{FF}$ is their mean-variance efficient combination.

2.2 Designing A Large Factor Model

The Fama-French model achieves its parsimony from rigid assumptions on the SDF weighting function—discarding all but a few conditioning variables and making detailed choices for the functional form that maps the selected variables into SDF weights.

The AIPT takes a different approach, founded on the philosophy of machine learning, and explores the benefits of flexible models that accommodate many potential conditioning variables and diverse functional forms. In our analysis, Z_t is an $N_t \times D$ matrix that collects a large number of D conditioning variables for each risky asset. We derive a complex factor pricing model by replacing detailed specification choices with a generic nonparametric

approximation to the unknown SDF:

$$w(Z_t) \approx \sum_{p=1}^P \lambda_p S_p(Z_t) = S_t \lambda, \quad (3)$$

where each $S_p(Z_t)$ is an $N_t \times 1$ nonlinear basis function of the input data Z_t and λ_p is a scalar basis coefficient. This basis expansion explores the space of potential shapes for $w(Z_t)$, then combines these building blocks with coefficients (λ_p) that best reconstruct the unknown function $w(Z_t)$. Collecting the P basis terms into an $N_t \times P$ matrix of characteristics $S_t = [S_{1,t}, \dots, S_{P,t}]$, and collecting parameters $\lambda = [\lambda_1, \dots, \lambda_P]'$, we arrive at $S_t \lambda$ as the nonparametric approximation to the true SDF weighting function. The larger is P , the broader the set of nonlinearities considered, the larger the number of model parameters, and the more flexible the approximation.

The nonparametric approximation in (3) takes the same basic form as the Fama-French SDF with S_t in place of Z_t^{FF} . The basis terms $S_p(Z_t)$ merely transform conditioning variables into a large number of new nonlinear asset “characteristics” analogous to those underpinning the Fama-French model. Substituting the approximation in (3) into the SDF definition (1), we arrive at

$$M_{t+1} \approx 1 - \lambda' S_t' R_{t+1} = 1 - \lambda' F_{t+1} = 1 - R_{t+1}^M. \quad (4)$$

In other words, like the Fama-French model, the SDF portfolio (denoted $R_{t+1}^M = \lambda' F_{t+1}$) is a linear combination of factor portfolios,

$$F_{t+1} = S_t' R_{t+1}. \quad (5)$$

Each factor $F_{p,t+1} = S_p(Z_t)' R_{t+1}$ is a characteristic-managed portfolio of risky assets that uses the nonlinear asset “characteristic” $S_p(Z_t)$ as the vector of portfolio weights, and again λ

is the mean-variance efficient combination of factors.¹⁰ The difference vis-a-vis Fama-French (and other parsimonious models) is that the factors in (4) make flexible and nonlinear use of conditioning data, and there are a large number P of these factors.

2.3 Random Fourier Factors

To operationalize this large factor model approximation of the SDF, our empirical approach adapts the method of random Fourier features, or RFF (Rahimi and Recht, 2007).¹¹ While (Rahimi and Recht, 2007) propose RFF for predictive regression, our approach embeds RFF in an SDF. The nonlinear basis functions in this method are trigonometric transformations of the original signals Z_t :

$$[S_{2p-1,t}, S_{2p,t}] \in \mathbb{R}^{N_t \times 2} = [\sin(\gamma Z_t \omega_p), \cos(\gamma Z_t \omega_p)]', \quad \omega_p \sim \text{i.i.d. } N(0, I), \quad p = 1, \dots, P/2. \quad (6)$$

Each RFF basis function forms a random linear combination (ω_p) of the raw characteristics Z_t , then feeds this through sine and cosine activation functions.¹² The transformation is applied stock-by-stock, converting stock i 's characteristics into new nonlinear characteristics that include not just transformations of each individual characteristic but general multi-way interactive effects involving all characteristics for stock i .

Upon closer inspection, we see that the SDF in (4) is, in fact, a two-layer neural network with input parameters given by ω_p , output parameters given by λ , and trigonometric activation neurons (see Figure 1). In a typical neural network, both the input and output parameters are estimated using computationally costly numerical methods. The fascinating insight of (Rahimi and Recht, 2007) is that, by using randomization to fix the input

¹⁰Because the size of the cross-section varies over time, our empirical analysis in fact scales the factors as $F_{t+1} = N_t^{-1/2} S_t' R_{t+1}$ to maintain a relatively consistent scale of factor returns over time.

¹¹In Section 3.8 we demonstrate the robustness of our findings to other complex model specifications.

¹²The parameter γ controls the Gaussian kernel bandwidth in generating random Fourier features. Following Kelly et al. (2022), we randomly chose γ from the grid $[0.5, 0.6, 0.7, 0.8, 0.9, 1.0]$ for each ω_i that we generate. This embeds varying degrees of nonlinearity in the generated feature set S_t .

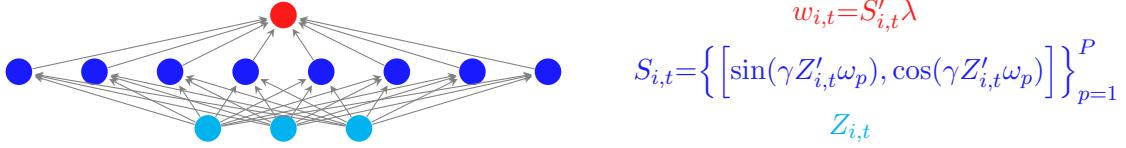


Figure 1: Neural Network Representation of Random Fourier Factors

Note. Illustration of nonlinear mapping from stock i 's characteristics $Z_{i,t}$ into hidden layer neurons $S_{i,t}$ and then into the conditional SDF weight $w_{i,t}$ during the construction of random Fourier factors.

parameters and only estimating the output parameters, random features regression can build nonparametric approximations in a computationally efficient manner.¹³ Building on this insight, our SDF model is estimable in closed form using only least squares regression (as discussed below).

RFF is an ideal tool for our analysis because, from a fixed input data set Z_t we can create asset pricing models with any desired factor dimension, P . For a low-dimensional model of, say, $P = 2$, one generates a single pair of RFFs. For a high-dimensional model of, say, $P = 10,000$, one can instead draw many random weight vectors ω_p , $p = 1, \dots, 5,000$. This gives us the ability to evaluate how asset pricing models are impacted by varying the number of factors *while holding the set of conditioning variables fixed*. It is important to emphasize that in our analysis, factor proliferation arises not from expanding the set of input data Z_t (which may or may not be low-dimensional) but instead from generating many nonlinear basis transformations of Z_t .

2.4 Estimator

The SDF coefficient vector λ represents the mean-variance efficient portfolio of factors. Note that the factor representation of the *conditional* SDF in essence produces an equivalent *unconditional* SDF. This allows us to estimate λ as the sample Markowitz portfolio of factors:

¹³Indeed, RFF is a special case of kernel regression. For applications of kernel regression in asset pricing settings, see Kozak (2020) and Kelly et al. (2024a).

$$\arg \max_{\lambda} \left\{ \hat{E}[\lambda' F_t] - \frac{1}{2} \hat{E}[(\lambda' F_t)^2] \right\} = \hat{E}[F_t F_t']^{-1} \hat{E}[F_t] \quad (7)$$

where $\hat{E}$ denotes the training sample mean. We denote the corresponding estimated SDF and SDF portfolio as

$$\hat{M}_t = 1 - \hat{R}_t^M \quad \text{and} \quad \hat{R}_t^M = \hat{\lambda}' F_t. \quad (8)$$

The AIPT explores large factor models in which the dimension of λ may exceed the number of time series observations, i.e., $P > T$. In this case, $\hat{E}[F_t F_t']^{-1}$ is rank-deficient, which we resolve by introducing a ridge penalty in the Markowitz problem:

$$\hat{\lambda}(z) = \arg \max_{\lambda} \left\{ \hat{E}[\lambda' F_t] - \frac{1}{2} \hat{E}[(\lambda' F_t)^2] - z \lambda' \lambda \right\} = \left(zI + \hat{E}[F_t F_t'] \right)^{-1} \hat{E}[F_t] \quad (9)$$

where z is the ridge parameter.

[Kozak et al. \(2020\)](#) point out that we can also interpret $\hat{\lambda}(z)$ in terms of “pricing errors.” Since the factors F_t are tradable assets, the true SDF M_t prices these factors with zero error due to the marginal investor’s first-order optimality condition:

$$E[M_t F_t] = 0. \quad (10)$$

[Hansen and Jagannathan \(1997\)](#) suggest a statistic for comparing the efficacy of SDF models in terms of their pricing error magnitudes. For a candidate SDF $\tilde{M}_t$, The Hansen-Jagannathan distance (HJD) is a weighted sum of squared pricing errors for the set of test

assets F_t :¹⁴

$$\mathcal{D}^{HJ} = E[\tilde{M}_t F_t]' E[F_t F_t']^{-1} E[\tilde{M}_t F_t]. \quad (11)$$

Substituting the SDF approximation model (4) for $\tilde{M}$ in the HJD and adding a ridge penalty to handle rank deficiency when $P > T$, we have

$$\hat{\lambda}(z) = \arg \min_{\lambda} \left\{ \hat{E}[(1 - \lambda' F_t) F_t]' \hat{E}[F_t F_t']^{-1} \hat{E}[(1 - \lambda' F_t) F_t] + z \lambda' \lambda \right\}. \quad (12)$$

In other words, $\hat{\lambda}(z)$ is also the regularized SDF estimator that minimizes pricing error (we discuss HJD in further detail in Section 2.6).

2.5 Data and Training

To make the conclusions from this analysis as easy to digest as possible, we perform our analysis in a conventional setting with conventional data. We use a comprehensive and standardized sample of monthly US stock returns and 153 stock characteristics from 1963 to 2023, compiled by [Jensen et al. \(2023\)](#) ([JKP](#) henceforth).¹⁵ As in [JKP](#), our universe includes NYSE/AMEX/NASDAQ securities with CRSP share code 10, 11, or 12, excluding “nano” stocks (i.e., stocks with market capitalization below the first percentile of NYSE stocks).

Some of the [JKP](#) characteristics have low coverage, especially in the early parts of the sample. To ensure that characteristic composition is fairly homogeneous over time and to avoid purging a large number of stock-month observations due to missing data, we reduce the 153 characteristics to a smaller set of $D = 130$ characteristics with the fewest missing values. We drop stock-month observations for which more than 30% of the 130 characteristic values are missing and use N_t to denote the number of the remaining stock observations at time t . Next, we cross-sectionally rank-standardize each characteristic and map it to the

¹⁴We discuss some attractive properties of the HJD in the following sections.

¹⁵To access the stock-level characteristic data and associated documentation, refer to jkpfactors.com.

$[-0.5, 0.5]$ interval, following [Gu et al. \(2020b\)](#). Ultimately, we obtain an $N_t \times D$ matrix of conditioning characteristics Z_t in each month.

In our empirical analyses, we study SDF performance as we vary two aspects of the model: the number of factors P and ridge penalty z . For each model of a given size and shrinkage, we conduct a rolling out-of-sample model performance analysis. Starting in January 1993, in each month t , we use the most recent 360 months of data to estimate the ridge SDF parameter in (9).¹⁶ We then track the out-of-sample SDF portfolio return in the subsequent month. From the sequence of monthly out-of-sample SDF realizations, we calculate out-of-sample SDF performance metrics.

2.6 Performance Metrics

We evaluate our SDF models with two standard out-of-sample performance metrics. A true SDF is the mean-variance efficient portfolio of risky assets. Thus, our first performance metric is the SDF portfolio Sharpe ratio. For a candidate SDF $\tilde{M}_t = 1 - \tilde{R}_t^M$, the annualized out-of-sample Sharpe ratio is

$$\widehat{SR} = \sqrt{12} \frac{\hat{E}_{OS}[\tilde{R}_t^M]}{\hat{\sigma}_{OS}(\tilde{R}_t^M)}$$

where $\hat{E}_{OS}$ and $\hat{\sigma}_{OS}$ denote sample average and standard deviation among the T_{OS} out-of-sample observations (to differentiate versus the notation $\hat{E}$ and T used for training samples). Since the SDF is constructed from excess returns on risky assets, the numerator need not be adjusted for the risk-free rate.

Our second performance metric is the pricing error for a large set of test assets. The test assets in our main analysis are the 153 [JKP](#) anomaly factors constructed by those authors

¹⁶The stochastic nature of RFF means that there is inherent variability in the estimated SDF model when P is small. To mitigate this variability, we repeat the RFF-based estimation 20 times with different random seeds and report average performance metrics across seeds. This has little qualitative effect on the plots shown. See [KMZ](#) for further discussion of this point.

and downloadable at jkpfactors.com (robustness analyses report pricing errors for other test asset sets as well). To aggregate test asset pricing errors into a single metric, we calculate the normalized sum of squared pricing errors according to the Hansen-Jagannathan distance (HJD) referenced in Section 2.4. For a set of test asset returns R_t^T , the out-of-sample HJD is

$$\hat{\mathcal{D}}^{HJ} = \hat{E}_{OS} \left[\tilde{M}_t R_t^T \right]' \hat{E}_{OS} \left[R_t^T R_t^{T'} \right]^+ \hat{E}_{OS} \left[\tilde{M}_t R_t^T \right], \quad (13)$$

where $^+$ is the Moore-Penrose quasi-inverse (necessary due to the degeneracy of $\hat{E}_{OS}[R_t^T R_t^{T'}]$ when $P > T_{OS}$).

The HJD has a number of attractive model comparison properties.¹⁷ It averages pricing errors among test assets using a common weighting matrix for all candidate models of M_t . This is important because it puts all models on equal footing for comparison, unlike other alpha or GMM-based comparisons. The weighting matrix is economically motivated since it gives the HJD a clear interpretation as the pricing error of the portfolio of test assets that is *most mispriced* by each model. Furthermore, the HJD is scale invariant.¹⁸ While typically used for in-sample comparison, the HJD easily generalizes for out-of-sample evaluation because it avoids the need to estimate out-of-sample time series alphas and betas for each test asset. Finally, because our theoretical derivations explicitly characterize the expected out-of-sample HJD for complex SDF models, the empirical HJD can be directly compared to theoretical predictions.

¹⁷In addition to Hansen and Jagannathan (1997), the literature including Kan and Robotti (2009), Chen and Ludvigson (2009), and Kelly and Xiu (2023) further advocates HJD as a pricing error metric.

¹⁸Differences in volatilities among factors can skew test statistics like average absolute alpha or the GRS statistic, and as a result, the volatility scaling choice for test assets can lead to different conclusions. HJD is invariant to changes in test asset volatility.

2.7 Benchmark Models

We provide reference points for our analysis by comparing large factor model performance to five benchmark factor pricing models in the literature:¹⁹

FF6: This is a six-factor model that includes the five factors of [Fama and French \(2015\)](#) plus their UMD momentum.

SY: This is a four-factor model that includes the “mispricing” factors of [Stambaugh and Yuan \(2017\)](#).

HXZ: This is a five-factor model that includes the q -factor model specification of [Hou et al. \(2015\)](#) augmented to include the expected growth factor of [Hou et al. \(2021\)](#).

DHS: This is a three-factor model that includes the long-horizon and short-horizon behavioral factors of [Daniel et al. \(2020\)](#).

BS: This is a composite model that selects factors from the literature based on the Bayesian methodology of [Barillas and Shanken \(2018\)](#). Its six factors include the market factor, the investment and profitability factors from [Hou et al. \(2015\)](#), the size and momentum factors from [Fama and French \(2015\)](#), and the value factor (HML^m) from [Asness et al. \(2013\)](#).

To compare models, we use the SDF corresponding to each of these factor pricing models, which requires estimating the Markowitz portfolio of factors in each model. We do this in rolling 360-month training windows and track benchmark performance out-of-sample. In other words, we use the same training/evaluation design for the benchmarks that we use for our large factor pricing models.²⁰

¹⁹All benchmark factor data are conveniently available from the respective authors’ websites.

²⁰When building the Markowitz portfolio for each benchmark model, we do not use ridge shrinkage because the number of assets is far smaller than the number of time series observations. Even with infeasible (ex-post optimal) ridge shrinkage, we find the out-of-sample SDF performance of benchmark models changes negligibly.

3 Empirical Results

3.1 The Virtue of Complexity in Factor Pricing Models

Our central empirical finding is that factor model performance is increasing in the number of factors. We illustrate this virtue of complexity by plotting out-of-sample model performance as a function of the number of model parameters (which, in our setting, corresponds to the number of factors). [KMZ](#) refer to such plots as “VoC curves.” Each curve applies a different ridge penalty $z = 10^{-5}, 10^{-3}, 10^{-1}, 1$, or 10 . Along each curve, we vary the complexity—i.e., the number of factors—in our SDF model. We consider specifications with $P = 36, \dots, 360,000$ random Fourier factors, corresponding to $c = 0.1, \dots, 1,000$. The VoC curves in [Figure 2](#) show four out-of-sample properties of each SDF specification: expected return (Panel A), volatility (Panel B), Sharpe ratio (Panel C), and pricing error (Panel D).

We find the following main patterns associated with model complexity. As we increase the number of factors, i) the average SDF return rises, ii) SDF volatility rises, iii) the SDF return increases faster than volatility, resulting in a rising SDF Sharpe ratio, iv) pricing errors decline, and v) performance gains accrue more quickly with less ridge shrinkage. The exceptions to these patterns occur when ridge shrinkage is small, and complexity is close to one, in which case we see a spike in out-of-sample SDF volatility that induces an abrupt dip in the Sharpe ratio and a spike in pricing error. This “double ascent” in Sharpe ratio and “double descent” in pricing error is reminiscent of previously documented phenomena in the machine learning literature (e.g. [Hastie et al., 2019](#); [Bartlett et al., 2020](#)).

Quantitatively, the differences in out-of-sample performance for high and low-complexity models are dramatic despite the fact that both model types use identical data inputs and differ only in the richness of their parameterizations. High-complexity models with light shrinkage earn out-of-sample Sharpe ratios near 3.7, roughly 2.6 times larger than the Sharpe ratios of low-complexity models with $c < 1$. Likewise, large factor models are much better

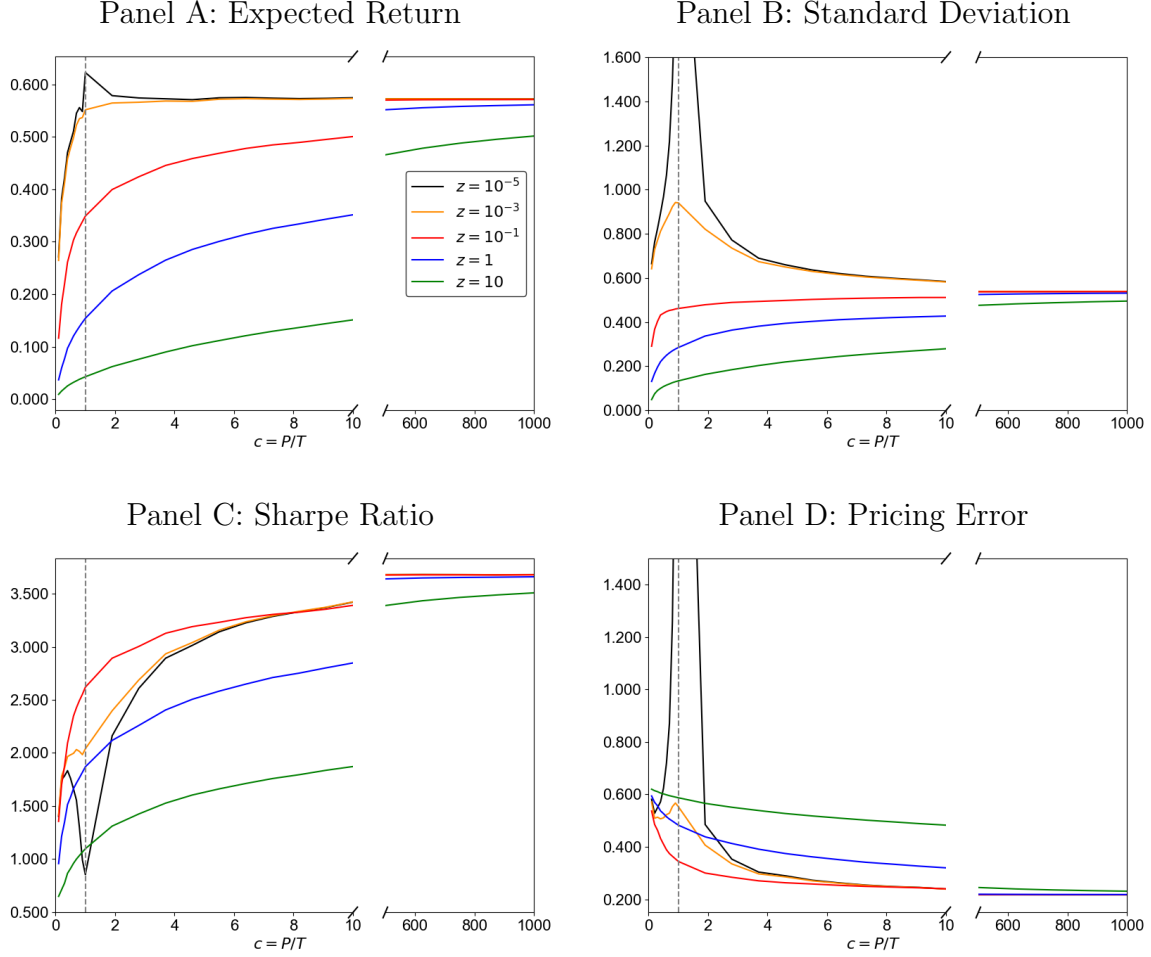


Figure 2: Out-of-sample Performance of Complex Factor Models

Note. Realized out-of-sample SDF portfolio average return, standard deviation, Sharpe ratio, and pricing error (HJD). The horizontal axis shows model complexity $c = P/T$, with P ranging from 36 to 360,000 and $T = 360$ months. Factors underlying the SDF are characteristic-managed portfolios constructed with random Fourier features derived from [JKP](#) stock characteristics. The test assets in Panel D are 153 anomaly factors from [JKP](#).

positioned to price our demanding test assets (153 anomaly portfolios collected from decades of finance research). Pricing errors for low-complexity models are more than twice as large as those for high-complexity models (0.54 vs 0.22).

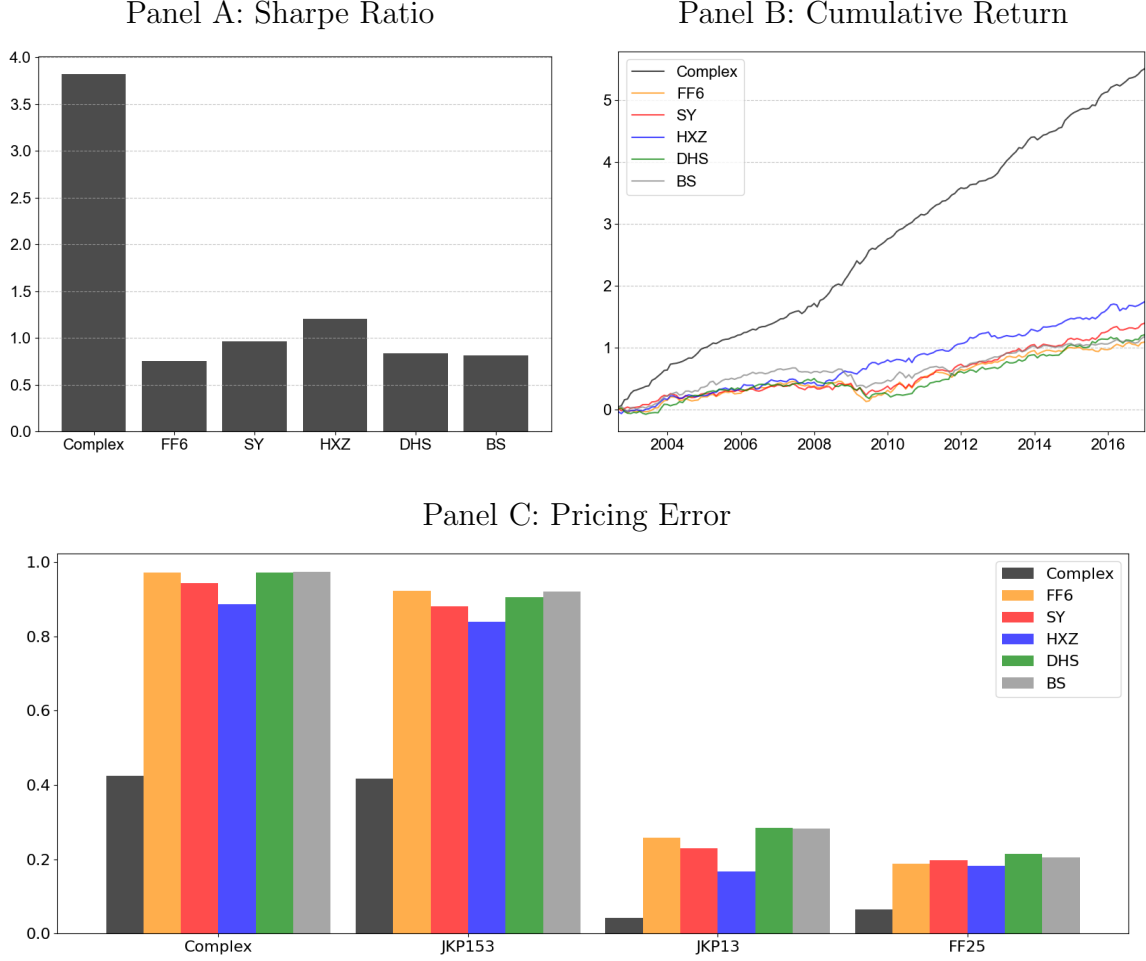


Figure 3: Performance Comparison of Complex and Benchmark Factor Models

Note. This figure reports the out-of-sample performance of SDF portfolios corresponding to various factor pricing models, including the complex model ($P = 360,000$ and $z = 10^{-5}$) as well as the simple benchmark models FF6, HXZ, SY, DHS, and BS. Panel A shows the out-of-sample Sharpe ratio for each model’s factor Markowitz portfolio. Panel B shows the simple cumulative sum of out-of-sample returns for each model’s factor Markowitz portfolio (we have standardized out-of-sample returns to 10% annualized volatility for all models in order to facilitate comparison in this panel). Panel C shows pricing errors (HJD) for four sets of test assets, including are the 360,000 random Fourier factors underlying the large factor model (denoted “complex”), the 153 **JKP** factors, the 13 **JKP** aggregate theme factors, and the 25 Fama-French size and book-to-market portfolios. Benchmark data sets are available for different sample periods, so we report statistics on the intersection of their out-of-sample ranges, which is from August 2002 to December 2016.

3.2 Comparison With Benchmarks

In Figure 2, we compare large and small variants of factor models using random Fourier factors. However, the conclusion that large factor models outperform simple models is not specific to the nonparametric framework in Figure 2.

Figure 3 compares the high complexity factor model ($P = 360,000$ and $z = 10^{-5}$) to five of the leading factor models in the literature: FF6, HXZ, SY, DHS, and BS. Panel A reports the out-of-sample Sharpe ratio for each model’s SDF portfolio, Panel B shows the corresponding cumulative return plot, and Panel C reports out-of-sample pricing errors.²¹

Panel A shows that the benchmark factor models are all roughly on par with each other in terms of Sharpe ratio, while the complex model rather dramatically outperforms every simple model. Its Sharpe ratio is around 3-6 times higher than the benchmarks. The cumulative returns in Panel B show that the portfolio performance comparison is not driven by any particular subsample and that the complex model does not exhibit any periods of unusually turbulent behavior.

Panel C shows that the large factor model also dominates in terms of pricing error. When the 153 JKP factors are test assets, HJD for the benchmark models hovers around 0.9, roughly twice that of the large factor model. This finding is not specific to this set of test assets. When the test assets are the 360,000 random Fourier factors themselves, the HJD hovers around 0.95 for benchmark models, again roughly twice the HJD of the complex factor model. Third, we study the JKP theme factors, which collect the individual anomaly factors into 13 aggregate theme portfolios. The benchmark HJD for this test set is around 0.2, which is roughly four times larger than the complex model HJD. Finally, we use the classic Fama-French 25 size and value portfolios. This test set has been the target for pricing models for forty years, and one can argue that benchmark models have been honed to price

²¹Figure 3 shows small differences in performance metrics for the high complexity factor model compared to Figure 2. This is because Figure 3 is restricted to a shorter time sample in which data is available for all benchmark models.

Table 1: Performance Comparison of Complex and Benchmark Factor Models

This table reports a comparison of out-of-sample SDF portfolio performance for various factor pricing models, including the complex model ($P = 360,000$ and $z = 10^{-5}$) as well as the simple benchmark models FF6, HXZ, SY, DHS, and BS. The first two rows show the annualized percentage alpha and associated t statistic of the complex SDF portfolios versus each benchmark SDF portfolio. The next two rows show each benchmark SDF portfolio’s alpha and t statistic versus the complex model. The last two rows show the ex-post tangency portfolio weights combining the SDF portfolios of all models with and without a non-negativity constraint (and constraining weights to sum to one). To facilitate alpha comparisons, out-of-sample SDF portfolio returns for all models are rescaled to have 10% annualized volatility. Benchmark data sets are available for different sample periods, so we report statistics on the intersection of their out-of-sample ranges, which is from August 2002 to December 2016.

	LHS					
	FF6	SY	HXZ	DHS	BS	Complex
Complex Alpha	42.8	40.7	39.4	43.5	44.4	–
vs. Benchmark	(13.9)	(13.6)	(13.0)	(13.7)	(13.8)	–
Benchmark Alpha	-2.8	-4.3	-2.3	1.50	3.90	–
vs. Complex	(-0.7)	(-1.2)	(-0.6)	(0.40)	(1.00)	–
Tangency Weights						
Unconstrained	-0.44	-0.24	-0.01	0.12	0.56	1.01
Non-negative	0.00	0.00	0.00	0.01	0.09	0.90

test assets like these. Even in this case, benchmark models deliver an out-of-sample pricing error of around 0.2, roughly three times larger than those from the complex pricing model.

Table 1 extends the benchmark comparison by reporting pairwise alphas between SDF portfolios of the large factor pricing model and each of the benchmark models. For this, we use the out-of-sample SDF portfolios but renormalize them ex post to a 10% annualized volatility to make alphas comparable for all models (unlike HJD, alphas lack scale invariance, so volatility scaling is helpful for comparison). The complex model produces a large and highly significant alpha versus all benchmarks, ranging from 39.4% to 44.4% per annum. The reverse is not true. Benchmark models all have small and statistically insignificant alphas versus the complex model. Finally, we estimate the “meta”-Markowitz portfolio that combines SDF portfolios of all individual models. The complex model dominates the meta

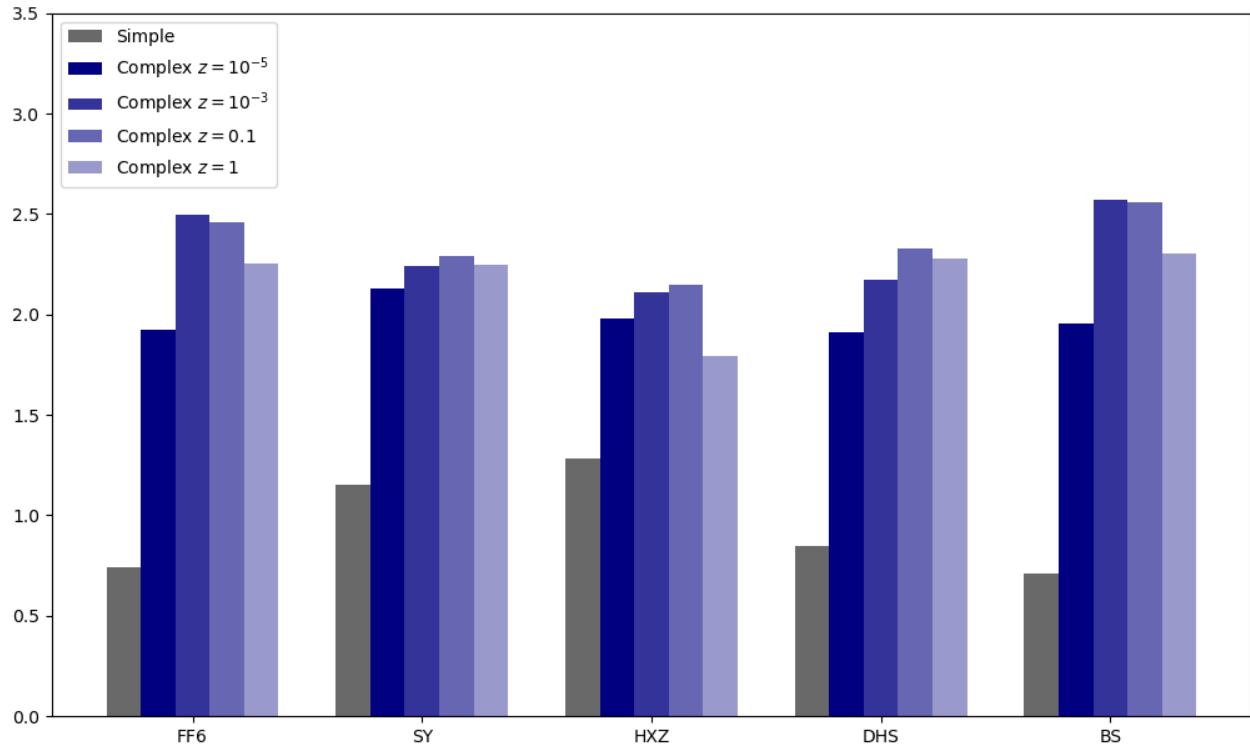


Figure 4: Complex Versions of Benchmark Models

Note. The figure reports out-of-sample Sharpe ratios of simple benchmark models compared to their nonlinear, high-complexity counterparts. For each benchmark model we construct its high complexity SDF with random Fourier factors derived from the characteristics underlying that model: size, value, investment, profitability, and momentum for FF6; size and two mispricing measures (MGMT and PERF) for SY; size, investment, profitability and expected growth for HXZ; size, PEAD, and two measures of issuance (CSI and NSI) for DHS; and asset growth, profitability, size, value and momentum for BS. Complex versions of each benchmark use $P = 360,000$ factors ($c = 1000$) and shrinkage of $z = 10^{-5}, 10^{-3}, 10^{-1}$ or 1.

portfolio, receiving a weight of 101% in the unconstrained portfolio and a weight of 90% in the portfolio that constrains weights on individual models to be non-negative.

3.3 The Nonlinear Fama-French Model

In the preceding analysis, we compare high complexity factor models derived from 130 [JKP](#) characteristics to the FF6 model and other benchmarks that are low complexity models derived from a small set of characteristics. Naturally, one may wonder about the benefits of complexity when we restrict the raw data inputs to match those used by a given benchmark

model. To investigate this, we construct complex SDFs using random features derived from *only* the characteristics that underly each benchmark. For example, in the case of FF6, we construct a complex factor model with a large number of nonlinear factors derived from the size, value, investment, profitability, and momentum characteristics.

Figure 4 reports the out-of-sample Sharpe ratio of high complexity variations of each benchmark model using $P = 360,000$ ($c = 1,000$) random Fourier factors and ridge shrinkage of $z = 10^{-5}, 10^{-3}, 10^{-1}$ or 1. The original formulation of the FF6 model has an out-of-sample Sharpe ratio of 0.74, while its high complexity version ranges from 1.93 to 2.50 depending on the degree of shrinkage. All benchmarks show the same pattern—if we hold the stock characteristics of each simple model fixed but enrich the model specification to incorporate highly parameterized nonlinear transformations of those characteristics, the out-of-sample performance of the benchmarks improves by a factor of 1.4 to 3.6, depending on the benchmark and the amount of shrinkage.

The conclusion is that the benefits of large factor models accrue even when starting from a small conditioning information set. Complexity may be applied to any set of conditioning variables to more flexibly reflect patterns in the underlying return-generating process.

3.4 Macroeconomics of the Complex SDF

What is responsible for the high degree of mean-variance efficiency of the complex SDF portfolio relative to benchmarks? At a statistical level, the answer is straightforward: The complex SDF model more successfully optimizes the maximum Sharpe ratio objective, indicating that it estimates the “true” efficient portfolio with greater accuracy than the benchmarks.

At an economic level, this improved measurement can help us better ascertain the economic nature of the “true” SDF. In equilibrium asset pricing models, an SDF reflects investors’ intertemporal marginal rate of substitution (IMRS), which is shaped by investors’

information and expectations about future macroeconomic conditions. As shown by [Hansen and Jagannathan \(1991\)](#), the mean-variance efficient portfolio is a projection of the IMRS onto the space of traded assets, resulting in a so-called “tradable” SDF like those studied above. If, as econometricians, we achieve a more efficient tradable SDF, then we are more successfully honing in on information and expectations about future economic conditions that investors are collectively impounding into current prices and that govern investors’ IMRS. As articulated by [Fama \(1990\)](#) in his study of market returns and the macroeconomy,

“the general hypothesis ... is that information about the production of a given period is spread across preceding periods and so affects the stock returns of preceding periods.”

We extend [Fama \(1990\)](#)’s hypothesis to the modern setting of SDF returns. If a complex SDF achieves superior mean-variance efficiency relative to benchmarks, it should better reflect investors’ IMRS, and it should in turn be a stronger predictor of future economic conditions. Or, stated more simply, portfolio efficiency and informational efficiency go hand-in-hand.

To pursue this hypothesis, we compare candidate SDFs in their ability to predict macroeconomic activity. Let y_t be a vector containing various measures of macro activity and conditions in log levels at the monthly frequency. In our baseline specification, y includes 10 variables: industrial production, macroeconomic uncertainty, the Federal Funds rate, employment, consumer sentiment from the Michigan survey, CPI, housing starts, the Baa-Aaa yield spread, oil prices, and the trade-weighted dollar exchange rate index.²²

The goal of our prediction analysis is to understand the extent to which information realized in SDF portfolio returns reflects expectations of future macroeconomic conditions. To accomplish this, we estimate impulse response functions using the method of local

²²Nine of these variables are from the FRED-MD database and correspond to FRED identifiers IND-PRO, FEDFUNDS, PAYEMS, UMCSENTx, CPIAUCSL, HOUST, AAA and BAA, OILPRICE_x, TWEX-AFEGSMTH_x. The only variable not from FRED-MD is the macroeconomic uncertainty index ($h = 1$ variant) from [Jurado et al. \(2015\)](#).

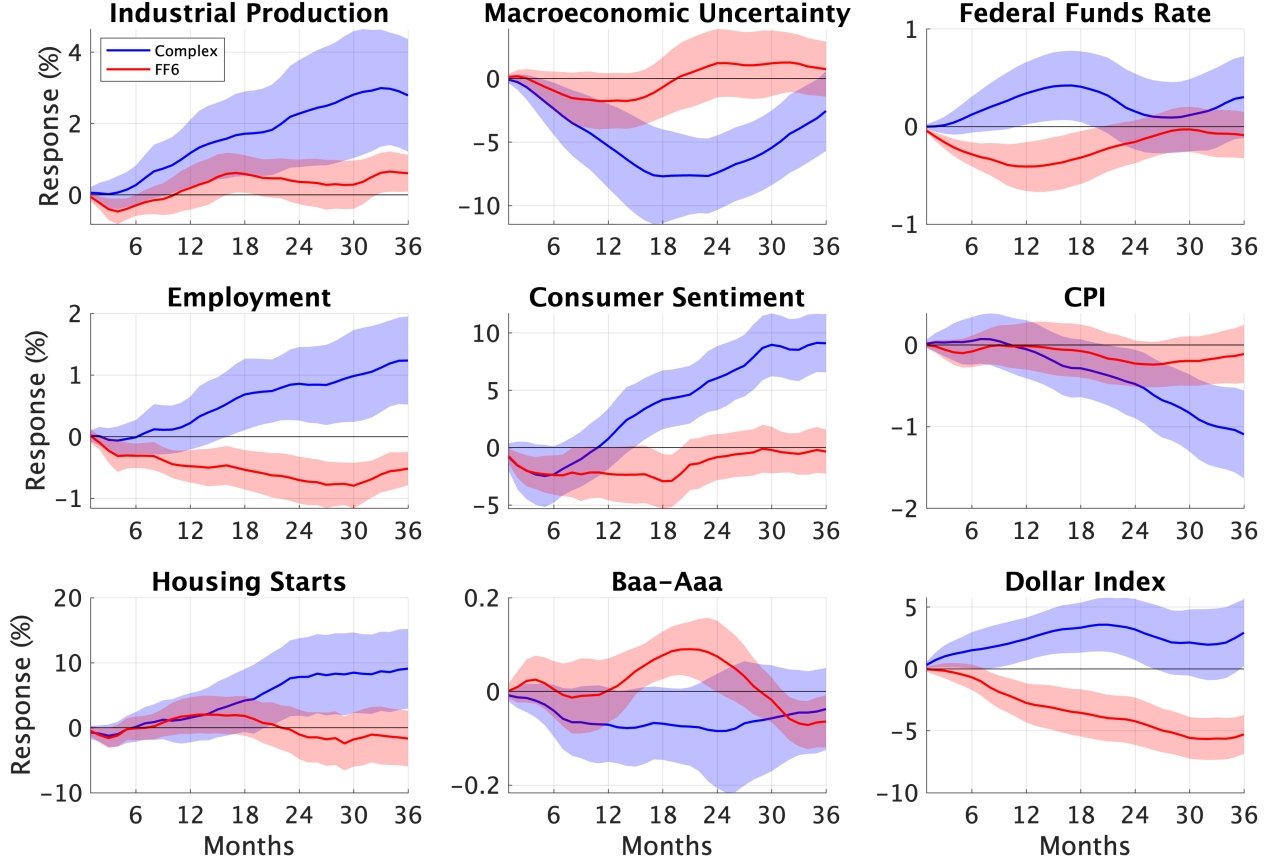


Figure 5: Macroeconomic Impulse Response Functions

Note. This figure reports SDF impulse responses estimated from equation (14). Dependent variables are shown in each figure title, and the impulse is a positive one standard deviation return to each SDF (the complex SDF in blue and the FF6 SDF in red). Shaded regions are pointwise 90% confidence intervals based on Newey-West standard errors with five lags.

projections (Jordà, 2005) using the “long difference” form (Jordà and Taylor, 2025):

$$y_{t+h} - y_t = a_h + b_{SDF,h} \left(\sum_{j=0}^{11} R_{t-j}^M \right) / \sigma^M + \sum_{j=0}^2 b_{j,h} y_{t-j} + e_{t+h}, \quad h = 1, \dots, 36. \quad (14)$$

The regression coefficient $b_{SDF,h}$ directly estimates, for each element of y , the expected h -period-ahead cumulative response (in percentage changes) to a positive one standard deviation realization of the annual SDF return, while controlling for predictive information

in y 's own lags.²³ We use a 12-month SDF return to accumulate information over a longer interval and reduce noise in the regressor.²⁴ We estimate regression (14) over the full out-of-sample period 1993-2023 using realized returns for the complex SDF and each benchmark.

Figure 5 reports the estimated impulse responses for the complex SDF. To provide a frame of reference Figure 5 also reports results for the FF6 SDF, which is representative of results for the other benchmarks. (Results for the full set of benchmarks are shown in Appendix L.)

The key conclusion from this analysis is that the complex SDF has strong predictive power for macroeconomic conditions, particularly at long horizons of two years or more. A positive one-standard-deviation annual return to the complex SDF predicts a roughly 3% (and highly significant) increase in industrial production over three years. It predicts a reduction in macroeconomic uncertainty of about 7.5% over two years. It predicts large and highly significant increases in employment (1.25%), consumer sentiment (9%), housing (9%), and the dollar (2.5%) over three years. A positive complex SDF return predicts a reduction in credit spreads of 0.08% over two years, which is a roughly 8% decrease in the credit spread relative to its long-term average level of 0.95%. Finally, we see a positive interest rate response of roughly 0.4% over 1.5 years (versus an average interest rate level of 2.5% in our sample) and a negative response in CPI of 1% over three years.

In Appendix Figures IA.2 and IA.3, we repeat this analysis replacing the last variable in y (Dollar index) with consumption. In this specification, we see that a positive one standard deviation in the complex SDF return predicts a significant 0.7% increase in consumption over three years. The remaining impulse responses are quantitatively unaffected, with the

²³Among these ten variables, there are two that we use in levels without taking logs, these are the Federal Funds rate and the Baa-Aaa credit spread, which are already reported in annualized percentages. In this case, the $b_{SDF,h}$ coefficient describes the expected simple (not log) change in the corresponding variable.

²⁴For this analysis, the annual SDF return is defined in standard deviation units, $(\sum_{j=0}^{11} R_{t-j}^M) / \sigma^M$, where σ^M is the sample standard deviation of on the rolling annual return.

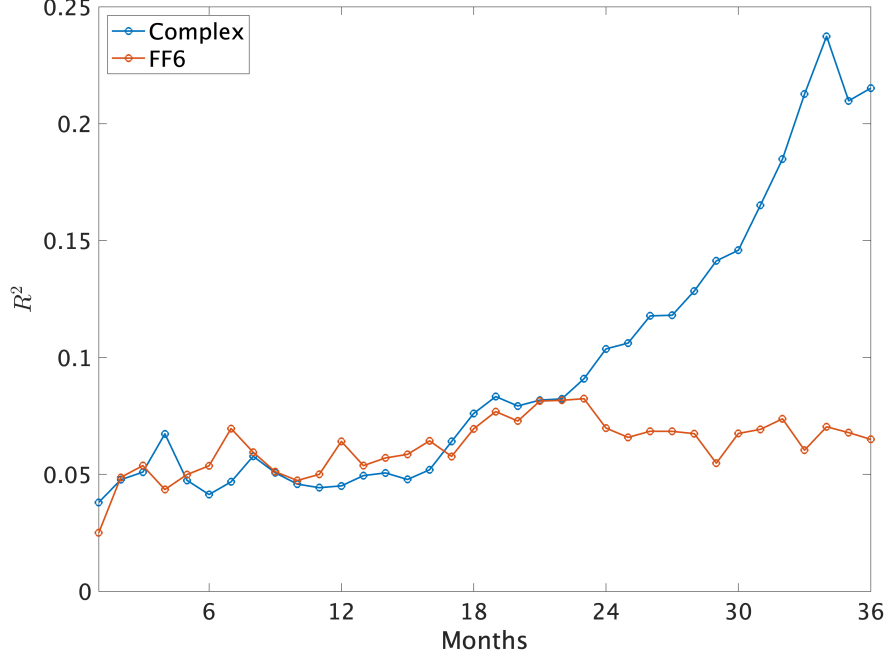


Figure 6: SDF Returns and Future Macroeconomic Outcomes

Note. This figure reports the R^2 from a regression of SDF returns on future macroeconomic outcomes per equation (15).

exception that the response of the credit spread drops more quickly, bottoming out after nine months rather than two years.

The results in Figure 5 illustrate the strength with which the complex SDF predicts a variety of macroeconomic conditions. But how much of the SDF's behavior can be attributed to the market's assimilation of information about future economic fundamentals collectively? We devise an answer to this question by running a somewhat unconventional regression of returns in month t on future realizations for all of the macroeconomic variables in y :

$$R_t^M = a_h + b_h'(y_{t+h} - y_t) + e_t, \quad h = 1, \dots, 36. \quad (15)$$

The R^2 of this regression describes the fraction of SDF return variation that is attributable to information about future macroeconomic conditions over horizons $h = 1$ to 36 months ahead. Figure 6 plots the resulting R^2 for the complex SDF and the FF6 benchmark SDF.

The complex SDF is predominantly driven by information about the macroeconomy over long horizons. Cumulative macroeconomic realizations over the subsequent three years explain 25% of the variation in complex SDF returns. This drops to around 5% if we restrict to economic realizations over the next year. In contrast, only about 5% of the benchmark SDF return is explained by future macroeconomic factors, regardless of the horizon we look at.²⁵

In summary, positive SDF returns predict rising production, employment, sentiment, and housing; they predict decreasing risk, credit spreads, and inflation. Our impulse responses control for lags of y , implying that this predictive information in the complex SDF is not spanned by standard macro variables. Meanwhile, the benchmark SDF shows either little significant predictive power for macro variables, or the sign of the prediction is counterintuitive (as in the case of employment).

The long-horizon impulse responses in Figure 5 are directionally consistent with equilibrium models incorporating recursive preferences and long-run risk, such as [Bansal and Yaron \(2004\)](#) and [Campbell et al. \(2018\)](#). In these models, the SDF reflects investors' expectations about future output uncertainty over the long run. Empirically, long-run risks are notoriously difficult to estimate and researchers continue to debate the reliability of proposed proxies (see e.g. [Liu and Matthies, 2022, 2024](#); [Maio, 2024](#)). Our complex SDF provides new empirical evidence of the long-run risk mechanism in equity market prices.

Our empirical findings about the joint cyclical variation in uncertainty, production, interest rates, and credit spreads are consistent with the equilibrium models of [Gilchrist and Zakrajšek \(2011\)](#); [Gilchrist et al. \(2014\)](#); [Christiano et al. \(2014\)](#). Their focus relates to the empirical fact that credit spreads and risk predict economic conditions ([Gilchrist and Zakrajšek, 2012](#); [Gilchrist et al., 2009](#)). Importantly, since we control for lags of y ,

²⁵Appendix L shows the same R^2 plot for all benchmark models. The patterns for all benchmarks are broadly similar to FF6. DHS is an exception in that it captures more information about future short-term macroeconomic phenomena than the other SDFs.

our analysis indicates that the complex SDF incorporates differential information about the macroeconomy above and beyond the information in prevailing credit spreads and risk.

Finally, our conclusions from the macroeconomic analysis above are robust to a variety of alternative design choices in the analysis of (14), including changing the constituent variable in y , changing the frequency of SDF return, and changing the number of y lags. These results are reported in Appendix L.

3.5 SDF Complexity or SDF Sparsity?

A recent line of financial machine learning research suggests that it is possible to estimate a successful factor pricing model through the imposition of sparsity.²⁶ The evidence indicates that an SDF model with a small number of factors can successfully price a wide variety of test assets. This is exemplified by Kozak et al. (2018), who study a collection of difficult-to-price anomaly portfolios that serve as their test assets. They then show that a simple linear SDF—comprised of just a few principal components of the anomaly portfolios—is powerful for pricing their entire anomaly cross-section.

The notion of SDF sparsity appears at odds with the benefits of complexity that we document above. While our results thus far demonstrate that a complex model can identify efficacious nonlinear pricing factors, is it possible that we also introduce unnecessary redundancy by using many thousands of factors? We investigate this possibility now.

In Figure 2, each point on each curve is a model with a particular number of factors, P , and a particular shrinkage parameter, z . To investigate the potential benefits of sparsity, we replace the P -factors in each model of Figure 2 with a small number K of principal components derived from those P factors. We then re-estimate the SDF coefficients on the K components and track the resulting out-of-sample factor model performance.

Figure 7 reports the results. In the left column, we consider a $K = 5$ component

²⁶This includes Gagliardini et al. (2016), Kelly et al. (2020), Kozak et al. (2020), Lettau and Pelger (2020), and Giglio and Xiu (2021), among others.

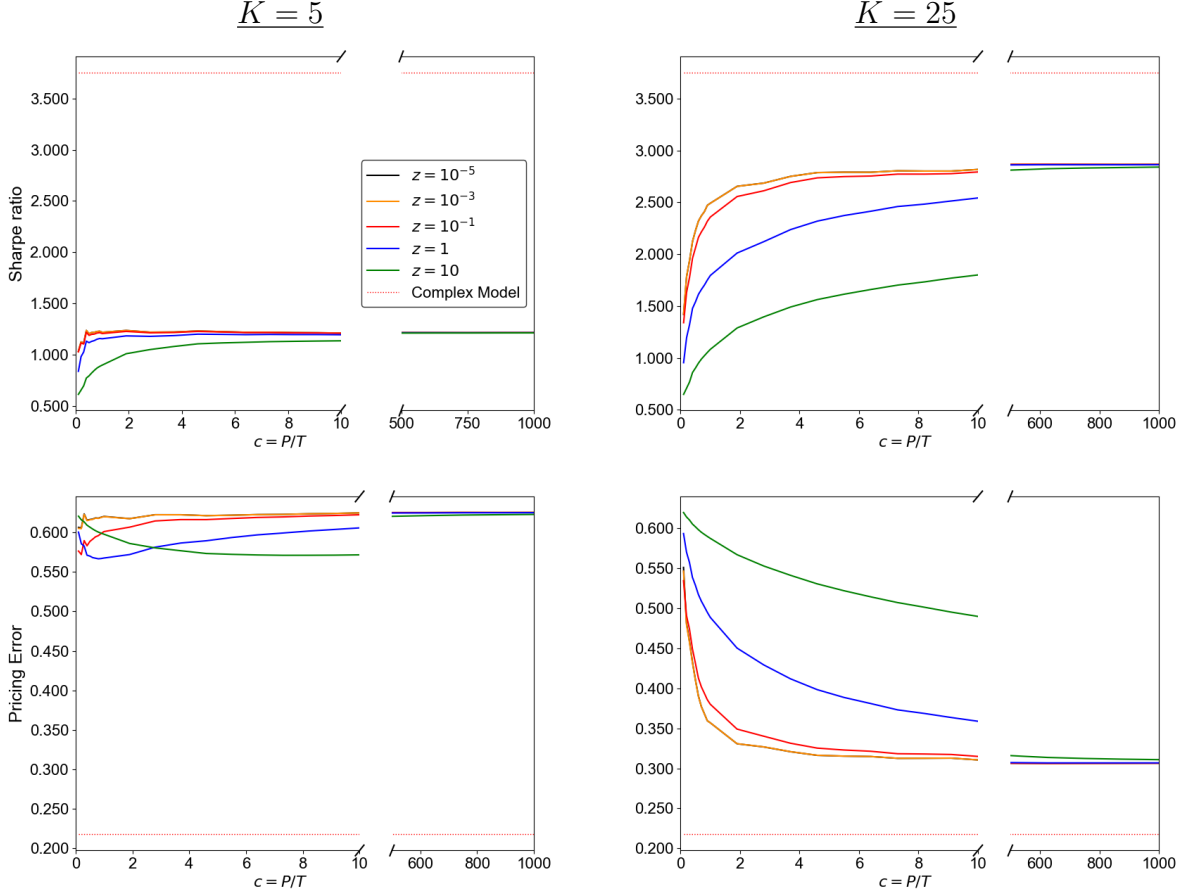


Figure 7: The Effect of Sparsity on Out-of-sample SDF Performance

Note. Realized out-of-sample SDF Sharpe ratio (top row) and pricing error (HJD, bottom row) for complex models with dimension reduction to $K = 5$ (left column) and $K = 25$ (right column) principal components. The horizontal axis shows model complexity $c = P/T$ with P ranging from 36 to 360,000 and $T = 360$ months. For ease of reference, “Complex Model” shows performance of the highest complexity model ($P = 360,000$, $z = 10^{-5}$) from Figure 2 without dimension reduction.

dimension reduction of each complex factor model (corresponding to the K used in some model specifications of Kozak et al., 2018; Kelly et al., 2020). The Sharpe ratio (top panel) of the dimension-reduced SDF is roughly 1.3, compared to 3.7 for the full complexity model. Likewise, pricing errors (bottom panel) are 0.57 for the $K = 5$ SDF versus 0.22 for the full complexity model. In other words, imposing sparsity on the SDF via dimension reduction inhibits SDF performance relative to the unreduced, high-complexity counterpart.

Kozak et al. (2018) consider larger numbers of principal components as well. The right

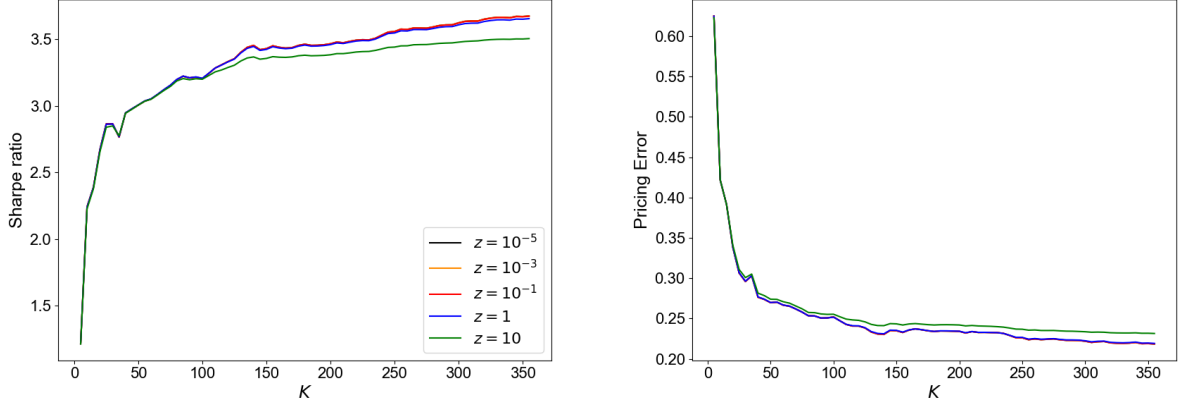


Figure 8: The Effect of Sparsity on Out-of-sample SDF Performance

Note. Realized out-of-sample SDF Sharpe ratio (top row) and pricing error (HJD, bottom row) for complex models with dimension reduction to $K = 5, \dots, 360$ principal components (x-axis).

column of Figure 7 shows that increasing model complexity by using $K = 25$ rather than $K = 5$ principal components leads to a large improvement in out-of-sample performance. But even in this case, the Sharpe ratio remains well below that of the full complexity model (3.7 vs 2.9), and the pricing error remains 41% higher (0.31 vs 0.22).

Two properties of high-dimensional models drive this effect. First, even if the true (unobservable) factor covariance matrix has a few large eigenvalues and, hence, a strong factor structure, the factors become difficult to detect when complexity is high.²⁷ Second and more surprisingly, even low-variance components have a significant Sharpe ratio, so excluding them leads to a large drop in out-of-sample SDF performance. This fact is illustrated in Figure 8. In this analysis, we hold the number of factors in the underlying model fixed at $P = 360,000$, but vary the number of components extracted from the model from $K = 5, 10, \dots, 360$. Adding low variance components typically helps, and never hurts, out-of-sample factor model performance. Based on this analysis, the AIPT’s large factor model conjecture appears better suited for describing market behavior than does the parsimony conjecture of the APT.

²⁷See, e.g., Lettau and Pelger (2020).

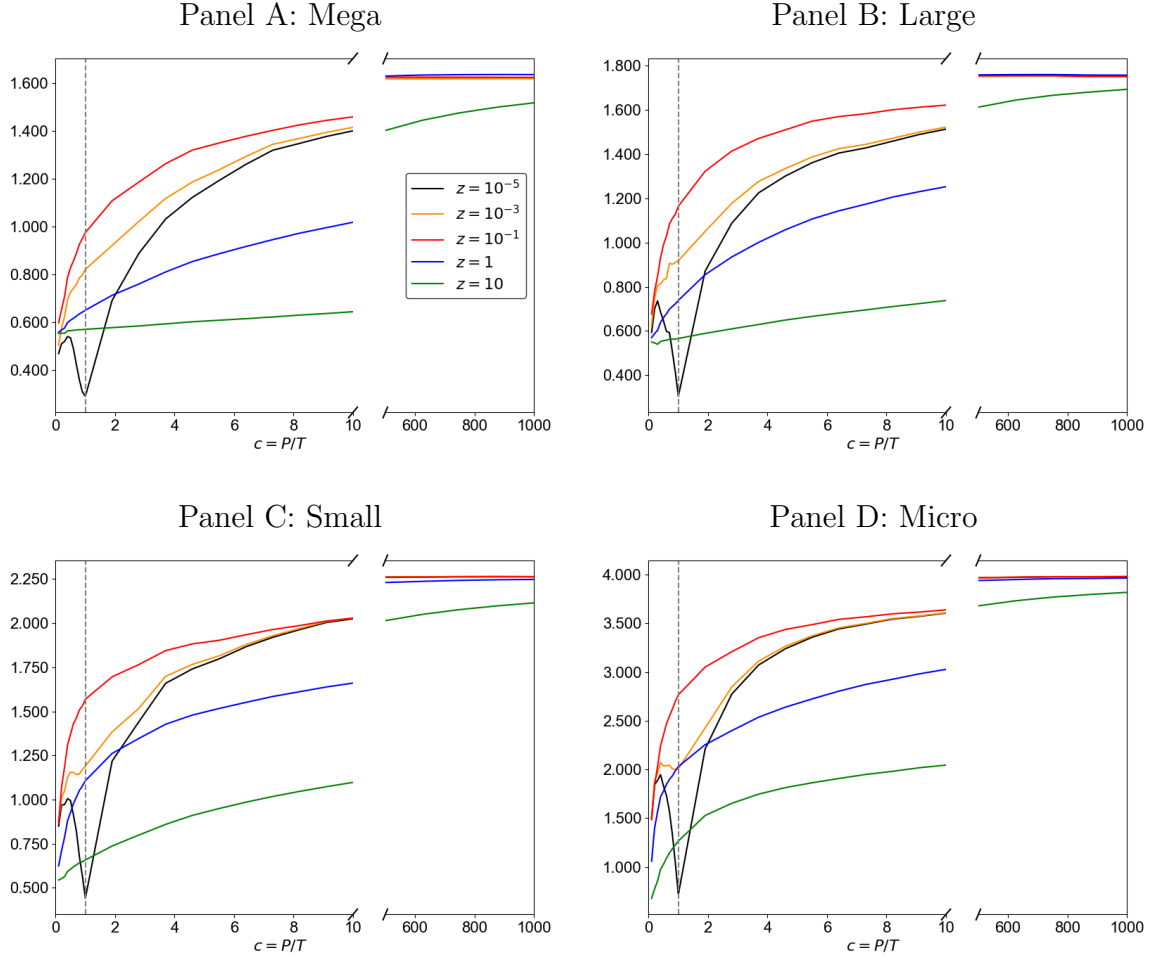


Figure 9: Out-of-sample Sharpe Ratios of Complex SDF Models By Size Group

Note. Realized out-of-sample SDF Sharpe ratio for SDFs estimated from subsamples based on market capitalization.

3.6 Sensitivity to Liquidity

A Sharpe ratio of 3.7 for the high-complexity SDF model suggests that the model is exploiting inefficiencies associated with illiquidity. To understand the role of complexity in factor pricing models while abstracting from the question of asset liquidity, we conduct our empirical analysis separately for different market capitalization groups. We study four size groups

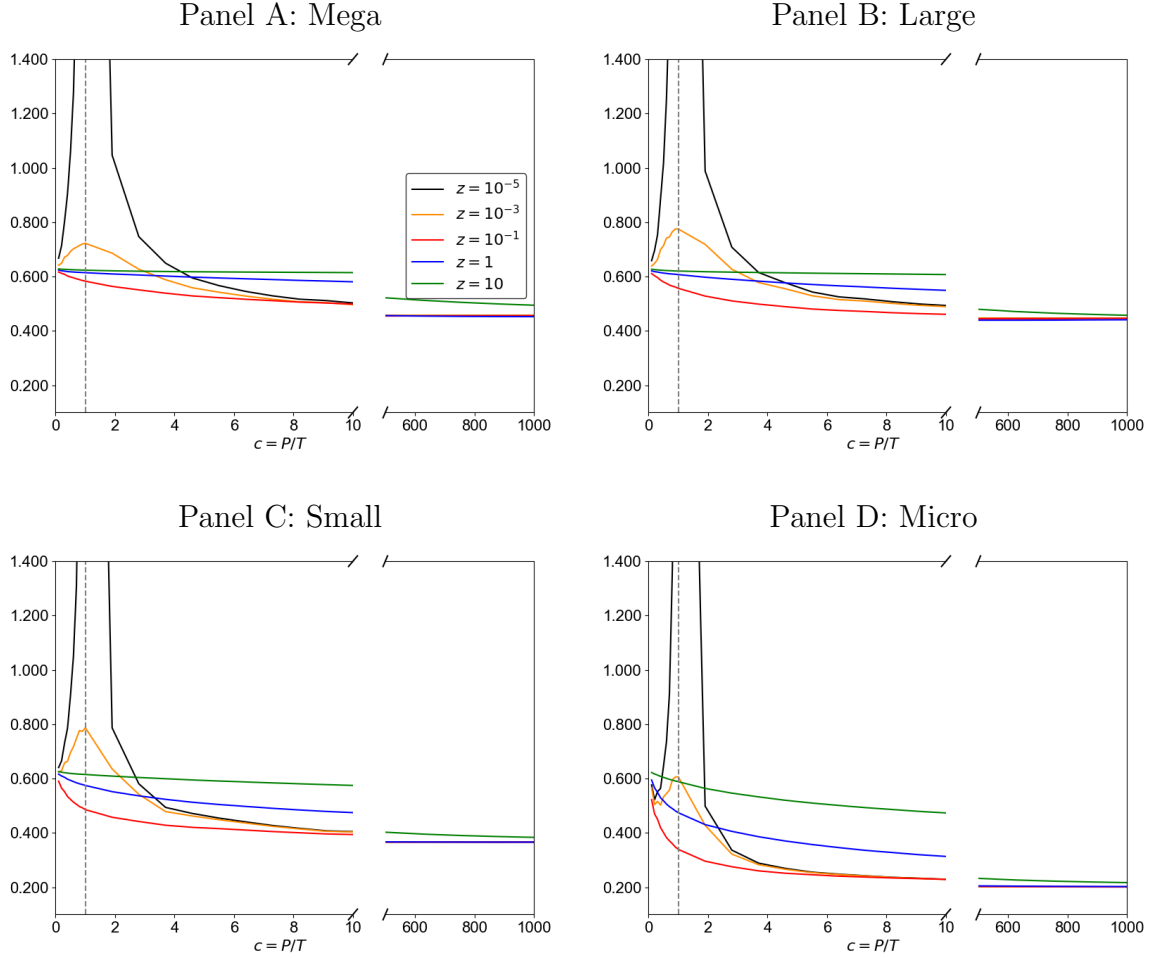


Figure 10: Out-of-sample Pricing Errors of Complex SDF Models By Size Group

Note. Realized out-of-sample SDF pricing error (HJD) for SDFs estimated from subsamples based on market capitalization.

from JKP: mega (largest 20% of stocks based on NYSE breakpoints each period), large (between 80% and 50%), small (between 50% and 20%), and micro (between 20% and 1%).

Figure 9 plots out-of-sample Sharpe ratios for SDFs estimated separately within each size group, while Figure 10 plots pricing errors. The central conclusion from these figures is that the virtue of complexity arises in all stock size groups. In terms of magnitude, Figure 9 indeed suggests that the SDF Sharpe ratios in Figure 2 are heavily influenced by microcap stocks. Yet the SDF Sharpe ratios based on mega or large stocks are on the order of 1.6 and

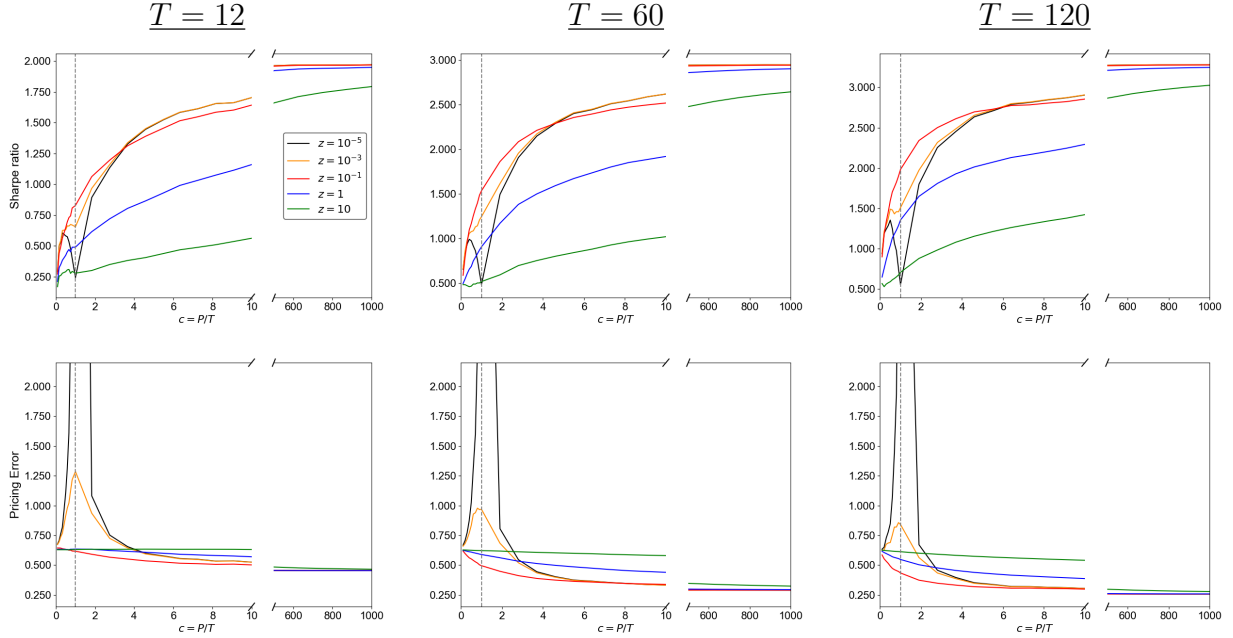


Figure 11: The Effect of Training Sample Size on Out-of-sample SDF Performance

Note. Realized out-of-sample SDF Sharpe ratio (top row) and pricing error (bottom row) training windows of $T = 12, 60$, and 120 months.

1.8, respectively. In other words, a complex factor model derived from mega stocks alone (roughly the 400 largest stocks in the US) produces an out-of-sample SDF Sharpe ratio well in excess of the FF6 model's Sharpe ratio of 0.74 (which uses the full cross-section of stocks). In summary, the patterns documented in our main analysis (Figure 2) do not appear to be driven by illiquidity and limits-to-arbitrage among the underlying assets. Instead, these findings suggest that models with a large number of factors are generally better suited to price assets in the cross-section, regardless of their liquidity.

3.7 Sensitivity to Sample Size

Our main analysis demonstrates the virtue of complexity when estimation is conducted in a rolling 360-month training sample. However, the benefits of complexity can accrue in smaller

training samples. In Figure 11, we plot VoC curves for training samples of $T = 12, 60$ and 120 months.

We find identical patterns in complex factor model behavior in training windows as short as a single year. Complex model performance weakens in shorter training windows as the training process has less information to learn from. The out-of-sample Sharpe ratio for our highest complexity factor model is 1.9, 2.9, 3.3, and 3.7 for $T = 12, 60, 120$, and 360 months, respectively. Likewise, pricing errors reduce from roughly 0.45 for $T = 12$ months to 0.22 for $T = 360$. In appendix Figure IA.9 we compare against simple benchmark SDFs that are also estimated in $T = 12$ month training windows. In this case, the complex model has a smaller margin of improvement over the benchmarks, since their small numbers of parameters make the simple benchmark SDFs resilient to estimation in fairly short samples. The complex Sharpe ratio is 1.9 versus 1.1 for the best simple benchmark (compared to the $T = 360$ case where the complex Sharpe ratio is 3.7 versus 1.2 for the best simple benchmark).

Finally, we show in appendix Figure IA.10 that the decay in empirical model performance from reducing the sample from $T = 360$ to $T = 12$ matches the theoretically expected decay from using small samples. In particular, we simulate data from the our theoretical model calibration in Section 4.7.1. When we reduce the simulated sample size from $T = 360$ to $T = 12$, the decline in simulated out-of-sample Sharpe ratio is quantitatively consistent with our empirical results. Theorem 3 below demonstrates that this decline is explicitly determined by the joint distribution of factor risk premia and the eigenvalues of their covariance matrix. See also Kelly and Malamud (2025) for an analogous result in the (simpler) single-asset case.

3.8 Sensitivity to Activation Functions

Our main empirical model uses random features built from sine and cosine activation functions following Rahimi and Recht (2007). There are, however, a wide variety of

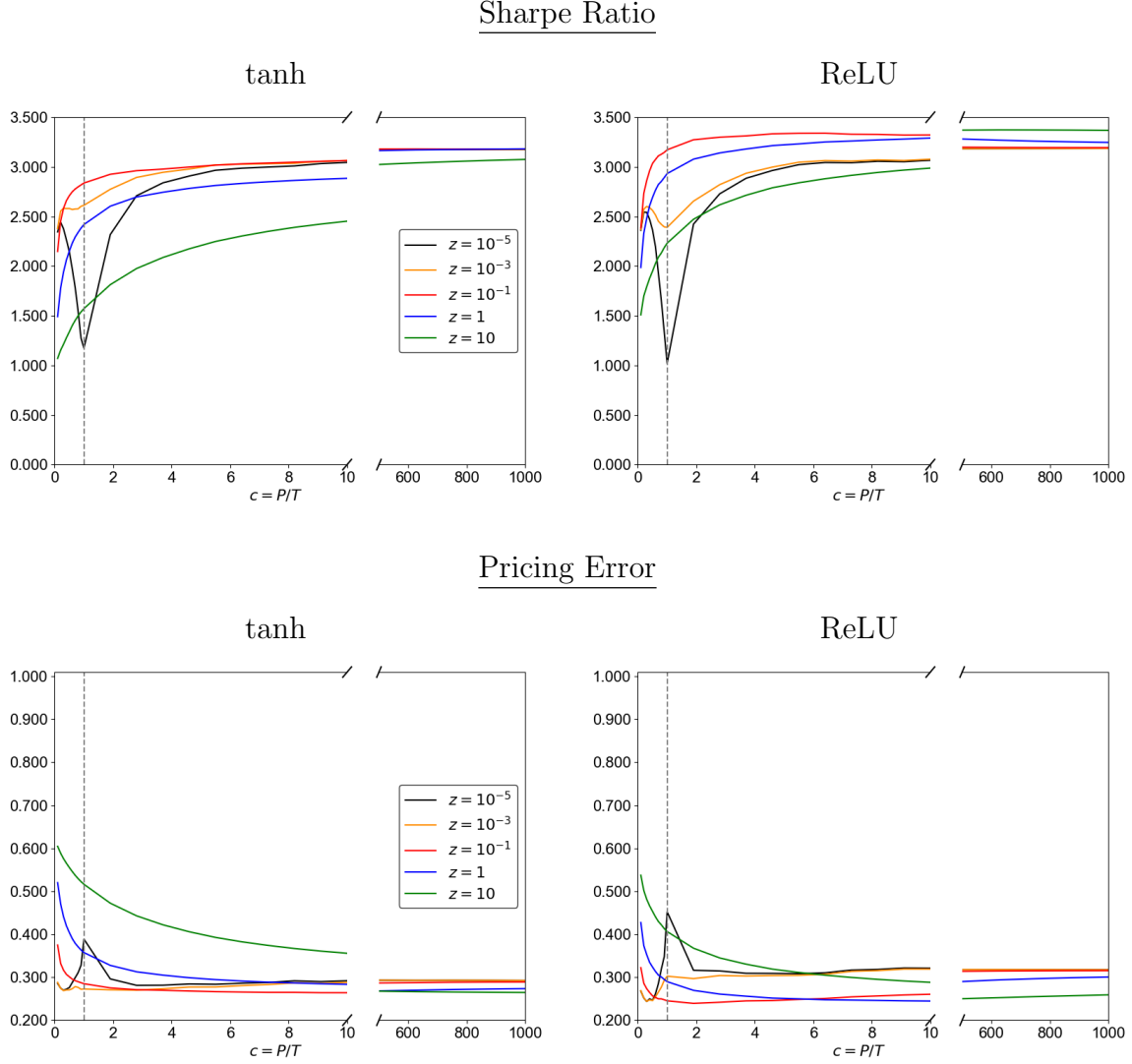


Figure 12: Complex SDF Performance With Alternative Activation Function

Note. This figure repeats the analysis of Figure 2 using either hyperbolic tangent activation (left column) or rectified linear unit activation (right column). Performance is measured as out-of-sample Sharpe ratio (top row) and pricing error (bottom row).

alternative activation functions used in the machine learning literature. In this section we examine the sensitivity of our results to this choice.

Other than the activation function, we hold all design choices from our main analysis in Figure 2 fixed. Specifically, we set $T = 360$, we draw random weights $\{\omega_p\}_{p=1}^P \sim \text{i.i.d. } N(0, I)$,

and we use the same grid of bandwidth parameters γ . We replace the sine/cosine activations in Equation (6) with two of the most common activation functions in the literature, hyperbolic tangent (“tanh”) and rectified linear unit (“ReLU”). Random features with these activations are

$$S_{p,t}^{(\text{tanh})} = \frac{e^{\gamma Z_t \omega_p} - e^{-\gamma Z_t \omega_p}}{e^{\gamma Z_t \omega_p} + e^{-\gamma Z_t \omega_p}} \quad \text{and} \quad S_{p,t}^{(\text{ReLU})} = \max[0, \gamma Z_t \omega_p].$$

Figure 12 shows tanh and ReLU model performance in terms of out-of-sample Sharpe ratio (top panels) and pricing errors (bottom panels), and are directly comparable to the results for the sine/cosine model in Figure 2. We see that the high complexity tanh and ReLU models achieve Sharpe ratios of roughly 3.0 to 3.5 and pricing errors of 0.25 to 0.30, slightly below but broadly in line with the performance of the main sine/cosine specification, and far in excess of benchmark model performance.

Why does the choice of activation function influence out-of-sample performance? Our analytical characterization in Theorem 3 indicates that the construction of random features affects Sharpe Ratios by shaping the alignment between factor risk premia and the eigenvalue distribution of their covariance matrix. This alignment, in turn, is governed by the capacity of non-linear factors to escape linearity and capture genuine non-linear relationships when the sample size T is small. These trade-offs between limited sample size and non-linear representation have been widely studied in prediction contexts and are known to depend on the smoothness of the activation function; see Karoui (2010), Cheng and Singer (2013), and Do and Vu (2013). Extending this understanding to the SDF setting remains an important avenue for future research.

4 AIPT: Theoretical Underpinnings of Complex Factor Models

This section presents artificial intelligence pricing theory. We first present assumptions that formalize the AIPT's core conjecture—that asset prices are determined by a high-dimensional factor structure. From these assumptions, we derive a statistical theory of large factor pricing models. Lastly, we show that the empirical behavior of factor pricing models in Section 3 bears a strikingly close correspondence to theoretically predicted behaviors of the AIPT.

4.1 Assets and Conditioning Information

Our theory begins with the following two assumptions:

Assumption 1 *There exist loadings $S_t \in \mathbb{R}^{N \times P}$, latent factors $\tilde{F}_{t+1}$, and idiosyncratic shocks ε_{t+1} such that returns $R_{t+1} \in \mathbb{R}^N$ satisfy*

$$R_{t+1} = S_t \tilde{F}_{t+1} + \varepsilon_{t+1}, \quad (16)$$

where $\varepsilon_{t+1} = \Sigma_{\varepsilon,t}^{1/2} \tilde{\varepsilon}_{t+1}$ with $E[\varepsilon_{i,t}]$ being i.i.d. with $E[\varepsilon_{i,t}] = 0$, $E[\varepsilon_{i,t}^2] = 1$, $E[\varepsilon_{i,t}^4] < \infty$. The latent factors satisfy $E_t[\tilde{F}_{t+1}] = \nu$, and the conditional covariance matrix $\Sigma_{F,t} = E_t[\tilde{F}_{t+1} \tilde{F}_{t+1}'] - \nu \nu'$ satisfies $\text{tr}(\Sigma_{F,t}) = O(1)$. We also assume that $\nu = P^{-1/2} \Sigma_\nu^{1/2} \hat{\nu}$, where $\hat{\nu} = (\hat{\nu}_i)_{i=1}^P \in \mathbb{R}^P$ is such that $\hat{\nu}_i$ are independent across i and $E[\hat{\nu}_i] = 0$, $E[\hat{\nu}_i^2] = 1$, $E[\hat{\nu}_i^4] \leq K$ for some $K > 0$, and are independent of all other variables; $\|\Sigma_\nu\| = O(1)$ as $P \rightarrow \infty$.

Assumption 2 *We have $S_t = \Sigma^{1/2} X_t \Psi^{1/2}$ for some positive definite matrices Σ, Ψ ; here, the random variables $X_{i,k,t}$ satisfy $E[X_{i,k,t}] = 0$, $E[X_{i,k,t}^2] = 1$, and $X_{i,k,t}$ are independent and $E[X_{i,k,t}^8] < \infty$. In the limit as $N, P \rightarrow \infty$, both Σ and Ψ stay uniformly bounded, $\text{tr}(\Sigma)$ is uniformly bounded, and $\lim_{N \rightarrow \infty} \text{tr}(\Sigma^2)/(\text{tr}(\Sigma))^2 = 0$. Without loss of generality, everywhere in the sequel we assume Σ is normalized so that $\text{tr}(\Sigma) = 1$.*

Factor pricing models summarize the risk-return tradeoff among N risky assets in the economy. To this end, Assumption 1 describes the dependence structure of assets and expected returns. It is a standard conditional factor structure that allows for an arbitrary number of factors P .²⁸ Assumption 1 defines the relevant conditioning information in this economy, which amounts to the conditional factor loadings in the $N \times P$ matrix S_t . We refer to S_t as “characteristics” or “signals” in connection with the empirical asset pricing literature.²⁹ The P -vector of factor risk prices, denoted ν , together with the conditional loadings determine asset expected returns.³⁰ Internet Appendix B further discusses Assumption 1 as a generic statistical specification arising in the context of machine learning factor pricing models.

Next, our theoretical derivations require assumptions about the covariance structure of the signals S_t . By Assumption 2,

$$E[S'_t S_t] = \text{tr}(\Sigma)\Psi \in \mathbb{R}^{P \times P} \quad \text{and} \quad E[S_t S'_t] = \text{tr}(\Psi)\Sigma \in \mathbb{R}^{N \times N}, \quad (17)$$

thus the matrix Ψ captures the covariance of signals across factors and Σ captures their covariance across assets. The assumption of bounded $\text{tr}(\Sigma)$ is a no-arbitrage condition to ensure that the predictable variation in returns stays bounded. The last two limits in Assumption 2 ensure that characteristic-managed portfolios offer a meaningful amount of diversification across stocks.³¹

²⁸The assumption $\text{tr}(\Sigma_{F,t}) = O(1)$ as $P \rightarrow \infty$ is the mathematical formalization of the idea of a factor structure. For example, if Σ_F has a finite rank K with bounded eigenvalues, then this condition is trivially satisfied. The assumption of constant conditional factor premia is without loss of generality since dynamic premia can be subsumed by S_t .

²⁹One can think of the loadings as some function of other underlying conditioning characteristics, or the characteristics could be loadings themselves as in the BARRA model popular among industry professionals.

³⁰The scaling $\nu = P^{-1/2}\Sigma_\nu^{1/2}\hat{\nu}$ with $P^{-1/2}$ is optimal and ensures that the vector of risk premia neither vanish nor explode for large P : $E[\|\nu\|^2] = P^{-1}\text{tr}(\Sigma_\nu) = O(1)$.

³¹For example, suppose that $\text{rank}(\Sigma) = 1$, so that $\Sigma^{1/2} = \pi\pi'$ for some $\pi \in \mathbb{R}^N$. Then, $S_t = \pi\pi'X_t\Psi^{1/2}$ and therefore all factors are given by

$$F_{t+1} = \Psi^{1/2}X'_t\pi(\pi'R_{t+1}), \quad (18)$$

4.2 Characteristic-managed Portfolios

From the equivalence of the conditional SDF and the mean-variance efficient portfolio

$$\tilde{M}_{t+1} = 1 - \pi'_t R_{t+1}, \quad \pi_t = E_t[R_{t+1} R'_{t+1}]^{-1} E_t[R_{t+1}], \quad (19)$$

Assumption 1 implies the following SDF representation.

Proposition 1 *Let $\Sigma_{F,\varepsilon,t} = E_t[\tilde{F}_{t+1} \varepsilon'_{t+1}]$. A conditional stochastic discount factor is*

$$\tilde{M}_{t+1} = 1 - \tilde{w}(S_t)' R_{t+1}, \quad (20)$$

where

$$\tilde{w}(S_t) = \Sigma_t^{-1} \mu_t \quad (21)$$

is the conditional mean-variance efficient portfolio with

$$\begin{aligned} \Sigma_t &= S_t(\Sigma_{F,t} + \nu\nu')S'_t + S_t\Sigma_{F,\varepsilon,t} + \Sigma'_{F,\varepsilon,t}S'_t + \Sigma_{\varepsilon,t} \\ \mu_t &= S_t\nu \end{aligned} \quad (22)$$

and

$$E_t[R_{i,t+1} \tilde{M}_{t+1}] = 0, \quad i = 1, \dots, N. \quad (23)$$

The SDF in (20) is stated as a conditional portfolio of the basic risky assets, R , as discussed in the empirical framework of Section 2. Estimating the conditional mean-variance portfolio

implying that all factor returns are proportional to returns on a single portfolio, $\pi' R_{t+1}$, with no idiosyncratic return. Assumption 2 ensures that such pathological situations cannot occur.

of basic assets is challenging. Not only does it require estimates of means and covariances for a large number of assets, but it also requires these moments in *conditional* terms.

To avoid the difficult task of modeling the conditional distribution of basic assets, it is common in the empirical literature to instead study characteristic-managed portfolios, or “factors,”³²

$$F_{t+1} = S_t' R_{t+1}. \quad (24)$$

The conjecture is that by studying the *unconditional* properties of factors, we can learn about the conditional properties of asset markets. For example, [Kozak et al. \(2020\)](#) approximate the conditional SDF using the unconditional mean-variance efficient portfolio of managed portfolios.³³ Yet it is easy to see that the mean-variance portfolio of F_t ,

$$\lambda = E[F_{t+1} F_{t+1}']^{-1} E[F_{t+1}] \quad (25)$$

is generally different from the conditionally efficient portfolio of basic assets that determines the true SDF in (20).

We prove in Proposition 2 that when the number of factors is large, the conjecture indeed holds, and the unconditional optimal portfolio of factors coincides with the true conditional SDF. While somewhat technical, the logic for this result is straightforward. Factors interact base asset returns with conditioning characteristics, and, in the large P limit, these interactions encode all available conditioning information. Therefore, we can rely on *unconditional* properties of factors to capture the conditional properties of the underlying assets.³⁴

³²The literature often refers to the managed portfolios F_t as “factors,” and we adhere to this slight abuse of nomenclature when the difference between F_t and the true factors $\tilde{F}_t$ is clear.

³³Relatedly, [Kozak and Nagel \(2023\)](#), [Filipovic and Schneider \(2024\)](#), and [Chernov et al. \(2025\)](#) discuss conditions under which managed portfolios “span” the conditional SDF.

³⁴See Internet Appendix E for technical details.

Proposition 2 (Characteristic-Managed Portfolios Span the SDF) *Let*

$$M_{t+1} = 1 - \lambda' F_{t+1} = 1 - w(S_t)' R_{t+1}, \text{ with } w(S_t) = S_t \lambda, \quad (26)$$

be the factor approximation for the SDF with λ given by (25). Suppose that $\Sigma_{\varepsilon,t} = \tilde{\Sigma}_{\varepsilon,t} + \sigma_\varepsilon^2 I$, where $\lim_{N \rightarrow \infty} \text{tr}(\Sigma \tilde{\Sigma}_{\varepsilon,t}) = 0$ and $\sigma_\varepsilon^2 \in \mathbb{R}_+$. Then, the difference $M_{t+1} - \tilde{M}_{t+1}$ converges to zero in L_1 and the Sharpe ratio of $w(S_t)' R_{t+1}$ converges to that of $\tilde{w}(S_t)' R_{t+1}$ as $N, P \rightarrow \infty$.

Proposition 2 implies that, when N, P are large enough,³⁵

$$\lambda' F_{t+1} = \lambda' S_t' R_{t+1} \approx \nu' S_t' \Sigma_t^{-1} R_{t+1}. \quad (27)$$

Thus, the surprising and very convenient implication of Proposition 2 is that model complexity, in fact, *simplifies* asset pricing. Much like continuous time limits simplify a variety of asset pricing calculations, large factor limits reduce conditional SDF modeling to an unconditional problem. Therefore, in the remaining theoretical development, we leave behind the base assets R_t and work directly with managed portfolios, F_t .³⁶

4.3 Large Models as Approximations

As the George Box adage goes, “All models are wrong, but some are useful.” In reality, the true data-generating process is unknown to the researcher. A core premise of modern

³⁵As noted by Chernov et al. (2025), the omission of Σ_t^{-1} in the construction of F_{t+1} introduces unhedged components, resulting in an efficiency wedge between the factor portfolio and the conditionally efficient portfolio. However, Proposition 2 shows that this wedge becomes asymptotically negligible.

³⁶Relatedly, Assumptions 1 and 2 govern the properties of R_t , which serves two purposes. First, the assumptions ensure that characteristic-managed portfolios span the SDF in the high-complexity limit, per Proposition 2. Second, they imply that characteristic-managed portfolios satisfy the technical conditions of random matrix theory (see Theorem 4 in the Internet Appendix). Once these conditions are established, the origin of factors is no longer relevant, and the theory can treat them in the abstract when proving our main result in Theorem 3. Formally, the idea that conditionally efficient portfolio of stocks can be represented as an unconditional portfolio of a rich enough set of nonlinear managed portfolios is discussed in (Hansen and Richard, 1987). In general, portfolio weights for each asset depend on the entire set of conditioning variables (e.g., characteristics of *all* assets). By contrast, Proposition 2 implies that we can restrict our attention to portfolio weights that depend only on own asset characteristics.

artificial intelligence (and nonparametric statistics more broadly) is that large models are useful because they provide flexible approximations of the unknown truth. However, while added complexity brings better approximations, it also brings the usual statistical challenges of heavy parameterization in the face of limited data.

In this section, we design a framework to understand the out-of-sample behavior of factor pricing models as we gradually expand their parameterization. To embed this investigation in our theoretical environment, we consider a true SDF with a large number of factors, P^* . We then consider an empirical model that approximates the SDF with only a fraction $q = \frac{P}{P^*} \leq 1$ of those factors.

Assumption 3 *A misspecified empirical model of size $P < P^*$ ($q = \frac{P}{P^*}$) uses a subset of factors, $F_{t+1}(q) = (F_{i,t+1})_{i=1}^P$ with covariance matrix $E[F_t(q)F_t(q)'] \in \mathbb{R}^{P \times P}$.*

We are interested in characterizing out-of-sample factor model behavior as we expand the subset of factors from $P = 1, \dots, P^*$, gradually reducing the degree of misspecification in the approximating model (when $P_1 = P$ the model is correctly specified) while increasing the number of parameters that must be estimated.

Our theoretical derivations are based on large P limits. The number of factors in the true, unattainable factor model is $P^* \rightarrow \infty$. We conceptualize the AIPT as a rich but nevertheless imperfect empirical model with $P \rightarrow \infty$ and $q \leq 1$. We now develop the limiting behavior of large empirical factor models using random matrix theory.

4.4 Feasible and Infeasible SDF Estimators

Proposition 2 directly motivates our empirical approximation to the SDF in equation (4) and the ridge estimator in equations (9) and (12). Following Assumption 3, the ridge estimator has access to a fraction q of the factors in the true model. Given a training sample of size T , the empirical model has complexity $c = P/T$. We define the *ridge SDF estimator* for an

empirical model of complexity c and specification q as

$$\hat{\lambda}(z; q; c) = (zI + \hat{E}[F_t(q)F_t(q)'])^{-1}\hat{E}[F_t(q)] \quad (28)$$

with corresponding portfolio return and SDF of

$$\hat{R}_{T+1}^M(z; q; c) = \hat{\lambda}(z; q; c)'F_{T+1}(q), \quad \hat{M}_{T+1}(z; q; c) = 1 - \hat{R}_{T+1}^M(z; q; c),$$

where $\hat{E}$ denotes sample average over the T training observations.

It will be useful to contrast $\hat{\lambda}(z; q; c)$ with the *infeasible ridge SDF estimator*

$$\lambda(z; q) = (zI + E[F(q)F(q)'])^{-1}E[F(q)] \quad (29)$$

and its return and SDF

$$R_{T+1}^M(z; q) = \lambda(z; q)'F_{T+1}(q), \quad M_{T+1}(z; q) = 1 - R_{T+1}^M(z; q). \quad (30)$$

This estimator is infeasible because it relies on the population mean and covariance of factors rather than their sample counterparts. The special case $\lambda = \lambda(0, 1)$ corresponds to the true SDF in (26). Naturally, as z increases from zero or as q decreases from one, the Sharpe ratio of $R_{T+1}^M(z)$ declines. The portfolio $\lambda(z; q)$ is an intermediate object between the true SDF $\lambda(0, 1)$ and the feasible estimator $\hat{\lambda}(z; q; c)$.

A remarkable aspect of the theoretical results below is that we can fully characterize the properties of $\hat{\lambda}(z; q; c)$ in terms of the infeasible estimator $\lambda(z; q)$. Our analysis focuses on the key properties that summarize the theoretical behavior of a factor-based SDF: its mean and variance (which in turn dictate its Sharpe ratio and pricing errors). Because they are central to our derivation below, we give special notation to these properties for the infeasible

ridge SDF:

$$\mathcal{E}(z; q) \equiv E[R_{T+1}^M(z; q)], \quad \mathcal{V}(z; q) \equiv \text{Var}[R_{T+1}^M(z; q)]. \quad (31)$$

4.5 The Ridge SDF and Random Matrix Theory

It is perhaps easiest to understand our large P theory for factor model behavior through the calibrations in Section 4.7, and readers interested in a non-technical discussion of the theory may proceed there directly. For interested readers, this section provides a brief overview of the random matrix theory concepts behind our analysis, and Section 4.6 gives a detailed presentation of our main result.

The central challenge to understanding $\hat{\lambda}(z; q; c)$ is the $P \times P$ matrix $\hat{E}[F_t(q)F_t(q)']$ whose dimension grows with the number of factors. The limiting properties of this object require the apparatus of random matrix theory (RMT).

One technical insight in our analysis is that incorporating ridge regularization in the SDF estimation problem allows us to extend core RMT results, such as the [Marčenko and Pastur \(1967\)](#) theorem, to asset pricing analysis. We derive the necessary extensions of [Marčenko and Pastur \(1967\)](#) to describe expected out-of-sample SDF behavior in terms of three objects: the eigenvalue distribution of the factor population covariance matrix, $E[FF']$, the number of parameters per training observation (i.e., complexity, c), and the extent of model misspecification (governed by q). Recall that our factors are defined as $F_{t+1} = S_t' R_{t+1}$, with signals and returns satisfy Assumption 1 and 2. As we show in the Internet Appendix, the only properties of the data-generating process that matter for portfolio performance in the high complexity limit are Ψ and ν . Let $\Psi(q) = \Psi[1 : P, 1 : P]$ be the Ψ sub-matrix corresponding to the first P signals.

We need the following regularity condition on the covariance matrix of factors.

Assumption 4 *In the limit as $P \rightarrow \infty$, the eigenvalue distribution of $\Psi(q)$ converges to a*

limit distribution supported on a bounded interval. Furthermore, $\lim_{P \rightarrow \infty} P^{-1} \text{tr}(\Sigma_\nu[1 : P, 1 : P] \Psi(q)^k)$ exists for each $k \in \mathbb{N}$. Furthermore, in addition to Assumptions 1 and 2, we have $E[\tilde{\varepsilon}_{i,t}^3] = E[X_{i,j,t}^3] = 0$ for all i, j , and the latent factors satisfy $\tilde{F}_t = \nu + \Sigma_{F,t}^{1/2} \tilde{X}_t$, where $\tilde{X}_t = (\tilde{X}_{i,t})_{i=1}^P$ with $\tilde{X}_{i,t}$ being i.i.d. with $E[\tilde{X}_{i,t}] = 0$, $E[\tilde{X}_{i,t}^2] = 1$.

The first part of Assumption 4 ensures that the joint structure of signal covariance and risk premia are well-behaved in the $P \rightarrow \infty$ limit. The second part of the assumption ensures that the covariance of latent factors fully captures their dependency structure (e.g., this is always the case when F_t are normally distributed).

The central object that dictates the behavior of large factor models is the eigenvalue distribution of the factor covariance matrix, which our derivations represent as a Stieltjes transform. This transform for the eigenvalue distribution of $E[F_t(q)F_t(q)']$ in the large P limit is denoted

$$m(-z; q) = \lim_{P \rightarrow \infty} \frac{1}{P} \text{tr} \left((E[F_t(q)F_t(q)'] + zI)^{-1} \right), \quad (32)$$

and, similarly, the infeasible moments (31) converge to well-defined limits

$$\mathcal{E}(z; q), \mathcal{V}(z; q).$$

We denote the limiting eigenvalue distribution of $\hat{E}[F_t(q)F_t(q)']$ as

$$m(-z; q; c) = \lim_{\substack{P \rightarrow \infty \\ P/T \rightarrow c \\ P/P^* \rightarrow q}} \frac{1}{P} \text{tr} \left(\left(zI + \hat{E}[F_t(q)F_t(q)'] \right)^{-1} \right) \quad (33)$$

The main challenge of the AIPT is that when the factor model is heavily parameterized ($c > 0$), $\hat{E}[F_t(q)F_t(q)']$ converges not to $E[F(q)F(q)']$ but to a distortion of it. In Theorem 4 in Internet Appendix D, we apply an extension of the Marčenko and Pastur (1967) theorem due to Bai and Zhou (2008) to our asset pricing environment in order to establish the explicit

mapping between $\hat{E}[F_t(q)F_t(q)']$ and $E[F_tF_t']$. Namely,

$$m(z; q; c) = \frac{1}{1 - c - c z m(z; q; c)} m\left(\frac{z}{1 - c - c z m(z; q; c)}; q\right). \quad (34)$$

The nonlinear master equation in (34) has a unique positive solution $m(z; q; c)$ for any $z < 0$. This solution is a function of only $m(z; q)$ and c . Thus, equation (34) links $m(-z; q; c)$ in (33) to $m(-z; q)$ in (32). When complexity is low and misspecification is small, i.e., when $c \approx 0$ and $q \approx 1$, (34) implies $m(z; q; c) \approx m(z; q)$, as predicted by the standard law of large numbers. However, for $c > 0$, the empirical Stieltjes transform and the “true” Stieltjes transform decouple in a *deterministic* manner. When complexity is high (e.g., $c > 1$), large parts of information about $m(z)$ are lost in finite samples even in the correctly specified ($q = 1$) case, and this information loss is, of course, further exacerbated by misspecification.

The last object we need to introduce is the particular formulation of the HJD that we use to characterize pricing errors. In the high complexity regime, the exact details of computing the out-of-sample HJD are important. We assume that the data sample is split into two sets: in-sample data indexed as $t \in [1, T]$ and out-of-sample data indexed as $t \in (T + 1, T + T_{OS}]$, where T_{OS} is the number of out-of-sample periods. For simplicity, we focus our theoretical analysis on the case in which the test assets are the factors $F_t(q)$ used to estimate the SDF.³⁷ We thus define the out-of-sample HJD as

$$\mathcal{D}_{OS}^{HJ}(z; q; c) = \mathcal{E}_{OS}(z; q; c)' B_{OS}^+ \mathcal{E}_{OS}(z; q; c),$$

Where $\mathcal{E}_{OS}(z; q; c) = \frac{1}{T_{OS}} \sum_{t \in (T, T+T_{OS}]} F_t(q) \hat{M}_t(z; q; c)$ is the out-of-sample pricing error vector and $B_{OS} = \frac{1}{T_{OS}} \sum_{t \in (T, T+T_{OS}]} F_t(q) F_t(q)'$ is the out-of-sample test asset second moment.

³⁷The pricing error properties of the SDF are particularly tractable to derive when the test assets are the P factors $F_t(q)$, but this point can be generalized this at the cost of further notation and derivations.

4.6 The Theoretical Behavior of Large Factor Models

Characterizing the behavior of factor portfolios in high dimensions requires tools from RMT. Established RMT results (see, e.g. [Hastie et al., 2019](#)) apply to a setting with $P, T \rightarrow \infty$ and are based on assumptions of independence of higher moments of the underlying high-dimensional variables. However, these assumptions are violated in our setting because factors $F_{t+1} = S'_t R_{t+1}$ exhibit highly complex non-linear dependencies. This makes our setting significantly more complex and requires developing novel mathematical tools for dealing simultaneously with three infinities: $N, T, P \rightarrow \infty$. We develop these tools in the Internet Appendix. They allow us to prove our main theoretical result below. One striking implication of this result is the fact that the rate at which $N \rightarrow \infty$ is irrelevant: The only thing that matters for the asymptotic behavior is complexity $c = \lim P/T$. The relative rate at which N goes to infinity versus P, T has no impact on the asymptotics.³⁸

We now state our main theoretical result, which describes the expected out-of-sample properties of large factor models.

Theorem 3 *In the limit as $N, P, T \rightarrow \infty$, $P/T \rightarrow c$, $P/P^* \rightarrow q$, the expected out-of-sample moments of the ridge SDF portfolio satisfy*

- i. $\lim E[\hat{R}_{T+1}^M(z; q; c)] = \mathcal{E}(Z^*(z; q; c); q)$ where
 $Z^*(z; q; c) = z(1 + \xi(z; q; c)) \in (z, z + c)$ and $\xi(z; q; c) = \frac{c(1 - m(-z; c; q)z)}{1 - c(1 - m(-z; c; q)z)}$,
- ii. $\lim \text{Var}[\hat{R}_{T+1}^M(z; q; c)] = \mathcal{V}(Z^*(z; q; c); q) + G(z; q; c) \mathcal{R}(Z^*(z; q; c); q)$ where
 $G(z; q; c) = (z\xi(z; q; c))' \in (0, cz^{-2}]$, $\mathcal{R}(z; q) \equiv (1 - \mathcal{E}(z; q))^2 + \mathcal{V}(z; q)$
- iii. $\lim \frac{\text{Var}[\hat{R}_{T+1}^M(z; q; c)]}{E[\hat{R}_{T+1}^M(z; q; c)]^2} = (1 + G(z; q; c)) \frac{1}{\mathcal{S}^2(Z^*; q)} + G(z; q; c) \left(\frac{1 - \mathcal{E}(Z^*; q)}{\mathcal{E}(Z^*; q)} \right)^2$,
where $\mathcal{S}^2(z; q) = \mathcal{E}^2(z; q)/\mathcal{V}(z; q)$
- iv. $\lim E[\mathcal{D}_{OS}^{HJ}(z; q; P; T) - \bar{F}'_{OS} B_{OS}^+ \bar{F}_{OS}] = -1 + (1 + G(z; q)) \mathcal{R}(Z^*(z; q; c); q)$.

³⁸In fact, as we show in the Internet Appendix, the finite- N error is controlled solely by $\text{tr}(\Sigma^2)/\text{tr}(\Sigma)^2$, the Herfindahl index of the eigenvalues of Σ .

The central theme in Theorem 3 is that the large number of factors relative to the number of training observations limits the estimator’s ability to learn the true parameters. When $c > 0$, there are too many parameters and too few data points for the estimator to converge to its population counterpart. This failure to fully hone in on the truth results in an asymptotic wedge between the out-of-sample performance of the trained model and that of the true model. We refer to this as “limits to learning.” Perhaps surprisingly, Theorem 3 shows that we can explicitly quantify the effects of limits to learning using only knowledge of the factors’ sample covariance in the training data.

Limits to learning manifest in two ways—implicit shrinkage and complexity risk—which impact the SDF’s out-of-sample properties. We describe these properties now.

4.6.1 Expected Return

Part *i.* of Theorem 3 describes how complexity inhibits the expected return of the SDF portfolio. It describes the expected return in terms of the infeasible ridge portfolio’s return. The key to this is the function $Z^*(z; q; c)$, which is the “implicit shrinkage” of the feasible estimator. Z^* is monotone increasing in c . In other words, high complexity imposes additional shrinkage on the SDF, above and beyond the explicit ridge shrinkage z .³⁹

If the model is correctly specified ($q = 1$) and with a complexity of zero, then $Z^*(z; 0; 1) = z$ and the feasible SDF’s expected return converges to the infeasible expected return, $E[\hat{\lambda}(z)'R_{T+1}] \rightarrow E[\lambda(z)'R_{T+1}] = \mathcal{E}(z)$. But holding z fixed, a rise in complexity to $c > 0$ induces additional bias in the estimator and drives down the expected return of the SDF portfolio. By how much? By the same amount that the expected return drops when the infeasible portfolio’s shrinkage rises from z to $Z^*(z; 1; c)$. In other words, the challenge of learning in a complex setting is equivalent to knowing the true factor moments but being

³⁹Complexity generates this additional implicit shrinkage in an intuitive way. Holding z fixed, if we increase P , we cannot raise $\|\lambda\|^2$ further due to the ridge penalty. By adding more parameters, we can only continue to satisfy the ridge constraint by shrinking the $\hat{\lambda}$ vector further.

forced to use an unduly large shrinkage. Remarkably, $Z^*(z; 1; c)$ is available in closed form thanks to the expression for $\xi(z; q; c)$ from RMT.

The monotonicity of $Z^*(z; q; c)$ in z means that out-of-sample expected returns are highest with minimal shrinkage. But even in the ridgeless limit when $z \rightarrow 0$, Theorem 3 shows there are limits to learning. In particular, there is an unavoidable reduction in expected return because $Z^*(z; q; c)$ is uniformly bounded away from zero in the high complexity regime ($c > 1$).

Thus, limits to learning hold even if the model is correctly specified. But this case is unrealistic; we can't ever expect to achieve an empirical model that nests all relevant conditioning information or uses that information in its proper nonlinear form. Furthermore, under correct specification, comparative statics for complexity change both the empirical *and the true model* as we vary c . Thus, these theoretical comparative statics cannot be taken to data.

The empirically relevant comparative statics must consider a single true DGP, in essence fixing the set of true factors and varying the number of those factors that the empirical model has access to. We can conceptualize these comparative statics by varying q . When the model is misspecified ($q < 1$), expected returns are hindered further because the model relies on incomplete information. But the effect of q is subtle. Holding T fixed, a higher q has two opposing effects on the expected SDF return. First, larger q means that there are more empirical factors, so c rises as well, which exacerbates the limits to learning. However, higher q also reduces specification bias and, therefore, improves the approximation power of the model. This shows up as a smaller implicit shrinkage, Z^* . Which effect dominates depends on the eigenvalue distribution of the factors. As we will see in the calibration below, when there is a concentrated eigenvalue distribution and thus a few dominant factors—as in the assumptions of the APT—the cost of complexity (i.e., more severe limits to learning) dominates the approximation benefits and high complexity may hurt expected SDF returns

on net. But when the eigenvalue distribution is dispersed, and there are many relevant factors—i.e., under the conditions of the AIP—approximation gains dominate, and high complexity leads to better out-of-sample SDF returns.

4.6.2 Variance

Part *ii.* of Theorem 3 regards the variance of the ridge SDF portfolio. On the right side of (35), the first term relates to the implicit shrinkage effect discussed above. In the large P limit, the feasible SDF portfolio with parameter z has the same volatility as the infeasible portfolio with a larger ridge parameter $Z^*(z; q; c) > z$. Due to this implicit shrinkage and the monotonicity of Z^* in c , SDF portfolio variance decreases with model complexity when $c > 1$.

While the first term in (35) reflects the implicit shrinkage in large factor models, the second term represents a different phenomenon that we call “complexity risk.” Complexity risk can be thought of as sampling variation that exists even in the large T limit. It is a pure-variance effect governed by the function $G(z; q; c)$, independent of expected factor returns, and only depends on the eigenvalue distribution of $E[F_t(q)F_t(q)']$. For a simple model, $c = 0$, there are infinitely more observations than parameters, so the SDF estimator converges to a non-random limit, thus $G(z; q; 0) = 0$, and there is no complexity risk. However, when $c > 0$, sampling variation survives even in the large T limit because the number of parameters is too large to be accurately informed by the data. Complexity risk is a second-moment manifestation of limits to learning. The behavior of SDF portfolio variance is driven primarily by c and is relatively insensitive to the extent of model misspecification, q .

4.6.3 Sharpe Ratio and Pricing Error

Part *iii.* of Theorem 3 describes the SDF portfolio’s limiting Sharpe ratio. First, focusing on the effect of complexity, consider the correct specification case, $q = 1$. Building on

i. and *ii.*, a rise in complexity unambiguously decreases the expected SDF return due to implicit shrinkage. The effect on the variance is mixed—implicit shrinkage lowers variance, but complexity risk raises it. The net effect of complexity is an unambiguous decrease in the Sharpe ratio in the correctly specified setting.

As emphasized above, while the $q = 1$ case is helpful for developing intuition around complexity effects, our real interest lies in understanding SDF Sharpe ratios in the misspecified setting. In this case, the effect of complexity on the Sharpe ratio is ambiguous and depends on the eigenvalue distribution of the factor covariance. When eigenvalues are concentrated, complexity has a small net effect on SDF Sharpe, and can potentially reduce it. However, when the eigenvalue distribution is relatively flat, the Sharpe ratio tends to increase with complexity, and the gains from complexity are potentially large.

In the high complexity limit, pricing error is the mirror image of Sharpe ratio. When conditions are such that Sharpe ratio increases with complexity, it is also the case that pricing errors are reduced by complexity.

In summary, Theorem 3 characterizes the subtle mechanisms through which complexity determines a model’s out-of-sample performance. Complexity introduces a tradeoff between the quality of a model’s approximation for the unknown truth on the one hand versus limits to learning (implicit shrinkage and complexity risk) on the other. In terms of the effect on the SDF expected return, a larger model reduces specification bias, which raises expected returns, but it also induces additional implicit shrinkage, which lowers returns. While complexity’s implicit shrinkage has a dampening effect on SDF volatility, it at the same time raises volatility by producing a sampling variation (i.e., complexity risk) that survives in the large T limit. When the marginal approximation benefits are large relative to the loss due to limits to learning, we find that the complexity raises the out-of-sample SDF Sharpe ratio and reduces pricing errors. Under these conditions, there is a virtue of complexity, and large factor models dominate simple, parsimonious models. Theorem 3 shows that the eigenvalue

distribution among factors is critical for determining whether complexity is a virtue or a vice.

4.7 Illustrating the Theory

Theorem 3 derives foundational results for understanding the role of complexity in factor pricing models. However, the analytical expressions are dense and can be difficult to parse. In this section, we attempt to bring Theorem 3 to life by visualizing the role of complexity in different data environments. We present two calibrations of our theory: one is an AIPT model with hundreds of common factors, and the other is an APT model with only a few common factors.

4.7.1 Calibration 1: The AIPT

As emphasized in the theory discussion above, the eigenvalue distribution of $E[FF']$ is a critical driver of out-of-sample SDF behavior. Our calibrations assume that the eigenvalue distribution of $E[FF']$ is diagonal with eigenvalues

$$\eta = \{1/p^\alpha\}_{p=1}^{P^*}, \alpha \geq 0. \quad (35)$$

The larger the calibrated parameter, α , the more concentrated the eigenvalue distribution, and the fewer prominent components there are among the factor portfolios F . The AIPT and APT calibrations differ in only their choice of α .

The AIPT calibration sets $\alpha = 0.5$. The concept of “effective rank” is a useful means of interpreting the α parameter. Defined as⁴⁰

$$\text{EffRank}(E[FF']) = (\text{tr } E[FF'])^2 / \text{tr}(E[FF']^2),$$

⁴⁰See, e.g., [Bartlett et al. \(2020\)](#).

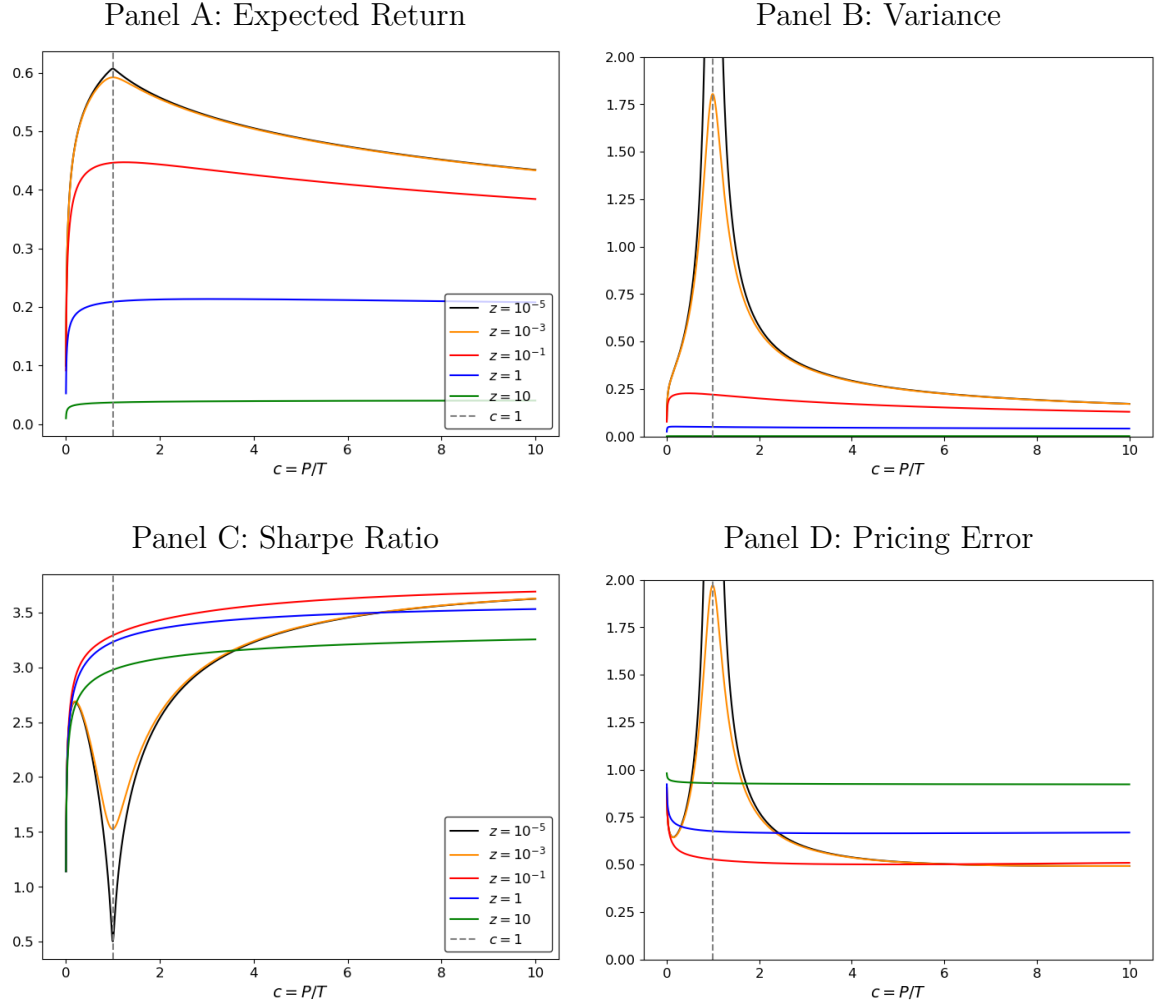


Figure 13: Out-of-sample Behavior of Large Factor Models in AIPT Calibration

Note. Limiting out-of-sample mean, variance, Sharpe ratio, and pricing error (HJD) of the SDF as a function of c and z from Theorem 3 with eigenvalues of $E[FF']$ given by $\eta = \{1/p^{0.5}\}_{p=1}^{P^*}$.

effective rank quantifies the number of dominant principal components among the factors. With $\alpha = 0.5$, the effective rank of $E[FF']$ in our AIPT calibration is roughly 500, capturing the spirit of the AIPT conjecture that the underlying DGP possesses a large number of distinct factors.

The remaining calibration choices do not vary across AIPT and APT calibrations. The

vector of mean returns behaves according to

$$E[F] \sim N\left(0, \kappa \frac{1}{P^*} \text{diag}(\eta^{1+\beta})\right).$$

The β parameter controls the degree of alignment between (squared) factor means and variances. We use $\beta = 1.5$, implying that higher variance factors have higher squared Sharpe ratios, as in the economic prior advocated by [KNS](#). The κ parameter controls the overall magnitude of the factor Sharpe ratios, and we set $\kappa = 400$ to match the magnitude of SDF Sharpe ratio in AIPT calibration with the empirical results in [Figure 2](#). In terms of dimensionality, our calibration assumes $T = 360$ training observations (as in our data analysis) and a true model complexity of P^*/T to 10. We use [Theorem 3](#) to draw theoretical VoC curves by gradually increasing the fraction of observable factors in the empirical model ($q = P/P^*$) from 0 to 1, thus varying empirical complexity $c = P/T$ from 0 to 10. These illustrations are the direct theoretical counterparts to the empirical VoC curves presented in [Section 3](#).

[Figure 13](#) plots the AIPT calibration results. Each curve corresponds to a different choice of ridge penalty z , and each point on a curve corresponds to a different number of factors P . The four panels show how complexity (on the x -axis) affects the expected out-of-sample mean, variance, and Sharpe of the SDF portfolio and the SDF's expected out-of-sample pricing errors.

Panel A shows that the out-of-sample expected return of the SDF portfolio is increasing in model complexity. The intuition for this result is that higher model complexity allows the SDF model to approximate the unknown true SDF more accurately. As the approximation improves and specification error shrinks, the estimated SDF is able to achieve a higher expected portfolio return. When the true DGP has many factors, the gains from better approximation overwhelm performance loss due to limits to learning. This is true for all ridge penalty levels and throughout the full range of complexity. Expected return curves are

flatter with higher explicit ridge shrinkage, z . This is because more shrinkage increases the estimator’s bias, reducing the approximating power of the SDF, which eats into its returns.

In Panel B, we see that SDF volatility is highly sensitive to model complexity. When c approaches unity, the variance of the low z SDFs spikes. The logic for this behavior follows from arguments in [KMZ](#). As $c \rightarrow 1$, the unregularized sample covariance matrix of factors becomes unstable, which causes the unregularized estimator to explode. Intuitively, when $c = 1$, the number of model parameters equals the number of time series observations, so there is an SDF estimate with an infinite in-sample Sharpe ratio. Without regularization, this estimator badly overfits the training data and produces disastrous out-of-sample behavior.

When $c \gg 1$, the ridge SDF estimator experiences implicit shrinkage due to complexity, which reduces SDF volatility. As the other curves in Panel B show, SDF volatility can also be controlled by raising the explicit ridge shrinkage, z . This is the low variance benefit of the shrinkage-induced bias.

The out-of-sample SDF Sharpe ratio is shown in Panel C. The ridgeless SDF estimator demonstrates “double ascent,” in analogy to the “double descent” MSE phenomenon studied in statistics literature.⁴¹ At low complexity ($c \ll 1$), the Sharpe ratio rises with complexity as larger models show improved approximation power benefits. But near $c = 1$, the Sharpe ratio collapses to zero due to the explosion in SDF variance. Finally, at high complexity ($c \gg 1$), variance comes under control, and the benefits of improved approximation again dominate and lead to an increasing Sharpe ratio. Panel C also demonstrates that, with appropriate explicit shrinkage z , the complex SDF estimator exhibits “permanent ascent” with an increasing Sharpe ratio throughout the full range of complexity.

Finally, Panel D illustrates the behavior of out-of-sample SDF pricing errors (HJD) as a function of complexity. Pricing errors in Panel D are a mirror image of the patterns for the Sharpe ratio in Panel C. The more complex the SDF, the better its ability to price assets out-of-sample. As long as the number of true factors (P) is large, these pricing errors

⁴¹See, for example, [Spigler et al. \(2019\)](#); [Belkin et al. \(2018, 2019, 2020\)](#); [Bartlett et al. \(2020\)](#).

never go to zero, even for very high-complexity empirical models. Higher complexity means the empirical SDF includes more true factors, improving its pricing ability. But at the same time, higher complexity also means more stringent limits to learning, so out-of-sample pricing errors are bounded away from zero (this is true even for high complexity models that are *correctly* specified).

The most important point regarding Figure 13 is that it closely matches the qualitative patterns in the data, as documented in Figure 2. The out-of-sample expected return and Sharpe ratio of the ridge SDF are gradually increasing in complexity while pricing errors are decreasing.

4.7.2 Calibration 2: The APT

To contrast with the AIPT in the previous calibration, we change the DGP assumption so that the cross-section of returns is dominated by only a few distinct factors, thus embodying the APT conjecture. We accomplish this by setting the eigenvalue distribution parameter α in (35) at $\alpha = 2.0$, which implies an effective rank of approximately 2.5. In this case, the cross-section of returns is dominated by only a few distinct factors, thus embodying the APT conjecture.

Figure 14 plots the APT calibration results. The main conclusion from this figure is that there are no benefits to complexity when the eigenvalue distribution is highly concentrated. When the ridge penalty is small (e.g., $z \leq 10^{-5}$), high-complexity models have lower Sharpe ratios and larger pricing errors than models with complexity near zero. The fact that peak SDF performance is reached with $c \ll 1$ is at odds with the patterns in Figure 2, where performance peaks at $c \gg 1$ regardless of the value of z .⁴²

⁴²Interestingly, while there are no gains to using large factor models in the APT calibration, there is surprisingly also *no cost* as long as moderate ridge shrinkage is used. This is because the overfit costs of adding additional (mostly redundant) parameters are offset by the implicit shrinkage that comes with complexity.

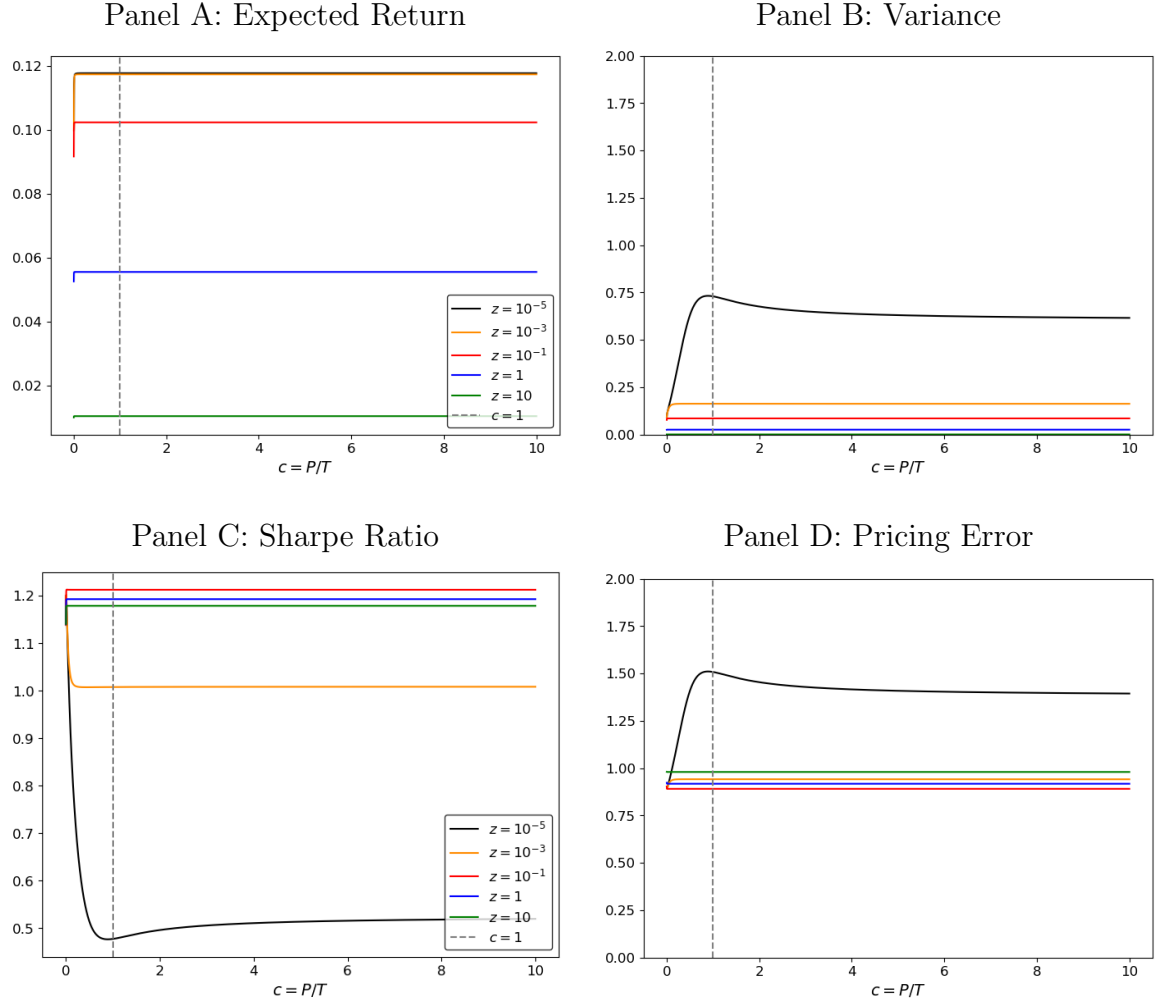


Figure 14: Out-of-sample Behavior of Large Factor Models in APT Calibration

Note. Limiting out-of-sample mean, variance, Sharpe ratio, and pricing error (HJD) of the SDF as a function of c and z from Theorem 3 with eigenvalues of $E[FF']$ given by $\eta = \{1/p^2\}_{p=1}^{P^*}$.

In summary, while the AIPT conjecture aligns closely with the evidence in Figure 2, the APT appears counterfactual.

4.7.3 Covariance Matrices For Discriminating Between APT and AIPT

The preceding calibrations compare the APT and AIPT to real-world data in terms of empirical asset pricing properties such as the Sharpe ratio. We can also compare APT and AIPT based on their predictions for the sample covariance matrix of factors.

Both DGPs assume a particular structure for the eigenvalue distribution of the true factor covariance matrix—either concentrated in a few dominant components in the case of the APT or distributed across a large number of components per the AIPT. One important consideration is that, in high-complexity environments, the sample covariance of factors will deviate markedly from the true covariance. This means we should not evaluate models based on how closely a model’s *true* eigenvalue distribution aligns with empirical eigenvalues. Instead, we must compare the *expected* sample eigenvalue distribution under each model assumption to the empirical distribution. Fortunately, we can use random matrix theory to make precise predictions for the expected sample eigenvalue distributions implied by the APT and AIPT models, and these predictions can then be compared directly to the empirical eigenvalue distribution.

The results of this comparison are shown in Figure 15. The calibrations are the same as those used in the preceding sections. We show the true eigenvalue distributions in our AIPT (dotted red line) and APT (dotted blue line) calibrations. Next, the solid red and blue lines show the expected sample eigenvalue distributions implied by RMT for each model for the calibrated high complexity sample (i.e., with $T = 360$ observations and $P = 3,600$ factors). In high-complexity settings, sample covariances are unavoidably polluted by noise. The RMT curves demonstrate how this noise obscures the true eigenvalue distribution and causes the sample eigenvalues from the two models to move toward one another.

Finally, the solid black line shows the realized empirical eigenvalue distribution in our empirical model with the same number of factors and observations. The main conclusion

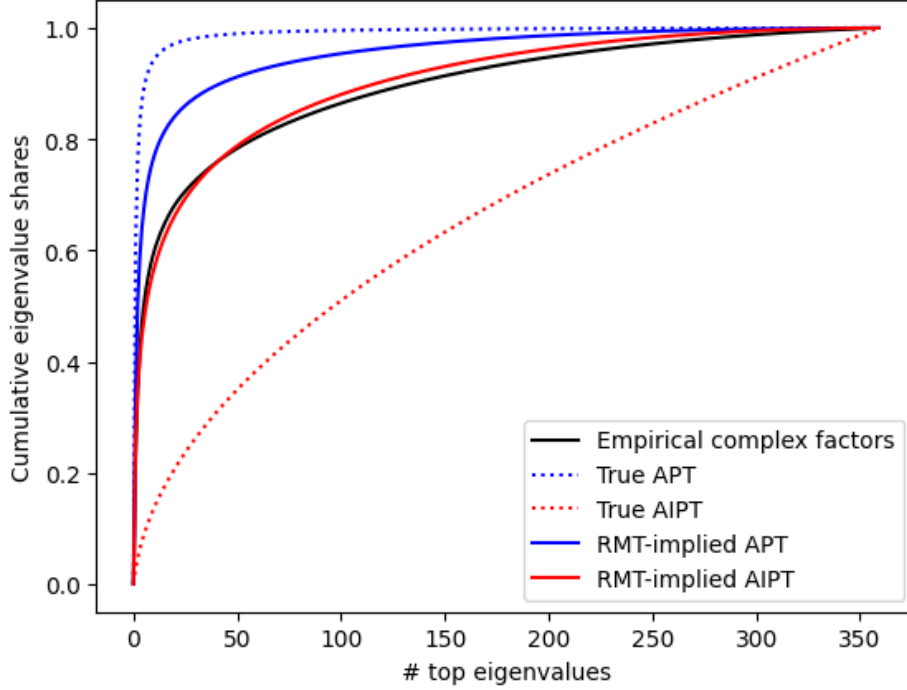


Figure 15: APT, AIPT, and Empirical Eigenvalue Distributions

Note. The figure shows eigenvalue distributions for the factor covariance matrix in our calibrations and in the data. The dotted red and blue lines show the true eigenvalue distributions in our AIPT ($\alpha = 0.5$) and APT ($\alpha = 2.0$) calibrations, respectively. The solid red and blue lines show the corresponding expected sample eigenvalue distributions for each model implied by RMT for the calibrated high complexity sample (i.e., with $T = 360$ observations and $P = 3,600$ factors). The solid black line shows the sample eigenvalue distribution in our empirical analysis with $T = 360$ and $P = 3,600$.

from this plot is that the AIPT provides a better match to the data not just in terms of Sharpe ratios and pricing errors but also in terms of the covariance structure among factors.

5 Conclusion

In this paper, we show that the out-of-sample performance of factor pricing models is increasing in terms of the number of factors. Larger models achieve higher risk-adjusted returns and lower pricing errors than smaller models, including standard low-dimensional factor models in the literature. Our novel machine learning empirical design allows us to compare models with varying degrees of statistical complexity while holding the raw data

inputs fixed. More heavily parameterized models outperform simpler models because they afford a better approximation to the unknown data-generating process. The cost of these parameters is higher out-of-sample model volatility, but this cost is more than justified by the gains in expected returns from using rich parameterizations.

We then develop a theory of machine learning SDF estimators founded on the concept of statistical model complexity. Among our key theoretical findings is the virtue of asset pricing model complexity: In essence, the out-of-sample performance of factor pricing models generally improves with the number of factors. This result holds as long as the covariance matrix of factors is not too concentrated. We also characterize the limits to learning that arise from the application of highly parameterized prediction models amid relative data scarcity. While heavy parameterization precludes consistent estimation of the SDF, the virtue of complexity arises from the improved approximation power of complex models, overwhelming the countervailing effect of limits to learning. The empirical patterns that we document bear a strikingly close resemblance to the predictions of our theory. Meanwhile, the theoretical predictions of the APT are contradicted by our empirical evidence.

References

- Asness, Clifford S, Tobias J Moskowitz, and Lasse Heje Pedersen, “Value and momentum everywhere,” *The Journal of Finance*, 2013, *68* (3), 929–985.
- Avramov, Doron, Si Cheng, and Lior Metzker, “Machine learning vs. economic restrictions: Evidence from stock return predictability,” *Management Science*, 2023, *69* (5), 2587–2619.
- Backus, David, Bryan Routledge, and Stanley Zin, “Asset prices in business cycle analysis,” *Unpublished manuscript, New York University*, 2007.
- Bai, Zhi-Dong and Jack W Silverstein, “No eigenvalues outside the support of the limiting spectral distribution of large-dimensional sample covariance matrices,” *The Annals of Probability*, 1998, *26* (1), 316–345.
- Bai, Zhidong and Wang Zhou, “Large sample covariance matrices without independence structures in columns,” *Statistica Sinica*, 2008, pp. 425–442.
- Bansal, Ravi and Amir Yaron, “Risks for the long run: A potential resolution of asset pricing puzzles,” *The journal of Finance*, 2004, *59* (4), 1481–1509.
- Barillas, Francisco and Jay Shanken, “Comparing asset pricing models,” *The Journal of Finance*, 2018, *73* (2), 715–754.
- Bartlett, Peter L, Philip M Long, Gábor Lugosi, and Alexander Tsigler, “Benign overfitting in linear regression,” *Proceedings of the National Academy of Sciences*, 2020, *117* (48), 30063–30070.
- Belkin, M, D Hsu, S Ma, and S Mandal, “Reconciling modern machine learning and the bias-variance trade-off. arXiv e-prints,” 2018.
- Belkin, Mikhail, Alexander Rakhlin, and Alexandre B Tsybakov, “Does data interpolation contradict statistical optimality?,” in “The 22nd International Conference on Artificial Intelligence and Statistics” PMLR 2019, pp. 1611–1619.

- , **Daniel Hsu**, and **Ji Xu**, “Two models of double descent for weak features,” *SIAM Journal on Mathematics of Data Science*, 2020, 2 (4), 1167–1180.
- Box, George EP and Gwilym Jenkins**, *Time Series Analysis: Forecasting and Control*, San Francisco: Holden-Day, 1970.
- Bryzgalova, Svetlana, Jiantao Huang, and Christian Julliard**, “Consumption in asset returns,” *The Journal of Finance*, Forthcoming.
- , **Markus Pelger**, and **Jason Zhu**, “Forest through the trees: Building cross-sections of stock returns,” *Available at SSRN 3493458*, 2020.
- Campbell, John Y, Stefano Giglio, Christopher Polk, and Robert Turley**, “An intertemporal CAPM with stochastic volatility,” *Journal of Financial Economics*, 2018, 128 (2), 207–233.
- Chen, Andrew Y and Tom Zimmermann**, “Open source cross-sectional asset pricing,” *Critical Finance Review*, Forthcoming, 2021.
- Chen, Luyang, Markus Pelger, and Jason Zhu**, “Deep learning in asset pricing,” *Management Science*, 2023.
- Chen, Xiaohong and Sydney C Ludvigson**, “Land of addicts? an empirical investigation of habit-based asset pricing models,” *Journal of Applied Econometrics*, 2009, 24 (7), 1057–1093.
- Cheng, Xiuyuan and Amit Singer**, “The spectrum of random inner-product kernel matrices,” *Random Matrices: Theory and Applications*, 2013, 2 (04), 1350010.
- Chernov, Mikhail, Magnus Dahlquist, and Lars A Lochstoer**, “Unpriced Risks: Rethinking Cross-Sectional Asset Pricing,” Technical Report, National Bureau of Economic Research 2025.
- , —, and **Lars Lochstoer**, “Pricing currency risks,” *The Journal of Finance*, 2023, 78 (2), 693–730.
- Chinco, Alex, Adam D Clark-Joseph, and Mao Ye**, “Sparse signals in the cross-section of returns,” *The Journal of Finance*, 2019, 74 (1), 449–492.

- Christiano, Lawrence J, Roberto Motto, and Massimo Rostagno**, “Risk shocks,” *American Economic Review*, 2014, *104* (1), 27–65.
- Cong, Lin William, Guanhao Feng, Jingyu He, and Xin He**, “Growing the efficient frontier on panel trees,” *NBER Working Paper*, 2022, (w30805).
- Connor, Gregory, Matthias Hagmann, and Oliver Linton**, “Efficient semiparametric estimation of the Fama–French model and extensions,” *Econometrica*, 2012, *80* (2), 713–754.
- Croce, Mariano M, Tatyana Marchuk, and Christian Schlag**, “The leading premium,” *The Review of Financial Studies*, 2023, *36* (8), 2997–3033.
- Da, Rui, Stefan Nagel, and Dacheng Xiu**, “The Statistical Limit of Arbitrage,” Technical Report, Technical Report, Chicago Booth 2022.
- Daniel, Kent, David Hirshleifer, and Lin Sun**, “Short-and long-horizon behavioral factors,” *The review of financial studies*, 2020, *33* (4), 1673–1736.
- Dew-Becker, Ian and Stefano Giglio**, “Asset pricing in the frequency domain: theory and empirics,” *The Review of Financial Studies*, 2016, *29* (8), 2029–2068.
- Do, Yen and Van Vu**, “The spectrum of random kernel matrices: universality results for rough and varying kernels,” *Random Matrices: Theory and Applications*, 2013, *2* (03), 1350005.
- Fama, Eugene F**, “Stock returns, expected returns, and real activity,” *The journal of finance*, 1990, *45* (4), 1089–1108.
- **and Kenneth R French**, “Common risk factors in the returns on stocks and bonds,” *Journal of financial economics*, 1993, *33* (1), 3–56.
- **and —**, “A five-factor asset pricing model,” *Journal of financial economics*, 2015, *116* (1), 1–22.
- Fan, Jianqing, Yuan Liao, and Weichen Wang**, “Projected principal component analysis in factor models,” *Annals of statistics*, 2016, *44* (1), 219.

- , Zheng Tracy Ke, Yuan Liao, and Andreas Neuhierl, “Structural Deep Learning in Conditional Asset Pricing,” *Available at SSRN 4117882*, 2022.
- Feng, Guanhao, Stefano Giglio, and Dacheng Xiu, “Taming the factor zoo: A test of new factors,” *The Journal of Finance*, 2020, 75 (3), 1327–1370.
- Filipovic, Damir and Paul Schneider, “Fundamental properties of linear factor models,” *arXiv preprint arXiv:2409.02521*, 2024.
- Freyberger, Joachim, Andreas Neuhierl, and Michael Weber, “Dissecting characteristics nonparametrically,” *The Review of Financial Studies*, 2020, 33 (5), 2326–2377.
- Gagliardini, Patrick, Elisa Ossola, and Olivier Scaillet, “Time-varying risk premium in large cross-sectional equity data sets,” *Econometrica*, 2016, 84 (3), 985–1046.
- Giglio, Stefano and Dacheng Xiu, “Asset pricing with omitted factors,” *Journal of Political Economy*, 2021, 129 (7), 1947–1990.
- , Bryan Kelly, and Dacheng Xiu, “Factor models, machine learning, and asset pricing,” *Annual Review of Financial Economics*, 2022, 14, 337–368.
- Gilchrist, Simon and Egon Zakrajšek, “Monetary policy and credit supply shocks,” *IMF Economic Review*, 2011, 59 (2), 195–232.
- and —, “Credit spreads and business cycle fluctuations,” *American economic review*, 2012, 102 (4), 1692–1720.
- , Jae W Sim, and Egon Zakrajšek, “Uncertainty, financial frictions, and investment dynamics,” Technical Report, National Bureau of Economic Research 2014.
- , Vladimir Yankov, and Egon Zakrajšek, “Credit market shocks and economic fluctuations: Evidence from corporate bond and stock markets,” *Journal of monetary Economics*, 2009, 56 (4), 471–493.
- Gu, Shihao, Bryan Kelly, and Dacheng Xiu, “Autoencoder Asset Pricing Models,” *Journal of Econometrics*, 2020.

- , — , and — , “Empirical asset pricing via machine learning,” *The Review of Financial Studies*, 2020, *33* (5), 2223–2273.
- Guijarro-Ordóñez, Jorge, Markus Pelger, and Greg Zanolli**, “Deep learning statistical arbitrage,” *arXiv preprint arXiv:2106.04028*, 2021.
- Han, Yufeng, Ai He, David Rapach, and Guofu Zhou**, “Expected stock returns and firm characteristics: E-LASSO, assessment, and implications,” *SSRN*, 2019.
- Hansen, Lars Peter and Kenneth J Singleton**, “Generalized instrumental variables estimation of nonlinear rational expectations models,” *Econometrica: Journal of the Econometric Society*, 1982, pp. 1269–1286.
- and **Ravi Jagannathan**, “Implications of security market data for models of dynamic economies,” *Journal of political economy*, 1991, *99* (2), 225–262.
- and — , “Assessing specification errors in stochastic discount factor models,” *The Journal of Finance*, 1997, *52* (2), 557–590.
- and **Scott F Richard**, “The role of conditioning information in deducing testable restrictions implied by dynamic asset pricing models,” *Econometrica: Journal of the Econometric Society*, 1987, pp. 587–613.
- Harvey, Campbell R, Yan Liu, and Heqing Zhu**, “... and the cross-section of expected returns,” *The Review of Financial Studies*, 2016, *29* (1), 5–68.
- Hastie, Trevor, Andrea Montanari, Saharon Rosset, and Ryan J Tibshirani**, “Surprises in high-dimensional ridgeless least squares interpolation,” *arXiv preprint arXiv:1903.08560*, 2019.
- He, Ai, Dashan Huang, Jiaen Li, and Guofu Zhou**, “Shrinking factor dimension: A reduced-rank approach,” *Management science*, 2023, *69* (9), 5501–5522.
- Hornik, Kurt, Maxwell Stinchcombe, and Halbert White**, “Multilayer feedforward networks are universal approximators,” *Neural networks*, 1989, *2* (5), 359–366.
- Hou, Kewei, Chen Xue, and Lu Zhang**, “Digesting anomalies: An investment approach,” *The Review of Financial Studies*, 2015, *28* (3), 650–705.

- , — , and — , “Replicating anomalies,” *The Review of Financial Studies*, 2020, *33* (5), 2019–2133.
- , **Haitao Mo**, **Chen Xue**, and **Lu Zhang**, “An augmented q-factor model with expected growth,” *Review of Finance*, 2021, *25* (1), 1–41.
- Jensen, Theis Ingerslev, Bryan Kelly, and Lasse Heje Pedersen**, “Is there a replication crisis in finance?,” *The Journal of Finance*, 2023, *78* (5), 2465–2518.
- Jordà, Òscar**, “Estimation and inference of impulse responses by local projections,” *American economic review*, 2005, *95* (1), 161–182.
- and **Alan M Taylor**, “Local projections,” *Journal of Economic Literature*, 2025, *63* (1), 59–110.
- Jurado, Kyle, Sydney C Ludvigson, and Serena Ng**, “Measuring uncertainty,” *American Economic Review*, 2015, *105* (3), 1177–1216.
- Kan, Raymond and Cesare Robotti**, “Model comparison using the Hansen-Jagannathan distance,” *The Review of Financial Studies*, 2009, *22* (9), 3449–3490.
- Karoui, Nouredine El**, “The spectrum of kernel random matrices,” *The Annals of Statistics*, 2010, *38* (1), 1–50.
- Kelly, Bryan and Dacheng Xiu**, “Financial machine learning,” *Foundations and Trends® in Finance*, 2023, *13* (3-4), 205–363.
- , **Boris Kuznetsov, Semyon Malamud, and Teng Andrea Xu**, “Large (and Deep) Factor Models,” *arXiv preprint arXiv:2402.06635*, 2024.
- , **Semyon Malamud, and Kangying Zhou**, “The virtue of complexity in return prediction,” *The Journal of Finance*, 2024, *79* (1), 459–503.
- , **Seth Pruitt, and Yinan Su**, “Characteristics are Covariances: A Unified Model of Risk and Return,” *Journal of Financial Economics*, 2020.
- Kelly, Bryan T and Semyon Malamud**, “Understanding The Virtue of Complexity,” *Available at SSRN 5346842*, 2025.

- , — , and **Kangying Zhou**, “The virtue of complexity everywhere,” *Available at SSRN*, 2022.
- Kozak, Serhiy**, “Kernel trick for the cross-section,” *Available at SSRN 3307895*, 2020.
- and **Stefan Nagel**, “When do cross-sectional asset pricing factors span the stochastic discount factor?,” Technical Report, National Bureau of Economic Research 2023.
- , — , and **Shrihari Santosh**, “Interpreting Factor Models,” *The Journal of Finance*, 2018, *73* (3), 1183–1223.
- , — , and — , “Shrinking the cross-section,” *Journal of Financial Economics*, 2020, *135* (2), 271–292.
- Lettau, Martin and Markus Pelger**, “Factors that fit the time series and cross-section of stock returns,” *The Review of Financial Studies*, 2020, *33* (5), 2274–2325.
- Liu, Yukun and Ben Matthies**, “Long-Run Risk: Is It There?,” *The Journal of Finance*, 2022, *77* (3), 1587–1633.
- and — , “Is Long Run Risk Really Priced? A Reply,” *A Reply (April 22, 2024)*, 2024.
- Maior, Paulo**, “Is Long-Run Risk Really Priced? Revisiting Liu and Matthies (2022),” *The Journal of Finance*, 2024, *79* (4), 2885–2900.
- Marčenko, Vladimir A and Leonid Andreevich Pastur**, “Distribution of eigenvalues for some sets of random matrices,” *Mathematics of the USSR-Sbornik*, 1967, *1* (4), 457.
- Martin, Ian WR and Stefan Nagel**, “Market efficiency in the age of big data,” *Journal of Financial Economics*, 2021.
- McLean, R David and Jeffrey Pontiff**, “Does academic research destroy stock return predictability?,” *The Journal of Finance*, 2016, *71* (1), 5–32.
- Parker, Jonathan A and Christian Julliard**, “Consumption risk and the cross section of expected returns,” *Journal of Political Economy*, 2005, *113* (1), 185–222.
- Preite, Massimo Dello, Raman Uppal, Paolo Zaffaroni, and Irina Zviadadze**, “What is Missing in Asset-Pricing Factor Models?,” 2022.

- Rahimi, Ali and Benjamin Recht**, “Random Features for Large-Scale Kernel Machines.,” in “NIPS,” Vol. 3 Citeseer 2007, p. 5.
- Rapach, David E and Guofu Zhou**, “Time-series and cross-sectional stock return forecasting: New machine learning methods,” *Machine learning for asset management: New developments and financial applications*, 2020, pp. 1–33.
- Ross, Stephen A.**, “The Arbitrage Theory of Capital Asset Pricing,” *Journal of Economic Theory*, 1976, *13*, 341–360.
- Santos, Tano and Pietro Veronesi**, “Conditional betas,” 2004.
- Spigler, Stefano, Mario Geiger, Stéphane d’Ascoli, Levent Sagun, Giulio Biroli, and Matthieu Wyart**, “A jamming transition from under-to over-parametrization affects generalization in deep learning,” *Journal of Physics A: Mathematical and Theoretical*, 2019, *52* (47), 474001.
- Stambaugh, Robert F and Yu Yuan**, “Mispricing factors,” *The review of financial studies*, 2017, *30* (4), 1270–1315.
- Tukey, John W**, “Discussion, emphasizing the connection between analysis of variance and spectrum analysis,” *Technometrics*, 1961, *3* (2), 191–219.
- White, Halbert**, *Estimation, inference and specification analysis* number 22, Cambridge university press, 1996.

Internet Appendix

A Road Map

The Appendix is organized as follows:

- Section B provides a Neural Network interpretation of the AIPT environment.
- Section C reports some useful properties of the infeasible portfolio.
- Section D presents some auxiliary results that we use in all of our proofs.
- Section E provides a proof of Proposition 2: The fact that, in the high complexity limit, characteristic-managed portfolios span the SDF.
- Section F provides some preliminaries on Random Matrix Theory, paving the road for all of our subsequent theoretical analyses. It provides a roadmap for the subsequent appendices and explains how this requires developing non-trivial arguments accounting for the complex dependency structure among factors $F_{t+1} = S'_t R_{t+1}$ that make classic results (such as, e.g., those from Hastie et al. (2019)) inapplicable to our setting.
- Section G provides the proof of one of the most important and technically challenging results in our paper: Lemma 11. This Lemma is the key ingredient to the proof of all Random Matrix Theory results. It establishes the striking result that the asymptotic limit as $P, T, N \rightarrow \infty$ is independent of the rate as which $N \rightarrow \infty$. In fact, the rate of convergence to this limit is controlled solely by $\text{tr}(\Sigma^2)$.
- Section H develops mathematical techniques for computing RMT quantities arising in AIPT.
- Section I computes $\lim E[\hat{R}_{T+1}^M(z; q; c)]$, proving item i. of Theorem 3.

- Section J computes $\lim E[(\hat{R}_{T+1}^M(z; q; c))^2]$, proving item ii. of Theorem 3.
- Section K computing the limit of HJD distance, proving item iv. of Theorem 3.

B Neural Network Interpretation of the AIPT Environment

The structure of returns in Assumption 1 has a clear machine-learning interpretation. Imagine for a moment that returns on asset i are generated by a low-dimensional factor model like those common in economic theory,⁴³

$$R_{i,t+1} = \beta(X_{i,t})' G_{t+1} + u_{i,t+1}, \quad (36)$$

where $X_{i,t}$ is a vector of J conditioning variables that determines i 's conditional betas on a small number K of latent factors, G_{t+1} . If one has no knowledge of the specific functional form for the conditional beta function, one can use a machine learning model to approximate it. For example, a shallow neural network could replace the $K \times 1$ vector $\beta(X_{i,t})$ with the approximation

$$\beta(X_{i,t}) \approx \sum_{p=1}^P \xi_p S_{i,t,p} = \underbrace{\Xi}_{K \times P} \underbrace{S_{i,t}}_{P \times 1}, \quad (37)$$

where

$$S_{i,t} = A(\Omega X_{i,t}) = (A(\omega_p' X_{i,t}))_{p=1}^P. \quad (38)$$

The neural network model approximates the unknown beta function with a linear combination of “generated conditioning variables” denoted $S_{i,t,p}$. Specifically, each $S_{i,t,p}$ is a basis function that captures nonlinear predictive information in the raw conditioning variables

⁴³Santos and Veronesi (2004) is an example asset pricing theory that generates a conditional beta formulation along these lines.

$X_{i,t}$. To build the basis functions, the neural network first generates a $J \times P$ matrix $\Omega = (\omega_p)_{p=1}^P$ of weights with rows ω_p to combine the elements of $X_{i,t}$ into P different linear combinations of $\Omega X_{i,t} \in \mathbb{R}^P$. Next, these linear combinations are transformed by a nonlinear activation function $A(x)$, so that we end up with nonlinear features $S_{i,t} = A(\Omega X_{i,t}) \in \mathbb{R}^P$. Then, equation (37) collects the P basis terms into a weighted sum in order to approximate $\beta(X_{i,t})$. The $K \times 1$ vectors ξ_p determine how each nonlinear basis term best contributes to the approximation of each of the K betas. We can write this sum in a matrix form by collecting the basis terms into a $P \times 1$ vector $S_{i,t}$ and the weights into the $K \times P$ matrix Ξ . Universal approximation theory, such as [Hornik et al. \(1989\)](#), ensures that the formulation in (37) can accurately approximate the true conditional beta function under regularity conditions.⁴⁴

To tie this back to Assumption 1, we may stack assets' beta coefficients into an $N \times K$ matrix and substitute (37) into (36) to deliver

$$\begin{aligned} R_{t+1} &\approx S_t \tilde{F}_{t+1} + u_{t+1}, \text{ with} \\ \tilde{F}_{t+1} &= \Xi' G_{t+1}, \quad \nu = E[\tilde{F}_{t+1}] = \Xi' E[G_{t+1}]. \end{aligned} \tag{39}$$

The key point of this neural network example is that, while Assumption 1 treats the factor loadings S_t as known and potentially high-dimensional, we interpret it as a generic statistical specification that arises from machine learning approximations to an unknown (and likely low-dimensional) factor pricing model.

C Properties of the Infeasible Portfolio

Let $\lambda = E[FF']^{-1}E[F]$. By a direct calculation,⁴⁵

$$\lambda = E[FF']^{-1}E[F] = \frac{1}{1 + MaxSR^2} \text{Var}[F]^{-1}E[F], \tag{40}$$

⁴⁴The approximating structure in (37) is analyzed by [Gu et al. \(2020a\)](#) and is a semi-nonparametric extension of the IPCA model in [Kelly et al. \(2020\)](#).

⁴⁵See the Sherman-Morrison formula (98).

where $\text{Var}[F]$ is the covariance matrix of factors, and where we have defined

$$MaxSR^2 = E[F]' \text{Var}[F]^{-1} E[F] \quad (41)$$

to be the maximal achievable unconditional squared Sharpe ratio. Most existing papers perform their analysis assuming that the population moments of the factors are directly observable and, hence, so is the vector of factor risk premia, λ . The corresponding portfolio satisfies

$$E[\lambda' F_{t+1}] = E[(\lambda' F_{t+1})^2] = E[F]' E[FF']^{-1} E[F] = \frac{MaxSR^2}{1 + MaxSR^2}. \quad (42)$$

It will be instructive for our subsequent analysis to decompose the maximal Sharpe ratio into the contributions coming from the factor principal components. Given the eigenvalue decomposition $\text{Var}[F] = U \text{diag}(\mu) U'$, we can define PC_i to be the i -th column of $U'F$. In the sequel, we will use

$$\theta = U' E[F] \quad (43)$$

to denote the vector of mean returns of the PCs. Then, we can rewrite the maximal Sharpe ratio (41) as

$$MaxSR^2 = \sum_i \frac{\theta_i^2}{\mu_i} = \sum_i (SR(PC_i))^2. \quad (44)$$

We will now use this representation to understand the effect of ridge shrinkage on the performance of the *infeasible* efficient portfolio,

$$R_{t+1}^{infeas}(z) = E[F]'(zI + \text{Var}[F])^{-1} F_{t+1}. \quad (45)$$

We call this portfolio *infeasible* because, in the big data regime, when $P > T$, neither $E[F] \in \mathbb{R}^P$ nor $E[FF'] \in \mathbb{R}^{P \times P}$ can be efficiently estimated from only T observations. By construction, $R_{t+1}^{infeas}(0) = \lambda' F_{t+1}$ achieves the *MaxSR*, and

$$\mathcal{E}(z) = E[R^{infeas}(z)] = E[F]'(zI + E[FF'])^{-1}E[F] = \frac{A(z)}{1 + A(z)}, \quad (46)$$

where we have defined

$$\begin{aligned} A(z) &= E[F]'(zI + \text{Var}[F])^{-1}E[F] \\ &= \sum_i (SR(PC_i))^2 \frac{\mu_i}{\mu_i + z} \\ &= \sum_i (SR(PC_i))^2 \frac{1}{1 + z/\mu_i} \approx \sum_{i:\mu_i > z} (SR(PC_i))^2 \end{aligned} \quad (47)$$

and

$$A'(z) = - \sum_i \theta_i^2 \frac{1}{(\mu_i + z)^2}. \quad (48)$$

The function $A(z)$ will be important in understanding ridge-regularization in the high-complexity case. It turns out that the risk of the efficient portfolio can be expressed in terms of the derivative of $A(z)$: Defining

$$(zA(z))' = \sum_i (SR(PC_i))^2 \left(\frac{\mu_i}{\mu_i + z} \right)^2, \quad (49)$$

a somewhat tedious calculation implies that

$$\mathcal{V}(z) = \text{Var}[R^{infeas}(z)] = \frac{(zA(z))'}{(1 + A(z))^2}. \quad (50)$$

and

$$\begin{aligned}
& E[(R^{infeas}(z))^2] \\
&= \frac{1}{(1+A(z))^2} E[(E[F]'(zI + \Psi)^{-1}F_t)^2] = \frac{1}{(1+A(z))^2} E[E[F]'(zI + \Psi)^{-1}F_t F_t'(zI + \Psi)^{-1}E[F]] \\
&= \frac{1}{(1+A(z))^2} E[E[F]'(zI + \Psi)^{-1}F_t F_t'(zI + \Psi)^{-1}E[F]] \\
&= E[F]'(zI + \Psi)^{-1}\Psi(zI + \Psi)^{-1}E[F] + \mathcal{E}(z)^2 \\
&= \frac{1}{(1+A(z))^2} \sum_i \theta_i^2(z + \mu_i)^{-2}\mu_i + \left(\frac{A(z)}{1+A(z)} \right)^2 \\
&= \frac{A(z) + zA'(z) + A^2(z)}{(1+A(z))^2} \\
&= \frac{(A(z) + zA'(z))(1+A(z)) - zA(z)A'(z)}{(1+A(z))^2} \\
&= \frac{d}{dz} \left(\frac{zA(z)}{1+A(z)} \right).
\end{aligned} \tag{51}$$

Since the weights $\frac{\mu_i}{\mu_i + z}$ are monotone increasing in μ_i , we see that all that the ridge shrinkage does is re-weights principal components, giving a larger weight to higher-variance PCs. The following is a simple but important observation, implying that ridge shrinkage is always detrimental to performance.

Lemma 1 *The Sharpe ratio $SR^{infeasible}(z) = SR(R^{infeasible}(z))$ is monotone decreasing in z .*

D Auxilliary Results

D.1 Concentration of Quadratic Forms

Roadmap: A key technical result in many RMT proofs is the concentration of quadratic forms. Heuristically, it can be stated as follows: When $X_{i,t}$ are i.i.d. with mean zero and

unit variance, then $E[X_t X_t'] = I$ and, hence,

$$P^{-1} X_t' C X_t = P^{-1} \text{tr}(C X_t X_t') \underset{\substack{\approx \\ \text{law of large numbers as } P \rightarrow \infty}}{\approx} P^{-1} \text{tr}(C E[X_t X_t']) = P^{-1} \text{tr}(C) \quad (52)$$

for any bounded matrix C . The next Lemma formalizes this intuition and characterizes the speed of convergence of $X_t' C X_t$ to $\text{tr}(C)$. **This lemma is a key element to many of our proofs below.**

Lemma 2 (Lemma 2.7 in (Bai and Silverstein, 1998)) *Suppose that $X_{i,t}$ are i.i.d. with mean zero and unit variance. Let C a $P \times P$ matrix. Then, for any $p \geq 2$, there exists a constant $K_p > 0$ such that*

$$E[|X_t' C X_t - \text{tr}(C)|^p] \leq K_p ((E[|X_{1,1}|^4] \text{tr}(C C'))^{p/2} + E[|X_{1,1}|^{2p}] \text{tr}((C C')^{p/2})). \quad (53)$$

In particular,

$$E[|X_t' C X_t - \text{tr}(C)|^p] \leq K_p (P^{p/2} (E[|X_{1,1}|^4] \|C\|^2)^{p/2} + E[|X_{1,1}|^{2p}] P \|C\|^p) \leq \tilde{K} \|C\|^p P^{p/2}. \quad (54)$$

Interestingly enough, this concentration result can be established under weaker conditions than independence. We will need the following definition.

Definition 1 (Strongly uncorrelated variables) *We say that f_i , $i = 1, \dots, K$ are strongly uncorrelated if $E[f_{i_1}] = E[f_{i_1} f_{i_2}] = 0$ for all $i_1 \neq i_2$, $E[f_{i_1} f_{i_2} f_{i_3}] = 0$ for any i_1, i_2, i_3 and $E[f_{i_1} f_{i_2} f_{i_3} f_{i_4}] = 0$ unless the set $\{i_1, i_2, i_3, i_4\}$ contains exactly two different elements. Furthermore, $E[f_i^2 f_j^2] = E[f_i^2] E[f_j^2]$ for $i \neq j$.*

The following is true.

Lemma 3 [*Weak Concentration of Quadratic forms*] Suppose that $X = (X_i)_{i=1}^P$ with X_i being strongly uncorrelated according to Definition 1. Suppose also that $E[X_i^2] = 1, E[X_i^4] \leq k$, and let A_P be random matrices independent of X and such that $\|A_P\|_2 = o(1)$. Let also

$$Y_t = X_t' A_P X_t. \quad (55)$$

Then,

$$(1) Y_t = \text{tr}(A_P X_t X_t') \quad (56)$$

$$(2) \lim_{P \rightarrow \infty} E[(Y_t - \text{tr}(A_P))^2 | A_P] = 0 \quad (57)$$

In particular, If $A_P = B_P/P$ where $\|B_P\| \leq K$, we have $\|A_P\|_2^2 \leq P\|B_P\|^2/P^2$ and hence

$$\lim_{P \rightarrow \infty} E[(X_t' B_P X_t - \text{tr}(B_P))^2 | B_P] / P^2 = 0. \quad (58)$$

Proof of Lemma 3.

(1):

$$\begin{aligned} X_t' A X_t &\in R \Rightarrow X_t' A X_t = \text{tr}(X_t' A X_t) \\ \text{tr}(AB) &= \text{tr}(BA) \Rightarrow \text{tr}(X_t' A X_t) = \text{tr}(A X_t X_t') \end{aligned}$$

(2): Define $Y_t = X_t' A_P X_t$. We have

$$E[Y_t] = E[\text{tr}(A_P (X_t X_t')) | A_P] = \text{tr}(A_P E[X_t X_t']) = \text{tr}(A_P),$$

and hence

$$E[(Y_t - \text{tr}(A_P))^2 | A_P] = \text{Var}[Y_t | A_P] = E[Y_t^2 | A_P] - E[Y_t | A_P]^2 \quad (59)$$

and hence it suffices to prove that

$$E[Y_t^2 | A_P] - (\text{tr}(A_P))^2 \rightarrow 0 \quad (60)$$

For simplicity, we assume from now on that A_P is deterministic, and write $A_P = (A_{i,j})_{i,j=1}^P$.

We also assume that A is *symmetric*. Then,

$$Y_t = \sum_{i,j} X_i X_j A_{i,j} \quad (61)$$

and therefore

$$Y_t^2 = \sum_{i_1, j_1, i_2, j_2} X_{i_1} X_{j_1} A_{i_1, j_1} A_{i_2, j_2} X_{i_2} X_{j_2} \quad (62)$$

Now, among all fourth-order moments, $E[X_{i_1} X_{j_1} X_{i_2} X_{j_2}]$, the only non-zero moments are those where either all are identical, $i_1 = i_2 = i_3 = i_4$, or when there are exactly two identical pairs. The latter can happen in exactly 3 ways. First, $(i_1 = i_2, j_1 = j_2)$, $(i_1 = j_2, j_1 = i_2)$ give rise to the terms A_{i_1, j_1}^2 because, by assumption, A is symmetric, so that $A_{i_1, j_1} = A_{j_1, i_1}$. Second, $(i_1 = j_1, i_2 = j_2)$ gives rise to $A_{i,i} A_{j,j}$. Note also that $E[X_i^2 X_j^2] = E[X_i^2] E[X_j^2] = 1$.

Thus,

$$\begin{aligned}
E[Y_t^2] &= \sum_{i_1, j_1, i_2, j_2} A_{i_1, j_1} A_{i_2, j_2} E[X_{i_1} X_{j_1} X_{i_2} X_{j_2}] \\
&= \sum_i A_{i,i}^2 E[X_i^4] + \sum_{i,j, i \neq j} (2A_{i,j}^2 + A_{i,i} A_{j,j}) \\
&= \sum_i A_{i,i}^2 E[X_i^4] + \sum_{i,j, i \neq j} 2A_{i,j}^2 - \sum_i A_{i,i}^2 + \left(\sum_i A_{i,i}\right)^2 \\
&= \sum_i A_{i,i}^2 E[X_i^4] - 2 \sum_i A_{i,i}^2 + \sum_{i,j} 2A_{i,j}^2 - \sum_i A_{i,i}^2 + \left(\sum_i A_{i,i}\right)^2 \\
&= \sum_i A_{i,i}^2 (E[X_i^4] - 3) + 2\|A\|_2^2 + (\text{tr}(A))^2
\end{aligned} \tag{63}$$

Thus, since $E[Y_t] = \text{tr}(A)$, we have

$$E[Y_t^2] - \text{tr}(A)^2 = E[Y_t^2] - E[Y_t]^2 = \sum_i A_{i,i}^2 (E[X_i^4] - 3) + 2\|A\|_2^2 \leq (k-1)\|A\|_2^2. \tag{64}$$

because

$$\sum_i A_{i,i}^2 \leq \sum_{i,j} A_{i,j}^2 = \|A\|_2^2, \tag{65}$$

The proof is complete. $\square$

D.2 Convergence of Expectations

Roadmap: In many of the proofs below, we would like to use the following intuitive argument: If $Y_P - E[Y_P] \rightarrow 0$, then

$$E[Z_P Y_P] \approx E[Z_P E[Y_P]] = E[Z_P] E[Y_P].$$

The next two lemmas make this intuition rigorous.

Lemma 4 Suppose that $Y = Y_P$ is such that $Y_P - E[Y_P] \rightarrow 0$ in L_2 , and $Y_P \geq 0$. Then,

$$E[X/(1+Y)] - E[X]/(1+E[Y]) \rightarrow 0 \quad (66)$$

if $E[|X|^2]$ is uniformly bounded.

Proof of Lemma 4. We have

$$\begin{aligned} |E[X/(1+Y)] - E[X]/(1+E[Y])| &= |E[\frac{X(Y-E[Y])}{(1+Y)(1+E[Y])}]| \\ &\leq E[\frac{|X||Y-E[Y]|}{(1+Y)(1+E[Y])}] \leq E[|X||Y-E[Y]|] \\ &\leq E[X^2]^{1/2} \text{Var}[Y]^{1/2}. \end{aligned} \quad (67)$$

□

Lemma 5 Suppose that $Y(j) \geq 0, j = 1, \dots, J$, are such that $\text{Var}[Y(j)] \leq KP^{-1}$ and $E[Y_j]$ are bounded. Then,

$$|E[\frac{X}{(1+Y_1)\cdots(1+Y_J)}] - \frac{E[X]}{(1+E[Y_1])\cdots(1+E[Y_J])}| \leq K(J)P^{-1/2}E[X^2]^{1/2} \quad (68)$$

for some $K(J) > 0$.

Proof of Lemma 5. Let Y be such that

$$1+Y_J = (1+Y_1)\cdots(1+Y_J). \quad (69)$$

Let also

$$1+Z_J = (1+E[Y_1])\cdots(1+E[Y_J]). \quad (70)$$

Then,

$$|E[\frac{X}{(1+Y_1)\cdots(1+Y_J)}] - \frac{E[X]}{1+Z}| = |E[\frac{Y-Z}{(1+Y)(1+Z)}X]| \leq E[(Y-Z)^2/(1+Y)^2]^{1/2} E[X^2]^{1/2}. \quad (71)$$

We now prove by induction in J that $E[(Y-Z_J)^2/(1+Y_J)^2] \leq \kappa(J)P^{-1}$. We have

$$(1+Y_1)\cdots(1+Y_J) = (1+E[Y_1])(1+Y_2)\cdots(1+Y_J) + (Y_1-E[Y_1])(1+Y_2)\cdots(1+Y_J). \quad (72)$$

Thus, defining $Z_J = (1+E[Y_1])\cdots(1+E[Y_J]) - 1$, $Z_{J-1} = (1+E[Y_2])\cdots(1+E[Y_J]) - 1$, we get

$$Y_J - Z_J = (1+E[Y_1])(Y_{J-1} - Z_{J-1}) + (Y_1 - E[Y_1])Y_{J-1}. \quad (73)$$

Hence, since $1 \leq Y_{J-1} \leq Y_J$, we get

$$\begin{aligned} E[(Y_J - Z_J)^2/(1+Y_J)^2] &\leq 2 E[(1+E[Y_1])^2(Y_{J-1} - Z_{J-1})^2/(1+Y_J)^2] \\ &+ E[(Y_1 - E[Y_1])^2 Y_{J-1}^2/(1+Y_J)^2] \\ &\leq 2(1+E[Y_1])^2 \kappa(J-1)P^{-1} + E[(Y_1 - E[Y_1])^2] \leq (2(1+E[Y_1])^2 \kappa(J-1) + K)P^{-1} \end{aligned} \quad (74)$$

and thus we have proved the induction step.

□

D.3 Computing Moments

Roadmap: In this section, we derive analytical expressions for high-order moments of many quantities appearing in our calculations.

We will need the following lemma, whose proof follows by direct calculation.

Lemma 6 *Suppose that $X_t \in \mathbb{R}^{N \times P}$ is a matrix with i.i.d. elements satisfying $E[X_{i,k}X_{j,l}] = \delta_{(i,k),(j,l)}$. Then,*

$$E[X_t' \Sigma X_t] = \text{tr}(\Sigma) I_{P \times P}.$$

Lemma 7 *Let A be a symmetric matrix. Then, we have*

$$\begin{aligned} E[S_t' S_t A S_t' S_t] &= ((\text{tr} \Sigma)^2 + \text{tr}(\Sigma^2)) \Psi A \Psi + \text{tr}(\Sigma^2) \text{tr}(\Psi A) \Psi \\ &+ \text{tr}(\Sigma \circ \Sigma) \Psi^{1/2} (\kappa - 3) \text{diag}(\Psi^{1/2} A \Psi^{1/2}) \Psi^{1/2} \end{aligned} \quad (75)$$

where $\text{diag}(\Psi^{1/2} A \Psi^{1/2})$ is the diagonal matrix with diagonal coinciding with that of $\text{diag}(\Psi^{1/2} A \Psi^{1/2})$, and $\Sigma \circ \Sigma$ is the elementwise product of Σ with itself.

Proof. Substituting $S_t = \Sigma^{1/2} X_t \Psi^{1/2}$, we get

$$E[S_t' S_t A S_t' S_t] = E[S_t' S_t A S_t' S_t] = \Psi^{1/2} E[X_t' \Sigma X_t \tilde{A} X_t' \Sigma X_t] \Psi^{1/2}, \quad (76)$$

where $\tilde{A} = \Psi^{1/2} A \Psi^{1/2}$. Thus, it suffices to consider the case $\Psi = I$ and just compute

$$E[X_t' \Sigma X_t \tilde{A} X_t' \Sigma X_t], \quad (77)$$

in which case the desired identity takes the form

$$\begin{aligned} E[X_t' \Sigma X_t A X_t' \Sigma X_t] &= ((\text{tr} \Sigma)^2 + \text{tr}(\Sigma^2)) A + \text{tr}(\Sigma^2) \text{tr}(A) I \\ &+ \text{tr}(\Sigma \circ \Sigma) (\kappa - 3) \text{diag}(A). \end{aligned} \quad (78)$$

Further, it suffices to prove the result for rank-one matrices because any symmetric matrix

A can be written as

$$A = \sum_i \lambda_i \beta_i \beta'_i$$

Thus, suppose that $A = \beta \beta'$. Then,

$$\begin{aligned} & E[X' \Sigma X \beta \beta' X' \Sigma X]_{j,k} \\ &= E\left[\sum_{i_1, i_2, i_3, i_4, i_5, i_6} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, k} \right]. \end{aligned} \quad (79)$$

By assumption, all elements of X are i.i.d. and have mean zero. Thus, the only non-zero terms come from two identical pairs or when all terms are identical. For two identical pairs, they are determined by what coincides with (i_1, j) . These are the possibilities: $(i_1, j) = (i_2, i_3)$ or $(i_1, j) = (i_5, i_4)$ or $(i_1, j) = (i_6, k)$. The latter can only happen when $j = k$.

- Suppose first that $j \neq k$. Then,

$$\begin{aligned} & E\left[\sum_{i_1, i_2, i_3, i_4, i_5, i_6} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, k} \right] \\ &= E\left[\sum_{i_1=i_2, j=i_3, i_5=i_6, i_4=k} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, k} \right] \\ &+ E\left[\sum_{i_1=i_5, j=i_4, i_2=i_6, i_3=k} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, k} \right] \\ &= (\text{tr}(\Sigma)^2 + \text{tr}(\Sigma^2)) \beta_j \beta_k = (\text{tr}(\Sigma)^2 + \text{tr}(\Sigma^2)) A_{j,k}. \end{aligned} \quad (80)$$

- When $j = k$, we get

$$\begin{aligned}
& E\left[\sum_{i_1, i_2, i_3, i_4, i_5, i_6} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, j}\right] \\
&= E\left[\sum_{i_1, i_2, i_3, i_4, i_5} X_{i_1, j}^2 \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_1}\right] \\
&+ E\left[\sum_{i_1, i_2, i_3, i_4, i_5, i_6 \neq i_1} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, j}\right] \\
&= E\left[\sum_{i_1, i_2, i_3} X_{i_1, j}^2 \Sigma_{i_1, i_2} X_{i_2, i_3}^2 \beta_{i_3}^2 \Sigma_{i_2, i_1}\right] \\
&+ E\left[\sum_{i_1, i_2, i_3, i_4, i_5, i_6 \neq i_1} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, j}\right] \\
&= \text{tr}(\Sigma^2) \sum_i \beta_i^2 + (\kappa - 1) \sum_i \Sigma_{i, i}^2 \beta_j^2 \\
&+ E\left[\sum_{i_1, i_2, i_3, i_4, i_5, i_6 \neq i_1} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, j}\right] \\
&= \text{tr}(\Sigma^2) \sum_i \beta_i^2 + (\kappa - 1) \sum_i \Sigma_{i, i}^2 \beta_j^2 \\
&+ E\left[\sum_{i_1=i_2, i_3=j, i_4=j, i_5=i_6 \neq i_1} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, j}\right] \\
&+ E\left[\sum_{i_1=i_5, i_3=j, i_4=j, i_2=i_6 \neq i_1} X_{i_1, j} \Sigma_{i_1, i_2} X_{i_2, i_3} \beta_{i_3} \beta_{i_4} X_{i_5, i_4} \Sigma_{i_5, i_6} X_{i_6, j}\right] \\
&= \text{tr}(\Sigma^2) \sum_i \beta_i^2 + (\kappa - 1) \sum_i \Sigma_{i, i}^2 \beta_j^2 \\
&+ \beta_j^2 \sum_{i_1} \Sigma_{i_1, i_1} (\text{tr}(\Sigma) - \Sigma_{i_1, i_1}) \\
&+ \beta_j^2 \sum_{i_1} \sum_{i_2 \neq i_1} \Sigma_{i_1, i_2}^2 \\
&= \text{tr}(\Sigma^2) \sum_i \beta_i^2 + (\kappa - 3) \sum_i \Sigma_{i, i}^2 \beta_j^2 + \beta_j^2 (\text{tr}(\Sigma))^2 \\
&+ \beta_j^2 \sum_{i_1} \left(\sum_{i_2} \Sigma_{i_1, i_2}^2 - \Sigma_{i_1, i_1}^2 \right) \\
&= \text{tr}(\Sigma^2) \sum_i \beta_i^2 + (\kappa - 3) \sum_i \Sigma_{i, i}^2 \beta_j^2 + \beta_j^2 ((\text{tr}(\Sigma))^2 + \text{tr}(\Sigma^2))
\end{aligned} \tag{81}$$

Since $j = k$, the latter can be rewritten as

$$\text{tr}(\Sigma^2) \text{tr}(A) \delta_{j,k} + (\kappa - 3) \text{tr}(\Sigma \circ \Sigma) A_{j,k} + ((\text{tr}(\Sigma))^2 + \text{tr}(\Sigma^2)) A_{j,k}. \quad (82)$$

□

Lemma 8 *We have*

$$E[\varepsilon Z' \varepsilon] = E[\Sigma_{\varepsilon,t} Z]$$

and

$$E[\varepsilon' Z \varepsilon'] = E[Z' \Sigma_{\varepsilon,t}]$$

for any random vector Z independent of $\varepsilon, \Sigma_{\varepsilon}$. Furthermore,

$$E[\varepsilon' A \varepsilon] = \text{tr}(A \Sigma_{\varepsilon,t})$$

for any matrix A . Furthermore, for any matrix B independent of $\varepsilon, \Sigma_{\varepsilon}$, we have

$$E[(\varepsilon' B \varepsilon)^2] = \sum_i \tilde{B}_{i,i}^2 (\kappa_{\varepsilon} - 3) + 2 \text{tr}(\tilde{B}^2) + (\text{tr}(\tilde{B}))^2 \quad (83)$$

where $\tilde{B} = 0.5 \Sigma_{\varepsilon,t}^{1/2} (B + B') \Sigma_{\varepsilon,t}^{1/2}$, and

$$E[\varepsilon_t \varepsilon'_t B \varepsilon_t \varepsilon'_t] = \Sigma_{\varepsilon,t}^{1/2} (2\tilde{B} + \text{diag}((\kappa_{\varepsilon} - 3) \text{diag}(\tilde{B}) + \text{tr}(\tilde{B}))) \Sigma_{\varepsilon,t}^{1/2}.$$

Proof. We have

$$E[\varepsilon Z' \varepsilon]_{i,j} = E[\varepsilon_i \sum_j Z_j \varepsilon_j] = \sum_j E[\Sigma_{\varepsilon,i,j} Z_j]$$

and the first claim follows. The second claim follows because

$$E[\varepsilon' Z \varepsilon'] = E[\varepsilon Z' \varepsilon]'.$$

For the third claim, we have

$$E[\varepsilon' A \varepsilon] = \text{tr } E[\varepsilon' A \varepsilon] = \text{tr } E[A \varepsilon \varepsilon'] = E[\text{tr}(A \Sigma_{\varepsilon, t})], \quad (84)$$

while (83) follows from (63). For the last claim, we make the observation that, for any matrix B ,

$$\varepsilon' B \varepsilon = 0.5 \varepsilon' (B + B') \varepsilon.$$

Identity (83) follows (63).

For brevity, we omit the t subindex below. Next,

$$E[\varepsilon \varepsilon' B \varepsilon \varepsilon'] = E[\Sigma_{\varepsilon}^{1/2} \tilde{\varepsilon} \tilde{\varepsilon}' \Sigma_{\varepsilon}^{1/2} B \Sigma_{\varepsilon}^{1/2} \tilde{\varepsilon} \tilde{\varepsilon}' \Sigma_{\varepsilon}^{1/2}]. \quad (85)$$

Define

$$\hat{B} = \Sigma_{\varepsilon}^{1/2} B \Sigma_{\varepsilon}^{1/2}. \quad (86)$$

For $j \neq k$, we have

$$\begin{aligned} & E[\tilde{\varepsilon} \tilde{\varepsilon}' \tilde{B} \tilde{\varepsilon} \tilde{\varepsilon}']_{j,k} \\ &= E[\tilde{\varepsilon}_j \sum_{i_1, i_2} \tilde{\varepsilon}_{i_1} \tilde{\varepsilon}_{i_2} \hat{B}_{i_1, i_2} \tilde{\varepsilon}_k] = \hat{B}_{j,k} + \hat{B}_{k,j}, \end{aligned} \quad (87)$$

while, for $j = k$, we have

$$\begin{aligned}
& E[\tilde{\varepsilon}\tilde{\varepsilon}'\hat{B}\tilde{\varepsilon}\tilde{\varepsilon}']_{j,j} \\
&= E[\tilde{\varepsilon}_j^2 \sum_{i_1, i_2} \tilde{\varepsilon}_{i_1} \tilde{\varepsilon}_{i_2} \hat{B}_{i_1, i_2}] \\
&= \kappa_\varepsilon \hat{B}_{j,j} + \sum_{i \neq j} \hat{B}_{i,i} = (\kappa_\varepsilon - 1) \hat{B}_{j,j} + \text{tr}(\hat{B}).
\end{aligned} \tag{88}$$

□

D.4 Moments of Managed Portfolios

Recall that

$$F_{t+1} = S'_t R_{t+1} \tag{89}$$

Lemma 9 (Expected Factor Moments) *Let $\Sigma_{F,t}^* = \Sigma_{F,t} + \nu\nu'$. Suppose a normalization $\text{tr}(\Sigma) = 1$ and let $\sigma_* = \text{tr}(\Sigma\Sigma_\varepsilon)$ and $E[X_{i,k}^4] = \kappa$ for all i, k . Suppose also that $E[X_{i,k}^3] = 0$. We have*

$$E[S'_t \Sigma_\varepsilon S_t] = \text{tr}(\Sigma\Sigma_\varepsilon)\Psi$$

and

$$\begin{aligned}
E[F_{t+1}F'_{t+1}] &= ((\text{tr} \Sigma)^2 + \text{tr}(\Sigma^2))\Psi E[\Sigma_{F,t}^*]\Psi \\
&+ \text{tr}(\Sigma \circ \Sigma)\Psi^{1/2}(\kappa - 3) \text{diag}(\Psi^{1/2}E[\Sigma_{F,t}^*]\Psi^{1/2})\Psi^{1/2} + \Psi \left(\text{tr}(\Sigma\Sigma_\varepsilon) + \text{tr}(\Psi E[\Sigma_{F,t}^*]) \text{tr}(\Sigma^2) \right)
\end{aligned} \tag{90}$$

Thus,

$$\|E[F_{t+1}F'_{t+1}] - (\Psi E[\Sigma_{F,t}^*]\Psi + \sigma_*\Psi)\| \rightarrow 0 \quad (91)$$

when $P \rightarrow \infty$.

Proof of Lemma 9. We abuse the notation and use Σ_F to denote $E[\Sigma_{F,t}^*]$.

Recall that, under the normalization $\text{tr}(\Sigma) = 1$, by Assumption 2, $\text{tr}(\Sigma^2) \rightarrow 0$ and $\text{tr}(\Sigma \circ \Sigma) = \sum_i \Sigma_{i,i}^2 \leq \sum_{i,j} \Sigma_{i,j}^2 = \text{tr}(\Sigma^2) \rightarrow 0$. Let $\Sigma_{F,\varepsilon,t} = E_t[\tilde{F}\varepsilon']$.

$$\begin{aligned} E[F_{t+1}F'_{t+1}] &= E[S'_t(S_t\tilde{F} + \varepsilon)(S_t\tilde{F} + \varepsilon)'S_t] \\ &= E[S'_tS_t\Sigma_{F,\varepsilon,t}S_t] + E[S'_t\Sigma'_{F,\varepsilon,t}S'_tS_t] + E[S'_tS_t\Sigma_F S'_tS_t] + E[S'_t\Sigma_\varepsilon S_t], \end{aligned} \quad (92)$$

Then, the cubic terms all vanish because, by assumption, $E[X_{i,j,t}^3] = 0$, and

$$E[S'_t\Sigma_\varepsilon S_t] = E[\Psi^{1/2}X'_t\Sigma^{1/2}\Sigma_\varepsilon\Sigma^{1/2}X_t\Psi^{1/2}] = \Psi^{1/2}E[X'_t\Sigma^{1/2}\Sigma_\varepsilon\Sigma^{1/2}X_t]\Psi^{1/2} = \Psi \text{tr}(\Sigma\Sigma_\varepsilon) = \Psi\sigma_*,$$

while Lemma 7 implies that

$$\begin{aligned} &E[S'_t(S_t\tilde{F} + \varepsilon)(S_t\tilde{F} + \varepsilon)'S_t] \\ &= ((\text{tr}\Sigma)^2 + \text{tr}(\Sigma^2))\Psi\Sigma_F\Psi \\ &+ \text{tr}(\Sigma \circ \Sigma)\Psi^{1/2}(\kappa - 3)\text{diag}(\Psi^{1/2}\Sigma_F\Psi^{1/2})\Psi^{1/2} + \Psi + \text{tr}(\Psi\Sigma_F)\text{tr}(\Sigma^2) \end{aligned} \quad (93)$$

The claim now follows because $\text{tr}(\Sigma \circ \Sigma)$ and $\text{tr}(\Sigma^2)$ converge to zero, and Ψ is uniformly bounded when $P \rightarrow \infty$ by assumption. The proof of Lemma 9 is complete. □

E Proof of Propositions 1 and 2: Characteristic-managed Portfolios and the Conditional SDF

Proof of Proposition 1. We have

$$\pi_t = E_t[R_{t+1}R'_{t+1}]^{-1} E_t[R_{t+1}]. \quad (94)$$

By assumption,

$$R_{t+1} = S_t \tilde{F}_{t+1} + \varepsilon_{t+1} \quad (95)$$

and, hence,

$$E_t[R_{t+1}] = S_t E_t[\tilde{F}_{t+1}] = S_t \nu \quad (96)$$

whereas

$$\begin{aligned} E_t[R_{t+1}R'_{t+1}] &= E_t[(S_t \tilde{F}_{t+1} + \varepsilon_{t+1})(S_t \tilde{F}_{t+1} + \varepsilon_{t+1})'] \\ &= S_t E_t[\tilde{F}_{t+1} \tilde{F}'_{t+1}] S'_t + S_t \Sigma_{F,\varepsilon,t} + \Sigma'_{F,\varepsilon,t} S'_t + \Sigma_{\varepsilon,t} \\ &= S_t (\Sigma_{F,t} + \nu \nu') S'_t + S_t \Sigma_{F,\varepsilon,t} + \Sigma'_{F,\varepsilon,t} S'_t + \Sigma_{\varepsilon,t}, \end{aligned} \quad (97)$$

and the claim follows. The proof is complete. $\square$

We will frequently be using the Sherman-Morrison formula

$$(A + xx')^{-1} = A^{-1} - A^{-1}xx'A^{-1}/(1 + x'A^{-1}x), \quad (A + xx')^{-1}x = A^{-1}x/(1 + x'A^{-1}x) \quad (98)$$

for any matrix $A \in \mathbb{R}^{P \times P}$ and any vector $x \in \mathbb{R}^P$.

Lemma 10 *Let A, B be two positive definite, symmetric matrices. Then,*

$$(A + B)^{-1} = B^{-1} - (A + B)^{-1}AB^{-1}, \quad (99)$$

and

$$(A + B)^{-1}AB^{-1} \leq B^{-1}AB^{-1} \quad (100)$$

in the sense of positive semi-definite order. Furthermore, if $A \leq A_1$, then

$$(A + B)^{-1}AB^{-1} \leq (A_1 + B)^{-1}A_1B^{-1} \quad (101)$$

Proof of Lemma 10. Let $\hat{A} = B^{-1/2}AB^{-1/2}$. Then, we have

$$(A + B)^{-1}AB^{-1} = B^{-1/2}(\hat{A} + I)^{-1}\hat{A}B^{-1/2} \leq B^{-1/2}\hat{A}B^{-1/2} = B^{-1}AB^{-1}. \quad (102)$$

The last claim follows from the matrix monotonicity of $(\hat{A} + I)^{-1}\hat{A}$ in $\hat{A}$: Indeed,

$$(\hat{A} + I)^{-1}\hat{A} \leq (\hat{A}_1 + I)^{-1}\hat{A}_1 \quad (103)$$

implies

$$B^{-1/2}(\hat{A} + I)^{-1}\hat{A}B^{-1/2} \leq B^{-1/2}(\hat{A}_1 + I)^{-1}\hat{A}_1B^{-1/2} \quad (104)$$

which is equivalent to (101). □

Proof of Proposition 2. Let $\Sigma_{F,t}^* = \Sigma_{F,t} + \nu\nu'$. Let also

$$\hat{\Sigma}_t = S_t \Sigma_{F,\varepsilon,t} + \Sigma'_{F,\varepsilon,t} S'_t + S_t \Sigma_F S'_t \quad (105)$$

Recall that, by Proposition 1,

$$\tilde{w}(S_t) = (\hat{\Sigma}_t + \Sigma_{\varepsilon,t})^{-1} S_t \nu \quad (106)$$

is the conditionally efficient portfolio with the return

$$R'_{t+1} \tilde{w}(S_t) = R'_{t+1} (\hat{\Sigma}_t + \Sigma_{\varepsilon,t})^{-1} S_t \nu. \quad (107)$$

For simplicity, in the sequel omit the t subindex for Σ_F , Σ_F^* , $\hat{\Sigma}$, Σ_ε . Hence, defining

$$Q_t = (\hat{\Sigma} + \Sigma_\varepsilon)^{-1} = \Sigma_\varepsilon^{-1} - (\hat{\Sigma} + \Sigma_\varepsilon)^{-1} \hat{\Sigma} \Sigma_\varepsilon^{-1}, \quad (108)$$

and noting that, by the Sherman-Morrison formula,

$$(\hat{\Sigma} + S_t \nu \nu' S_t' + \Sigma_\varepsilon)^{-1} = Q_t - Q_t S_t \nu \nu' S_t' Q_t (1 + \nu' S_t' Q_t S_t \nu)^{-1}. \quad (109)$$

Then,

$$\begin{aligned} & E[R'_{t+1} \tilde{w}(S_t)] \\ &= E[\nu' S_t' (\hat{\Sigma} + \nu \nu' + \Sigma_\varepsilon)^{-1} S_t \nu] \\ &= E[\nu' S_t' (S_t \nu \nu' S_t' + \hat{\Sigma} + \Sigma_\varepsilon)^{-1} S_t \nu] \\ &= E[\nu' S_t' ((S_t \nu)(S_t \nu)' + (\hat{\Sigma} + \Sigma_\varepsilon))^{-1} S_t \nu] \\ &\stackrel{(109)}{=} E[\nu' S_t' (Q_t - Q_t S_t \nu \nu' S_t' Q_t (1 + \nu' S_t' Q_t S_t \nu)^{-1}) S_t \nu] \\ &= E[Z_t - Z_t^2 (1 + Z_t)^{-1}] = E[Z_t / (1 + Z_t)], \end{aligned} \quad (110)$$

where we have defined

$$Z_t = \nu' S_t' Q_t S_t \nu = \nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu - q, \quad (111)$$

with

$$q = \nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu - \nu' S_t' Q_t S_t \nu. \quad (112)$$

Now, for any $\delta > 0$, we have

$$\text{Var}_t[\delta^{-1} S_t \tilde{F}_{t+1} - \delta \varepsilon_{t+1}] \geq 0$$

and, hence,

$$S_t \Sigma_{f,\varepsilon} + \Sigma_{f,\varepsilon}' S_t' \leq \delta^{-2} S_t \Sigma_F S_t' + \delta^2 \Sigma_\varepsilon \quad (113)$$

so that

$$\widehat{\Sigma} \leq (1 + \delta^{-2}) S_t \Sigma_F S_t' + \delta^2 \Sigma_\varepsilon \quad (114)$$

By Lemma 10,

$$(\widehat{\Sigma} + \Sigma_\varepsilon)^{-1} \widehat{\Sigma} \Sigma_\varepsilon^{-1} \leq \Sigma_\varepsilon^{-1} ((1 + \delta^{-2}) S_t \Sigma_F S_t' + \delta^2 \Sigma_\varepsilon) \Sigma_\varepsilon^{-1} \quad (115)$$

and hence

$$0 \leq q \leq \nu' S_t' \Sigma_\varepsilon^{-1} ((1 + \delta^{-2}) S_t \Sigma_F S_t' + \delta^2 \Sigma_\varepsilon) \Sigma_\varepsilon^{-1} S_t \nu. \quad (116)$$

Our goal is to show that $\lim_{N,P \rightarrow \infty} q = 0$ in L_1 . To this end, we show below that

$$\lim E[\nu' S_t' \Sigma_\varepsilon^{-1} S_t \Sigma_F S_t' \Sigma_\varepsilon^{-1} S_t \nu] = 0, \quad (117)$$

while

$$E[\nu' S'_t \Sigma_\varepsilon^{-1} S_t \nu] \leq K \quad (118)$$

is uniformly bounded. As a result,

$$\limsup q \leq \delta^2 K \quad (119)$$

and the desired claim follows because $\delta > 0$ is arbitrary.

For simplicity, we will assume that $X_{i,k,t}$ all have the same fourth moment κ (otherwise, the identity needs to be replaced by an inequality). Then, we have that, by Lemma 7,

$$\begin{aligned} E[\nu' S'_t \Sigma_\varepsilon^{-1} S_t A S'_t \Sigma_\varepsilon^{-1} S_t \nu] &= \nu' \left(((\text{tr } \hat{\Sigma})^2 + \text{tr}(\hat{\Sigma}^2)) \Psi A \Psi + \text{tr}(\hat{\Sigma}^2) \text{tr}(\Psi A) \Psi \right. \\ &\quad \left. + \text{tr}(\hat{\Sigma}^2) (\kappa - 3) \Psi^{1/2} \text{diag}(\Psi^{1/2} A \Psi^{1/2}) \Psi^{1/2} \right) \nu \\ &= (\text{tr } \hat{\Sigma})^2 \nu' \left(\left(1 + \frac{\text{tr}(\hat{\Sigma}^2)}{(\text{tr } \hat{\Sigma})^2} \right) \Psi A \Psi + \frac{\text{tr}(\hat{\Sigma}^2)}{(\text{tr } \hat{\Sigma})^2} \text{tr}(\Psi A) \Psi \right. \\ &\quad \left. + \frac{\text{tr}(\hat{\Sigma} \circ \hat{\Sigma})}{(\text{tr } \hat{\Sigma})^2} (\kappa - 3) \Psi^{1/2} \text{diag}(\Psi^{1/2} A \Psi^{1/2}) \Psi^{1/2} \right) \nu \end{aligned} \quad (120)$$

with

$$A = \Sigma_F \quad (121)$$

and

$$\hat{\Sigma} = \Sigma^{1/2} \Sigma_\varepsilon^{-1} \Sigma^{1/2}. \quad (122)$$

Recall, that, by assumption, $\Sigma_{\varepsilon,t} = \tilde{\Sigma}_{\varepsilon,t} + \sigma_\varepsilon^2 I$. For simplicity, we assume the normalization

$\sigma_\varepsilon^2 = 1$. Then, $\Sigma_\varepsilon^{-1} = I - \tilde{\Sigma}_{\varepsilon,t}(I + \tilde{\Sigma}_{\varepsilon,t})^{-1}$ and, hence

$$\text{tr}(\hat{\Sigma}) = \text{tr}(\Sigma^{1/2}\Sigma_\varepsilon^{-1}\Sigma^{1/2}) = \text{tr}(\Sigma) - \text{tr}(\Sigma^{1/2}\tilde{\Sigma}_{\varepsilon,t}(I + \tilde{\Sigma}_{\varepsilon,t})^{-1}\Sigma^{1/2}) \quad (123)$$

Now,

$$\Sigma^{1/2}\tilde{\Sigma}_{\varepsilon,t}(I + \tilde{\Sigma}_{\varepsilon,t})^{-1}\Sigma^{1/2} \leq \Sigma^{1/2}\tilde{\Sigma}_{\varepsilon,t}\Sigma^{1/2}$$

and, hence,

$$\text{tr}(\Sigma^{1/2}\tilde{\Sigma}_{\varepsilon,t}(I + \tilde{\Sigma}_{\varepsilon,t})^{-1}\Sigma^{1/2}) \leq \text{tr}(\Sigma^{1/2}\tilde{\Sigma}_{\varepsilon,t}\Sigma^{1/2}) = \text{tr}(\Sigma\tilde{\Sigma}_{\varepsilon,t}) \rightarrow 0$$

by assumption. Thus, we can assume that $\Sigma_\varepsilon = I$ and, by Assumption 2, $\frac{\text{tr}(\hat{\Sigma}^2) + \text{tr}(\hat{\Sigma} \circ \hat{\Sigma})}{(\text{tr} \hat{\Sigma})^2} \rightarrow 0$ and, since ν and Ψ and A and $\text{tr}(\hat{\Sigma})$ are uniformly bounded, we get that

$$E[\nu' S_t' \Sigma_\varepsilon^{-1} S_t A S_t' \Sigma_\varepsilon^{-1} S_t \nu] \approx (\text{tr} \hat{\Sigma})^2 \nu' \Psi A \Psi \nu. \quad (124)$$

Since, by Assumption 1, $\text{tr}(A)$ is uniformly bounded, we also get that $\text{tr}(\Psi A \Psi) \leq \|\Psi\|^2 \text{tr}(A)$ is uniformly bounded and, hence,

$$E[\nu' \Psi A \Psi \nu] = E[\text{tr}(\Psi A \Psi \nu \nu')] = P^{-1} \text{tr}(\Psi A \Psi \Sigma_\nu) \rightarrow 0. \quad (125)$$

Thus, $E[q_t] \rightarrow 0$ and hence $q_t \rightarrow 0$ in probability. Now,

$$E[\nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu] = \text{tr}(\hat{\Sigma}) \nu' \Psi \nu \quad (126)$$

whereas, by Lemma 7,

$$\begin{aligned}
E[(\nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu)^2] &= E[\nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu \nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu] \\
&= \nu' \left(((\text{tr} \hat{\Sigma})^2 + \text{tr}(\hat{\Sigma}^2)) \Psi \nu \nu' \Psi + \text{tr}(\hat{\Sigma}^2) \text{tr}(\Psi \nu \nu') \Psi \right. \\
&\quad \left. + \text{tr}(\hat{\Sigma} \circ \hat{\Sigma}) (\kappa - 3) \Psi^{1/2} \text{diag}(\Psi^{1/2} \nu \nu' \Psi^{1/2}) \Psi^{1/2} \right) \nu
\end{aligned} \tag{127}$$

and the same argument as in (120) implies that

$$E[(\nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu)^2] \approx \text{tr}(\hat{\Sigma})^2 (\nu' \Psi \nu)^2. \tag{128}$$

Thus, $\text{Var}[\nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu] \rightarrow 0$ and, hence, $\nu' S_t' \Sigma_\varepsilon^{-1} S_t \nu \rightarrow \text{tr}(\hat{\Sigma}) \nu' \Psi \nu$ in probability.

As a result, $Z_t - \text{tr}(\hat{\Sigma}) \nu' \Psi \nu \rightarrow 0$ in probability, and hence

$$\frac{Z_t}{1 + Z_t} - \frac{\text{tr}(\hat{\Sigma}) \nu' \Psi \nu}{1 + \text{tr}(\hat{\Sigma}) \nu' \Psi \nu} \rightarrow 0 \tag{129}$$

in probability, and the dominated convergence theorem implies that the same holds in expectation. Similarly, for the second moment, we have

$$\begin{aligned}
&E[(\pi_t^{MV})' R_{t+1}' R_{t+1} \pi_t^{MV}] \\
&= E[\lambda' S_t' (S_t (\Sigma_F) S_t' + \Sigma_\varepsilon)^{-1} (S_t (\Sigma_F) S_t' + \Sigma_\varepsilon) (S_t (\Sigma_F) S_t' + \Sigma_\varepsilon)^{-1} S_t \lambda] \\
&= E[R_{t+1}' \pi_t^{MV}] \rightarrow \frac{\text{tr}(\hat{\Sigma}) \nu' \Psi \nu}{1 + \text{tr}(\hat{\Sigma}) \nu' \Psi \nu}.
\end{aligned} \tag{130}$$

Now, for the factor portfolios, we have

$$\begin{aligned}
E[F_t] &= E[S'_t R_{t+1}] = E[S'_t (S_t \tilde{F}_{t+1} + \varepsilon_{t+1})] \\
&= E[S'_t S_t \tilde{F}_{t+1}] = E[\Psi^{1/2} X'_t \Sigma X_t \Psi^{1/2} \tilde{F}_{t+1}] \\
&= E[\Psi^{1/2} X'_t \Sigma X_t \Psi^{1/2}] \nu \\
&= \text{tr}(\Sigma) E[\Psi^{1/2} \Psi^{1/2}] \nu \\
&= \text{tr}(\Sigma) \Psi \nu,
\end{aligned} \tag{131}$$

and, again by Lemma 7 and the same argument as in (120), we have

$$\begin{aligned}
E[F_t F'_t] &= E[S'_t (S_t \tilde{F}_{t+1} + \varepsilon_{t+1}) (S_t \tilde{F}_{t+1} + \varepsilon_{t+1})' S_t] = E[S'_t (\hat{\Sigma} + \nu \nu' + \Sigma_\varepsilon) S_t] \\
&\approx \text{tr}(\Sigma \Sigma_\varepsilon) \Psi + \text{tr}(\Sigma)^2 \Psi \nu \nu' \Psi + E[S'_t \hat{\Sigma} S_t].
\end{aligned} \tag{132}$$

(Note that we cannot apply Lemma 9 directly because we are not assuming $E[X_{j,k,t}^3] = 0$ here. In particular, our $\hat{\Sigma}$ contains $\Sigma_{F,\varepsilon}$.) Our goal is to show that the contribution of $A = E[S'_t \hat{\Sigma} S_t]$ is negligible. Then, defining

$$Q = (\text{tr}(\Sigma \Sigma_\varepsilon) \Psi + A)^{-1}, \tag{133}$$

we get that the efficient portfolio of factors is given by

$$\begin{aligned}
\pi_F &= (\text{tr}(\Sigma \Sigma_\varepsilon) \Psi + \text{tr}(\Sigma)^2 \Psi \nu \nu' \Psi + A)^{-1} \Psi \nu \\
&= (\text{tr}(\Sigma \Sigma_\varepsilon) \Psi + \text{tr}(\Sigma)^2 \Psi \nu \nu' \Psi + A)^{-1} \Psi \nu \\
&\stackrel{(98)}{=} \frac{1}{1+Z} Q \Psi \nu,
\end{aligned} \tag{134}$$

where

$$Z = \text{tr}(\Sigma)^2 \nu' \Psi Q \Psi \nu. \tag{135}$$

By the same argument as above,

$$\nu'(\Psi Q \Psi - \Psi(\text{tr}(\Sigma \Sigma_\varepsilon) \Psi)^{-1} \Psi) \nu \rightarrow 0 \quad (136)$$

because Σ_F has a bounded trace and we can bound A by a combination of $(1 + \delta^2) \Sigma_F + \delta^2 \Sigma_\varepsilon$ with an arbitrarily small δ . Thus,

$$Z \approx \frac{\text{tr}(\Sigma)^2}{\text{tr}(\Sigma \Sigma_\varepsilon)} \nu' \Psi \nu \quad (137)$$

and

$$E[\pi'_F F_{t+1}] = E[\nu' \frac{1}{1+Z} \Psi Q \Psi \nu] \approx \frac{Z}{1+Z}, \quad (138)$$

while

$$E[\pi'_F F_{t+1} F'_{t+1} \pi_F] = E[\pi'_F F_{t+1}], \quad (139)$$

and the proof is complete because

$$\text{tr}(\Sigma \Sigma_\varepsilon^{-1}) \nu' \Psi \nu = \frac{\text{tr}(\Sigma)^2}{\text{tr}(\Sigma \Sigma_\varepsilon)} \nu' \Psi \nu \quad (140)$$

when $\Sigma_\varepsilon = I$.

Finally, the fact that $E[(\pi'_F F_{t+1} - R'_{t+1} \tilde{w}(S_t))^2] \rightarrow 0$ follows because, otherwise, one could construct a better-diversified portfolio by combining the two, which is impossible. Namely, we know that

$$E[R'_{t+1} \tilde{w}(S_t) - 0.5((R'_{t+1} \tilde{w}(S_t))^2)] \geq E[R'_{t+1} b_t - 0.5((R'_{t+1} b_t^2)] \quad (141)$$

for any portfolio policy b_t because $\tilde{w}(S_t)$ is optimal for the agent with quadratic utility. Let

$b_t = \tilde{w}(S_t) + \varepsilon S_t \pi_F$. Then,

$$\max_{\varepsilon} E[R'_{t+1} b_t - 0.5((R'_{t+1} b_t)^2)] \quad (142)$$

is achieved with

$$\varepsilon = \frac{E[\pi'_F F_{t+1}] - E[\pi'_F F_{t+1} R'_{t+1} \tilde{w}(S_t)]}{E[(\pi'_F F_{t+1})^2]}. \quad (143)$$

Since the corresponding utility gain is zero, we must have

$$E[\pi'_F F_{t+1}] - E[\pi'_F F_{t+1} R'_{t+1} \tilde{w}(S_t)] = 0. \quad (144)$$

Thus,

$$\begin{aligned} E[(\pi'_F F_{t+1} - R'_{t+1} \tilde{w}(S_t))^2] &= E[(\pi'_F F_{t+1})^2] + E[(R'_{t+1} \tilde{w}(S_t))^2] - 2E[\pi'_F F_{t+1} R'_{t+1} \tilde{w}(S_t)] \\ &= E[(\pi'_F F_{t+1})^2] + E[(R'_{t+1} \tilde{w}(S_t))^2] - 2E[\pi'_F F_{t+1}] \\ &\approx \frac{Z}{1+Z} + \frac{Z}{1+Z} - 2\frac{Z}{1+Z} = 0. \end{aligned} \quad (145)$$

The proof is complete. $\square$

F A Proof that Managed Portfolio Returns Satisfy the Assumptions of RMT

We will use the notation

$$B_T = \frac{1}{T} \sum_{t=1}^T F_t F'_t \quad (146)$$

The goal of this section is to prove the following theorem.

Theorem 4 *The eigenvalue distribution of $E[F_t F_t']$ converges to that of $\Psi \sigma_*$ where $\sigma_* = \lim \text{tr}(\Sigma \Sigma_\varepsilon)$ in the limit as $N, P, T \rightarrow \infty$, $P/T \rightarrow c$, so that*

$$\lim \frac{1}{P} \text{tr}((zI + E[F_t F_t'])^{-1}) = \lim \sigma_*^{-1} m_\Psi(-z/\sigma_*) = \lim m_{\sigma_* \Psi}(-z) = \lim \frac{1}{P} \text{tr}((zI + \sigma_* \Psi)^{-1}), \quad (147)$$

whereas

$$\frac{1}{P} \text{tr}((zI + B_T)^{-1}) \rightarrow m(-z; c), \quad (148)$$

where, for each $z < 0$, we have that $m(z; c)$ is the unique positive solution to the nonlinear master equation

$$m(z; c) = \frac{1}{1 - c - c z m(z; c)} m_{\sigma_* \Psi} \left(\frac{z}{1 - c - c z m(z; c)} \right). \quad (149)$$

where we abuse the notation and use $m_{\sigma_* \Psi}$ to denote the limit $\lim \frac{1}{P} \text{tr}((zI + \sigma_* \Psi)^{-1})$.

The master equation (149) for the asymptotic Stieltjes transform has been established in many papers under various technical conditions. For example, [Hastie et al. \(2019\)](#) establish it under the classic assumption that $F_t = \Psi^{1/2} \tilde{F}_t$, where $\tilde{F}_t$ have i.i.d. coordinates. Such an i.i.d. assumption implies that there are no dependencies in the higher-order moments of F_t , all dependence structure is captured by Ψ . This assumption is clearly violated for managed portfolios $F_{t+1} = S_t' R_{t+1} = S_t' (S_t \tilde{F}_{t+1} + \varepsilon_{t+1})$ because F_{t+1} depends *quadratically* on S_t .

Here, we rely on a powerful result of [\(Bai and Zhou, 2008\)](#) who show that (149) holds as long as we have the *concentration of quadratic forms*: $E[(F_{t+1}' B F_{t+1} - \text{tr}(A_P B_P))^2] = o(T^2)$. Establishing this concentration result for managed portfolios is one of the key technical results in our paper. The following is true.

Lemma 11 (Managed Portfolios Satisfy The RMT Conditions) *Suppose that $P, T \rightarrow \infty, P/T \rightarrow c > 0$. Let A_P be a sequence of symmetric $P \times P$ matrices such that $\|A_P\| \leq K$ and A_P are independent of F_t . Then, $E[F_t F_t']$ is uniformly bounded and*

$$\text{Var}\left[\frac{1}{T} F_t' A_P F_t\right] = T^{-2}(O(P) + O(P^2 \text{tr}(\Sigma^2))), \quad (150)$$

so that

$$\frac{1}{T} (F_t' A_P F_t - \text{tr}(A_P \sigma_{*,t} \Psi)) \rightarrow 0$$

in probability, where $\sigma_{*,t} = \text{tr}(\tilde{\Sigma}_{\varepsilon,t} \Sigma)$. That is, averaging across P factors leads to constant risk, no matter which matrix A we use to measure it.

We can now prove Theorem 4.

Proof of Theorem 4. The first claim follows because, by Lemma 19, the other contributions do not impact eigenvalue distribution.

Let $P = P^*$. To prove the claim about the eigenvalue distribution of B_T , we use a Theorem of (Bai and Zhou, 2008). According to (Bai and Zhou, 2008), defining $F_{t+1} = S_t' R_{t+1}$, we need to verify the following technical conditions:

- (1) $E[F_{t+1} F_{t+1}'] = A_P$ for some matrix A_P
- (2) $E[(F_{t+1}' B F_{t+1} - \text{tr}(A_P B_P))^2] = o(T^2)$ for any bounded matrix sequence B_P , $P > 0$.
- (3) The norm of A_P is uniformly bounded, and its eigenvalue distribution converges as $P \rightarrow \infty$.

The only non-trivial claim here is item (2), which in turn follows from Lemma 11.

Then, recall that, under Assumption 4, the limit as $P \rightarrow \infty$, the eigenvalue distribution of $\Psi(q) = U(q) D(q) U(q)'$, $D(q) = \text{diag}(D_i(q))$, converges to a limit distribution $dH(x; q)$

supported on a bounded interval. Defining $F(q) = F[:P]$, we get from the above arguments that $E[F(q)] = E[F][:P] = (\Psi\nu)[:P]$. Similarly, for any matrix $B_P \in \mathbb{R}^{P \times P}$, we can list if to a $B_{P^*} \in P^* \times P^*$ by filling zeros and get

$$E[(F_{t+1}(q)' B_P F_{t+1}(q) - \text{tr}(A_P B_P))^2] = E[(F_{t+1}' B_{P^*} F_{t+1} - \text{tr}(A_{P^*} B_{P^*}))^2], \quad (151)$$

so that our argument above applies as well to $F_t(q)$.

The proof of Theorem 4 is complete. $\square$

G Proof of Lemma 11: Concentration of Quadratic Forms for Managed Portfolios

Roadmap: The goal of this section is to prove Lemma 11, which is the key ingredient to the proof of Theorem 4. This proof is done by direct calculation, carefully evaluating the expectation $T^{-2}E[(F_t' A F_t)^2]$ and showing that most of the terms are asymptotically negligible. The number of terms involved is very large: since F_t depend quadratically on S_t , we get that $(F_t' A F_t)^2$ involves moments of S_t or degree eight. The proof is based on the following technical lemmas.

Lemma 12 *Under Assumption 4, $E[(\tilde{F}_{t+1}' A \tilde{F}_{t+1})^2 | A] \leq K \|A\|^2$ for any symmetric matrix A , where K is a constant.*

Proof of Lemma 12. Without loss of generality, we may assume that A is positive definite (otherwise, we just decompose into positive and negative spectral components). Then,

$$\tilde{F}_{t+1} = \nu + \tilde{\Sigma}_{F,t}^{1/2} \tilde{X}_{t+1}$$

and the inequality

$$|x' Ay| \leq 2(x' Ax + y' Ay) \quad (152)$$

implies

$$0.5E[(\tilde{F}'_{t+1} A \tilde{F}_{t+1})^2 | A] \leq E[(\nu' A \nu)^2 | A] + E[(\tilde{X}'_{t+1} \tilde{\Sigma}_{F,t}^{1/2} A \tilde{\Sigma}_{F,t}^{1/2} \tilde{X}_{t+1})^2 | A] \quad (153)$$

By (64), using that $\nu' A \nu = P^{-1} \hat{\nu}' \Sigma_\nu^{1/2} A \Sigma_\nu^{1/2} \hat{\nu}$, as well as the simple inequality

$$\text{tr}(B^2) = \|B\|_2^2 \leq P \|B\|^2 \text{ for } B = B',$$

we get

$$\begin{aligned} & E[(\nu' A \nu)^2 | A] + E[(\tilde{X}'_{t+1} \tilde{\Sigma}_{F,t}^{1/2} A \tilde{\Sigma}_{F,t}^{1/2} \tilde{X}_{t+1})^2 | A] \\ & \leq (E[\nu' A \nu | A])^2 + (E[(\tilde{X}'_{t+1} \tilde{\Sigma}_{F,t}^{1/2} A \tilde{\Sigma}_{F,t}^{1/2} \tilde{X}_{t+1}) | A])^2 \\ & + k P^{-1} \|\Sigma_\nu^{1/2} A \Sigma_\nu^{1/2}\|_2^2 + k \|\tilde{\Sigma}_{F,t}^{1/2} A \tilde{\Sigma}_{F,t}^{1/2}\|_2^2 \\ & \leq \tilde{K} (P^{-2} \text{tr}(\Sigma_\nu A)^2 + \text{tr}(\tilde{\Sigma}_{F,t} A)^2 + \|A\|_2^2) \\ & \leq K \|A\|^2 \end{aligned} \quad (154)$$

for some $K > 0$ that only depends on $\|\Sigma_\nu\|$, $\|\tilde{\Sigma}_{F,t}\|$, $\text{tr}(\tilde{\Sigma}_{F,t})$. Here, we have used the *factor structure*: The fact that, by assumption, $\text{tr}(\tilde{\Sigma}_{F,t})$ stays uniformly bounded. $\square$

An important observation is that by Lemma 9,

$$\frac{1}{T} \text{tr}(A_P E[F_t F_t']) = \frac{1}{T} \text{tr}(A_P (\Psi \Sigma_F \Psi + \sigma_* \Psi)) + O(\text{tr}(\Sigma^2)) + O(1/T) \quad (155)$$

However, since Σ_F has a uniformly bounded trace norm, we have

$$\frac{1}{T} \text{tr}(A_P(\Psi \Sigma_F \Psi + \sigma_* \Psi)) = \frac{1}{T} \text{tr}(A_P(\sigma_* \Psi)) + O(1/T). \quad (156)$$

Lemma 13 *Let $\{X_{i,j}\}$ be independent random variables satisfying*

$$E[X_{i,j}] = 0, \quad E[X_{i,j}^2] = 1, \quad E[X_{i,j}^4] = \mu_4,$$

and let $\tilde{A}$ and Σ be symmetric matrices. Define

$$S = \sum_{i_1, i_2, i_3, i_4} \tilde{A}_{i_1, i_2} X_{i_3, i_2} \Sigma_{i_3, i_4} X_{i_4, i_1}.$$

Then,

$$E[S^2] = 2 \text{tr}(\tilde{A}^2) \text{tr}(\Sigma^2) + \text{tr}(\tilde{A})^2 \text{tr}(\Sigma)^2 + (\mu_4 - 3) \text{tr}(\tilde{A} \circ \tilde{A}) \text{tr}(\Sigma \circ \Sigma).$$

Proof of Lemma 13. We start by writing

$$S = \sum_{i_1, i_2, i_3, i_4} \tilde{A}_{i_1, i_2} X_{i_3, i_2} \Sigma_{i_3, i_4} X_{i_4, i_1},$$

so that

$$E[S^2] = E \left[\left(\sum_{i_1, i_2, i_3, i_4} \tilde{A}_{i_1, i_2} X_{i_3, i_2} \Sigma_{i_3, i_4} X_{i_4, i_1} \right)^2 \right].$$

Expanding the square and interchanging summation with expectation gives

$$E[S^2] = \sum_{\substack{i_1, i_2, i_3, i_4 \\ i'_1, i'_2, i'_3, i'_4}} \tilde{A}_{i_1, i_2} \tilde{A}_{i'_1, i'_2} \Sigma_{i_3, i_4} \Sigma_{i'_3, i'_4} E[X_{i_3, i_2} X_{i_4, i_1} X_{i'_3, i'_2} X_{i'_4, i'_1}].$$

Since the $X_{i,j}$ are independent with mean zero and unit variance,

$$E[X_{i,j}X_{k,l}] = \delta_{ik}\delta_{jl},$$

and in general the fourth moment for non-Gaussian variables can be written as

$$\begin{aligned} E[X_{i_3,i_2} X_{i_4,i_1} X_{i'_3,i'_2} X_{i'_4,i'_1}] = & \\ & \underbrace{\delta_{i_3,i'_3} \delta_{i_2,i'_2} \delta_{i_4,i'_4} \delta_{i_1,i'_1}}_{\text{Pairing 1}} + \underbrace{\delta_{i_3,i'_4} \delta_{i_2,i'_1} \delta_{i_4,i'_3} \delta_{i_1,i'_2}}_{\text{Pairing 2}} \\ & + \underbrace{\delta_{i_3,i_4} \delta_{i_2,i_1} \delta_{i'_3,i'_4} \delta_{i'_2,i'_1}}_{\text{Pairing 3}} + (\mu_4 - 3) \underbrace{\delta_{i_3,i_4} \delta_{i_2,i_1} \delta_{i_3,i'_3} \delta_{i_2,i'_2} \delta_{i_4,i'_4} \delta_{i_1,i'_1}}_{\text{Non-Gaussian correction}}. \end{aligned}$$

We now treat each contribution separately.

Term 1 (Pairing 1): The deltas

$$\delta_{i_3,i'_3} \delta_{i_2,i'_2} \delta_{i_4,i'_4} \delta_{i_1,i'_1}$$

force $i'_3 = i_3$, $i'_2 = i_2$, $i'_4 = i_4$, and $i'_1 = i_1$, yielding

$$T_1 = \sum_{i_1,i_2,i_3,i_4} \tilde{A}_{i_1,i_2}^2 \Sigma_{i_3,i_4}^2 = \left(\sum_{i_1,i_2} \tilde{A}_{i_1,i_2}^2 \right) \left(\sum_{i_3,i_4} \Sigma_{i_3,i_4}^2 \right).$$

Term 2 (Pairing 2): The deltas

$$\delta_{i_3,i'_4} \delta_{i_2,i'_1} \delta_{i_4,i'_3} \delta_{i_1,i'_2}$$

imply $i'_3 = i_4$, $i'_4 = i_3$, $i'_1 = i_2$, and $i'_2 = i_1$, so that

$$T_2 = \sum_{i_1,i_2,i_3,i_4} \tilde{A}_{i_1,i_2} \tilde{A}_{i_2,i_1} \Sigma_{i_3,i_4} \Sigma_{i_4,i_3}.$$

Term 3 (Pairing 3): The deltas

$$\delta_{i_3, i_4} \delta_{i_2, i_1} \delta_{i'_3, i'_4} \delta_{i'_2, i'_1}$$

force, separately, $i_3 = i_4$, $i_2 = i_1$ (in the first copy) and $i'_3 = i'_4$, $i'_2 = i'_1$ (in the second copy).

Hence,

$$T_3 = \left(\sum_{i_1, i_3} \tilde{A}_{i_1, i_1} \Sigma_{i_3, i_3} \right)^2 = \left(\sum_i \tilde{A}_{i, i} \sum_j \Sigma_{j, j} \right)^2.$$

Term 4 (Non-Gaussian Correction): The additional term is

$$(\mu_4 - 3) \delta_{i_3, i_4} \delta_{i_2, i_1} \delta_{i_3, i'_3} \delta_{i_2, i'_2} \delta_{i_4, i'_4} \delta_{i_1, i'_1},$$

which again forces $i_3 = i_4$, $i_2 = i_1$, and further $i'_3 = i_3$, $i'_2 = i_2$, $i'_4 = i_4$, and $i'_1 = i_1$. Thus,

$$T_4 = (\mu_4 - 3) \sum_{i_1, i_3} \tilde{A}_{i_1, i_1}^2 \Sigma_{i_3, i_3}^2.$$

Combining all terms, we obtain

$$\begin{aligned} E[S^2] &= T_1 + T_2 + T_3 + T_4 \\ &= \left(\sum_{i_1, i_2} \tilde{A}_{i_1, i_2}^2 \right) \left(\sum_{i_3, i_4} \Sigma_{i_3, i_4}^2 \right) + \sum_{i_1, i_2, i_3, i_4} \tilde{A}_{i_1, i_2} \tilde{A}_{i_2, i_1} \Sigma_{i_3, i_4} \Sigma_{i_4, i_3} \\ &\quad + \left(\sum_i \tilde{A}_{i, i} \sum_j \Sigma_{j, j} \right)^2 + (\mu_4 - 3) \sum_{i_1, i_3} \tilde{A}_{i_1, i_1}^2 \Sigma_{i_3, i_3}^2 \\ &= 2 \operatorname{tr}(\tilde{A}^2) \operatorname{tr}(\Sigma^2) + \operatorname{tr}(\tilde{A})^2 \operatorname{tr}(\Sigma)^2 + (\mu_4 - 3) \operatorname{tr}(\tilde{A} \circ \tilde{A}) \operatorname{tr}(\Sigma \circ \Sigma). \end{aligned}$$

This completes the proof. □

Corollary 5 *Let $Z_t = S'_{t-1}S_{t-1}$. Then,*

$$E[(\text{tr}(Z_t A))^2] \leq K(\|A\|_1^2 + \|A\|_2^2 \text{tr}(\Sigma^2)) \quad (157)$$

for some $K > 0$.

Note that $\text{tr}(A_P F_t F'_t) = F'_t A_P F_t$. For simplicity, we will assume that A_P is deterministic.⁴⁶ We can also assume that A_P is symmetric because $F'_t A_P F_t = F'_t 0.5(A_P + A'_P) F_t$. And we can assume it is positive definite (otherwise, we can decompose $A = A_+ - A_-$). We need to prove that

$$\frac{1}{T^2} E[F'_t A_P F_t F'_t A_P F_t] - \left(\frac{1}{T} E[F'_t A_P F_t] \right)^2 \rightarrow 0$$

We have by Lemma 9 that

$$\begin{aligned} E[F_t F'_t] &= ((\text{tr} \Sigma)^2 + \text{tr}(\Sigma^2)) \Psi \Sigma_F \Psi \\ &+ \text{tr}(\Sigma \circ \Sigma) \Psi^{1/2} (\kappa - 3) \text{diag}(\Psi^{1/2} \Sigma_F \Psi^{1/2}) \Psi^{1/2} + \Psi \left(\text{tr}(\Sigma \Sigma_\varepsilon) + \text{tr}(\Psi \Sigma_F) \text{tr}(\Sigma^2) \right) \end{aligned} \quad (158)$$

and, with Σ_F having uniformly bounded traces and Assumption 2, we get

$$\begin{aligned} \frac{1}{T} E[F'_t A_P F_t] &= \frac{1}{T} \text{tr} E[A_P F_t F'_t] \\ &= \frac{1}{T} \text{tr} \left(A_P \left((\text{tr} \Sigma)^2 \Psi \Sigma_F \Psi + \text{tr}(\Sigma \circ \Sigma) \Psi^{1/2} (\kappa - 3) \text{diag}(\Psi^{1/2} \Sigma_F \Psi^{1/2}) \Psi^{1/2} \right. \right. \\ &\quad \left. \left. + \Psi \left(\text{tr}(\Sigma \Sigma_\varepsilon) + \text{tr}(\Psi \Sigma_F) \text{tr}(\Sigma^2) \right) \right) \right) \\ &= T^{-1} \text{tr}(A_P \Psi) \sigma_* + o(1) \end{aligned} \quad (159)$$

⁴⁶Otherwise, we replace all expectations below by the expectations conditional on A_P .

since

$$\frac{1}{T} \text{tr}(\Psi A_P \Psi \Sigma_F) = O(1/T),$$

and the kurtosis term does not matter because it has a uniformly bounded trace. From now on, for simplicity, we perform calculations with the normalization $\sigma_* = 1$.

Throughout the proof, we will use the notation

$$\beta \equiv \tilde{F}_t, \tag{160}$$

so that

$$F_t = S'_{t-1} R_t = S'_{t-1} (S_{t-1} \beta + \varepsilon_t). \tag{161}$$

Note that, by Lemma 12 with $A = I$,

$$E[\|\beta\|^4] \tag{162}$$

is uniformly bounded.

Then, we have

$$\begin{aligned} F_t F'_t &= S'_{t-1} (S_{t-1} \beta \beta' S'_{t-1} + \varepsilon_t \beta' S'_{t-1} + S_{t-1} \beta \varepsilon'_t + \varepsilon_t \varepsilon'_t) S_{t-1} \\ &= Z_t \beta \beta' Z_t + S'_{t-1} \varepsilon_t \beta' Z_t + Z_t \beta \varepsilon'_t S_{t-1} + S'_{t-1} \varepsilon_t \varepsilon'_t S_{t-1} \end{aligned} \tag{163}$$

with $Z_t = S'_{t-1} S_{t-1}$.

Roadmap. Our goal is to show that

$$\frac{1}{T^2} E[F'_t A F_t F'_t A F_t] = \frac{1}{T^2} E[F'_t A F_t]^2 + O(T^{-1}) + O(\text{tr}(\Sigma^2)). \tag{164}$$

(Recall that we normalize $\text{tr}(\Sigma) = 1$ and assume that $\lim_{N \rightarrow \infty} \text{tr}(\Sigma^2)/(\text{tr}(\Sigma))^2 = 0$ in Assumption 2). In order to do that, we need to analytically compute the expectation involving fourth-order moments of all factors and then show that all the interaction terms are $O(T^{-1}) + O(\text{tr}(\Sigma^2))$. Due to the three infinities problem (large N, T, P), this requires long and tedious calculations presented below.

By assumption $E[\tilde{\varepsilon}_{i,t}] = E[\tilde{\varepsilon}_{i,t}^3] = 0$ and, hence, all terms that are linear or cubic in ε vanish, and we get

$$\begin{aligned}
\frac{1}{T^2} E[F_t' A F_t F_t' A F_t] &= \frac{1}{T^2} \text{tr} E[F_t F_t' A F_t F_t' A] \\
&= \frac{1}{T^2} \text{tr} E[(Z_t \beta \beta' Z_t + S_{t-1}' \varepsilon_t \beta' Z_t + Z_t \beta \varepsilon_t' S_{t-1} + S_{t-1}' \varepsilon_t \varepsilon_t' S_{t-1}) A \\
&\quad (Z_t \beta \beta' Z_t + S_{t-1}' \varepsilon_t \beta' Z_t + Z_t \beta \varepsilon_t' S_{t-1} + S_{t-1}' \varepsilon_t \varepsilon_t' S_{t-1}) A] \\
&= \frac{1}{T^2} \text{tr} E[Z_t \beta \beta' Z_t A Z_t \beta \beta' Z_t A] \\
&\quad + \frac{1}{T^2} 2 \text{tr} E[Z_t \beta \beta' Z_t A S_{t-1}' \varepsilon_t \varepsilon_t' S_{t-1} A] \\
&\quad + \frac{1}{T^2} 2 \text{tr} E[S_{t-1}' \varepsilon_t \beta' Z_t A S_{t-1}' \varepsilon_t \beta' Z_t A] \\
&\quad + \frac{1}{T^2} 2 \text{tr} E[S_{t-1}' \varepsilon_t \beta' Z_t A Z_t \beta \varepsilon_t' S_{t-1} A] \\
&\quad + \frac{1}{T^2} \text{tr} E[S_{t-1}' \varepsilon_t \varepsilon_t' S_{t-1} A S_{t-1}' \varepsilon_t \varepsilon_t' S_{t-1} A]
\end{aligned} \tag{165}$$

We will use Lemma 8, applying it to ε and β and rewriting each term in (165).

There is a slight complication because $\beta = P^{-1/2} \Sigma_\nu^{1/2} \hat{\nu} + \Sigma_{F,t}^{1/2} \tilde{X}$. Therefore, all expressions involving β quadratic form $\beta' Q \beta$ will involve a cross term of $\nu, \tilde{X}$. However, using (152), we can proceed as in (153) and bound the cross-term with terms so that, for any positive-definite matrix Q we have

$$(\beta' Q \beta)^2 \leq 2(\nu' Q \nu)^2 + (\tilde{X}' \Sigma_{F,t}^{1/2} Q \Sigma_{F,t}^{1/2} \tilde{X})^2 \tag{166}$$

Now, both ν and $\Sigma_{F,t}^{1/2} \tilde{X}$ have the property that $\text{tr} E[\nu \nu']$ and $\text{tr} E[\Sigma_{F,t}^{1/2} \tilde{X} \tilde{X}' \Sigma_{F,t}^{1/2}]$ are uniformly

bounded. This is the key property we will use below to show that all terms involving β are negligible.

Everywhere in the sequel, we use

$$\mathring{\Sigma} = P^{-1}\Sigma_\nu + \tilde{\Sigma}_{F,t}. \quad (167)$$

G.1 The Leading Term in (165): Dealing with 8–th order moments.

The goal of this subsection is to prove the following lemma.

Lemma 14 *We have*

$$\frac{1}{T^2} \text{tr} E[Z_t \beta \beta' Z_t A Z_t \beta \beta' Z_t A] \leq T^{-2} K (P + (\text{tr}(\Sigma^2))^2 P^2) (\text{tr}(\mathring{\Sigma}))^2 \|\Psi\|^4 \quad (168)$$

for some constant $K > 0$.

Proof of Lemma 14. The proof of Lemma 14 is based on a long calculation, carefully bounding the numerous terms appearing in the expression.

We start by noting that, for any positive definite matrix A ,

$$\text{tr}(A^2) \leq (\text{tr}(A))^2 \quad (169)$$

and, as a result,

$$\text{tr}((B^{1/2} A B^{1/2})^2) \leq (\text{tr}(AB))^2 \quad (170)$$

for any positive definite A, B . This allows us to bound the first term in (165) as

$$\begin{aligned}
& \frac{1}{T^2} \text{tr} E[Z_t \beta \beta' Z_t A Z_t \beta \beta' Z_t A] \\
&= \frac{1}{T^2} \text{tr} E[\beta' Z_t A Z_t \beta \beta' Z_t A Z_t \beta] = \frac{1}{T^2} \text{tr} E[(\beta' Z_t A Z_t \beta)^2] \\
&\quad \underbrace{\leq}_{(154) + (170)} \frac{K}{T^2} \text{tr} E[(\text{tr}(Z_t A Z_t \overset{\circ}{\Sigma}))^2] = \frac{K}{T^2} E[(\text{tr}(S'_{t-1} S_{t-1} A S'_{t-1} S_{t-1} \overset{\circ}{\Sigma}))^2] \\
&= \frac{K}{T^2} E[(\text{tr}(\Psi^{1/2} X'_{t-1} \Sigma X_{t-1} \Psi^{1/2} A \Psi^{1/2} X'_{t-1} \Sigma X_{t-1} \Psi^{1/2} \overset{\circ}{\Sigma}))^2]
\end{aligned} \tag{171}$$

Lemma 15 *We have*

$$E[(\text{tr}(X'_{t-1} A X_{t-1} B X'_{t-1} A X_{t-1} C))^2] \leq \|B\| \|\mathbb{C}\|_1^2 K(P \text{tr}(A)^4 + P^2(\text{tr}(A^2))^2) \tag{172}$$

for some $K > 0$.

Proof of Lemma 15. Without loss of generality, we may assume that A, B, C are positive definite. Let also

$$Y_{i,j} = X'_i A X_j, \quad Y = (X' A X) \tag{173}$$

where X_i is the i 'th column of X . Then,

$$\text{tr}(Y B Y C) = \text{tr}(C^{1/2} Y B Y C^{1/2}) \leq \|B\| \text{tr}(C^{1/2} Y^2 C^{1/2}) = \|B\| \text{tr}(Y^2 C) \tag{174}$$

Thus, from now on, we assume that $B = I$. We will now proceed as follows. First, we will derive the proof for a rank-one $C = \gamma \gamma'$. Then, we will use spectral decomposition of C to derive the bounds for general C with a bounded trace.

We thus have

$$T = \sum_{i_1, i_2, i_4} Y_{i_1, i_2} Y_{i_2, i_3} \gamma_{i_1} \gamma_{i_3} \tag{175}$$

and

$$T^2 = \sum_{i_1, i_2, i_3} Y_{i_1, i_2} Y_{i_2, i_3} \gamma_{i_1} \gamma_{i_3} \sum_{j_1, j_2, j_3} Y_{j_1, j_2} Y_{j_2, j_3} \gamma_{j_3} \gamma_{j_1} \quad (176)$$

Since Y is symmetric, we can rewrite it as

$$T^2 = \sum_{i_1, i_2, i_3} \sum_{j_1, j_2, j_3} Y_{i_1, i_2} Y_{i_3, i_2} \gamma_{i_1} \gamma_{i_3} Y_{j_1, j_2} Y_{j_3, j_2} \gamma_{j_3} \gamma_{j_1}. \quad (177)$$

Our goal is bound $E[T^2]$. First, we derive inequalities for moments of Y .

Lemma 16 *We have*

•

$$\begin{aligned} E[Y_{i,j}^2] &= E[X_i' A X_j X_j' A X_i] = \text{tr}(A^2), \quad i \neq j \\ E[Y_{i,i}^2] &\leq K \text{tr}(A)^2 \end{aligned} \quad (178)$$

for some $K > 0$. Furthermore,

$$E[Y_{i,i}^4] \leq K \text{tr}(A)^4 \quad (179)$$

by Lemma 2. Furthermore, by (64) and since $\text{tr}(A^2) \leq (\text{tr}(A))^2$ for positive definite matrices, we get for $i \neq j$ that

$$\begin{aligned} E[Y_{i,j}^4] &= E[X_i' A X_j X_j' A X_i X_i' A X_j X_j' A X_i] = E[(X_i' (A X_j X_j' A) X_i)^2] \\ &\leq E[(\text{tr}(A X_j X_j' A))^2 + K \text{tr}((A X_j X_j' A)^2)] \\ &\leq (K+1) E[(\text{tr}(A X_j X_j' A))^2] = (K+1) E[(\text{tr}(X_j' A A X_j))^2] = (K+1) E[(X_j' A^2 X_j)^2] \\ &\leq (K+1)^2 (\text{tr}(A^2))^2 \end{aligned} \quad (180)$$

Suppose first that $i_1 = i_3$. Then, the only non-zero contributions come from $j_1 = j_3$, and we get

$$\sum_{i_1, j_1} \gamma_{i_1}^2 \gamma_{j_1}^2 \sum_{i_2, j_2} E[Y_{i_1, i_2}^2 Y_{j_1, j_2}^2]. \quad (181)$$

Our goal is to bound $\sum_{i_2, j_2} E[Y_{i_1, i_2}^2 Y_{j_1, j_2}^2]$. From Lemma 16, we know that all moments and cross-moments of Y are uniformly bounded: By Cauchy-Schwarz,

$$E[Y_{i_1, i_2}^2 Y_{j_1, j_2}^2] \leq (E[Y_{i_1, i_2}^4] E[Y_{j_1, j_2}^4])^{1/2}. \quad (182)$$

The number of summands where $i_2 = i_1$ or $j_1 = j_2$ is $O(P \operatorname{tr}(A)^4)$, and the rest $O(P^2)$ summands are bounded by $K(\operatorname{tr}(A^2))^2$, so that

$$\begin{aligned} & \sum_{i_1, j_1} \gamma_{i_1}^2 \gamma_{j_1}^2 \sum_{i_2, j_2} E[Y_{i_1, i_2}^2 Y_{j_1, j_2}^2] \\ & \leq \sum_{i_1, j_1} \gamma_{i_1}^2 \gamma_{j_1}^2 K(P \operatorname{tr}(A)^4 + P^2(\operatorname{tr}(A^2))^2) \\ & \leq \|\gamma\|^4 K(P \operatorname{tr}(A)^4 + P^2(\operatorname{tr}(A^2))^2). \end{aligned} \quad (183)$$

It remains to bound the terms

$$\sum_{i_1 \neq i_3} \sum_{j_1 \neq j_3} \gamma_{i_1} \gamma_{i_3} \gamma_{j_3} \gamma_{j_1} \sum_{i_2, j_2} E[Y_{i_1, i_2} Y_{i_3, i_2} Y_{j_1, j_2} Y_{j_3, j_2}]. \quad (184)$$

Since i_2 and j_2 already come in pairs, the only non-zero terms are where inside (i_1, i_3, j_1, j_3) we have two identical pairs of integers (or, all of them coincide). Thus, we need to consider $(i_1 = j_1, i_3 = j_3)$ and $(i_1 = j_3, i_3 = j_1)$. It is easy to see that the same argument as above applies. For example,

$$\sum_{i_1, j_1} \sum_{j_1 \neq j_3} \gamma_{i_1}^2 \gamma_{j_3}^2 \sum_{i_2, j_2} E[Y_{i_1, i_2}^2 Y_{i_3, i_2}^2]. \quad (185)$$

is bounded from above by $\|\gamma\|^4(P \operatorname{tr}(A)^4 + P^2(\operatorname{tr}(A^2))^2)$ by the same Cauchy-Schwarz inequality. Now, let $C = \sum_i \gamma_i \gamma_i'$ be the spectral decomposition, with $\operatorname{tr}(C) = \sum_i \|\gamma_i\|^2$. Since T is linear in C , we get from the above

$$E[T^2] \leq \sum_{i,j} \|\gamma_i\|^2 \|\gamma_j\|^2 K(P \operatorname{tr}(A)^4 + P^2(\operatorname{tr}(A^2))^2) \quad (186)$$

The proof of Lemma 14 is complete. $\square$

G.2 Second Term in (165)

We have

$$\begin{aligned} & \frac{1}{T^2} 2 |\operatorname{tr} E[Z_t \beta \beta' Z_t A S'_{t-1} \varepsilon_t \varepsilon_t' S_{t-1} A]| \\ &= \frac{1}{T^2} 2 |\operatorname{tr} E[Z_t \beta \beta' Z_t A S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A]| \\ &= \frac{1}{T^2} 2 |\operatorname{tr} E[\mathring{\Sigma} Z_t A S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A Z_t]| \\ &= \frac{1}{T^2} 2 |\operatorname{tr} E[\mathring{\Sigma}^{1/2} Z_t A S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A Z_t \mathring{\Sigma}^{1/2}]| \\ &\leq \|\Sigma_{\varepsilon,t}\| \frac{1}{T^2} 2 |\operatorname{tr} E[\mathring{\Sigma}^{1/2} Z_t A S'_{t-1} S_{t-1} A Z_t \mathring{\Sigma}^{1/2}]| \\ &= \|\Sigma_{\varepsilon,t}\| \frac{1}{T^2} 2 |\operatorname{tr} E[\mathring{\Sigma}^{1/2} Z_t A Z_t A Z_t \mathring{\Sigma}^{1/2}]| \\ &= \|\Sigma_{\varepsilon,t}\| \frac{1}{T^2} 2 \operatorname{tr} E[Z_t A Z_t A Z_t \mathring{\Sigma}] \\ &= \|\Sigma_{\varepsilon,t}\| \frac{1}{T^2} 2 \operatorname{tr} E[Y \Psi^{1/2} A \Psi^{1/2} Y \Psi^{1/2} A \Psi^{1/2} Y \Psi^{1/2} \mathring{\Sigma} \Psi^{1/2}]. \end{aligned} \quad (187)$$

Thus, defining $C = \Psi^{1/2} \mathring{\Sigma} \Psi^{1/2}$ and replacing A with $\Psi^{1/2} A \Psi^{1/2}$, we may assume (for the purpose of establishing asymptotic bounds) that $\Psi = I$. Then,

$$\operatorname{tr} E[Y A Y A Y C] = \sum_{i_1, i_2, i_3, i_4, i_5, i_6} E[Y_{i_1, i_2} A_{i_2, i_3} Y_{i_3, i_4} A_{i_4, i_5} Y_{i_5, i_6} C_{i_6, i_1}] \quad (188)$$

As above, we may assume that $C = \gamma\gamma'$ with $\|\gamma\|$ bounded. First, we deal with $i_1 = i_6$ terms. This gives

$$\sum_{i_1} \sum_{i_2, i_3, i_4, i_5} E[Y_{i_1, i_2} A_{i_2, i_3} Y_{i_3, i_4} A_{i_4, i_5} Y_{i_5, i_1}] \gamma_{i_1}^2 \quad (189)$$

As above, we have $O(P)$ terms that are $O(1)$ (because $\text{tr}(\Sigma) = 1$) and $O(P^2)$ non-zero terms that are $O(\text{tr}(\Sigma^2))$.

Now, for $i_1 \neq i_6$, we get that the indices $i_1, i_2, i_3, i_4, i_5, i_6$ need to be coupled to get non-zero terms. We consider these cases separately. First, if $i_2 = i_1$, then (i_3, i_4) need to be coupled with (i_5, i_6) and hence, for each (i_1, i_6) , we only have one free index, for example,

$$\sum_{i_1, i_3, i_6} E[Y_{i_1, i_1} A_{i_2, i_3} Y_{i_3, i_6}^2 A_{i_4, i_5} \gamma_{i_6} \gamma_{i_1}] \quad (190)$$

Only the term with $i_3 = i_6$ will be $O(1)$, but all other terms will be $O(\text{tr}(\Sigma^2))$ by Lemma 16. A similar argument applies to all other couplings. Thus, we end up with a bound

$$\sum_{i_1, i_6} |\gamma_{i_6} \gamma_{i_1}| K(1 + P \text{tr}(\Sigma^2)) = \left(\sum_i |\gamma_i| \right)^2 K(1 + P \text{tr}(\Sigma^2)). \quad (191)$$

Now, by the Jensen inequality

$$\left(\sum_i |\gamma_i| \right)^2 \leq P \|\gamma\|^2, \quad (192)$$

and we get the final bound

$$\frac{1}{T^2} 2 |\text{tr} E[Z_t \beta \beta' Z_t A S'_{t-1} \varepsilon_t \varepsilon'_t S_{t-1} A]| \leq T^{-2} K P (1 + P \text{tr}(\Sigma^2)) = o(1). \quad (193)$$

for some constant $K > 0$.

G.3 Third Term in (165)

$$\begin{aligned}
& \frac{1}{T^2} 2 \operatorname{tr} E[S'_{t-1} \varepsilon_t \beta' Z_t A S'_{t-1} \varepsilon_t \beta' Z_t A] \\
&= \frac{1}{T^2} 2 \operatorname{tr} E[S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A Z_t \beta \beta' Z_t A] \\
&= \frac{1}{T^2} 2 \operatorname{tr} E[S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A Z_t \Sigma_{F,t}^* Z_t A] \\
&= \frac{1}{T^2} 2 \operatorname{tr} E[(\Sigma_{F,t}^*)^{1/2} Z_t A S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A Z_t (\Sigma_{F,t}^*)^{1/2}] \\
&\leq \|\Sigma_{\varepsilon,t}\| \frac{1}{T^2} 2 \operatorname{tr} E[(\Sigma_{F,t}^*)^{1/2} Z_t A Z_t A Z_t (\Sigma_{F,t}^*)^{1/2}]
\end{aligned} \tag{194}$$

which coincides with the bound in (187), and hence, the term is negligible by the previous subsection.

G.4 Fourth Term in (165)

We have

$$\begin{aligned}
& \frac{1}{T^2} 2 \operatorname{tr} E[S'_{t-1} \varepsilon_t \beta' Z_t A Z_t \beta \varepsilon'_t S_{t-1} A] \\
&= \frac{1}{T^2} 2 \operatorname{tr} E[\beta' Z_t A Z_t \beta \operatorname{tr}(\Sigma_{\varepsilon,t} S_{t-1} A S'_{t-1})] \\
&= \frac{1}{T^2} 2 \operatorname{tr} E[\operatorname{tr}(\Sigma_{F,t}^* Z_t A Z_t) \operatorname{tr}(\Sigma_{\varepsilon,t} S_{t-1} A S'_{t-1})].
\end{aligned} \tag{195}$$

Both traces are positive because $\operatorname{tr}(AB) > 0$ when both A, B are positive definite. Furthermore,

$$\operatorname{tr}(\Sigma_{\varepsilon,t} S_{t-1} A S'_{t-1}) \leq \|\Sigma_{\varepsilon,t}\| \operatorname{tr}(S_{t-1} A S'_{t-1}) \leq K \operatorname{tr}(Z_t A). \tag{196}$$

As above, replacing A with $\Psi^{1/2} A \Psi^{1/2}$, we can assume that $\Psi = I$ for the sake of deriving bounds.

Using the Cauchy-Schwarz inequality, it is possible to show that the term is negligible, of the order $(\operatorname{tr}(\Sigma^2))^{1/2}$. Here, we instead perform the direct calculation show that, in fact,

it is bounded from above by $\text{tr}(\Sigma^2)$. We need to bound

$$E[|\text{tr}(CYAY) \text{tr}(YA)|] \quad (197)$$

and, as above, we can assume $C = \gamma\gamma'$. Similarly, using $|A| \leq \|A\|I$, we can assume that $A = I$.

$$E[\text{tr}(CY^2) \text{tr}(Y)] = E\left[\sum_{i_1, i_2, i_3} \gamma_{i_1} \gamma_{i_2} Y_{i_2, i_3} Y_{i_3, i_1} \sum_{i_5} Y_{i_5, i_5}\right] \quad (198)$$

The terms with $i_1 = i_2$ give

$$E\left[\sum_{i_1, i_3} \gamma_{i_1}^2 Y_{i_1, i_3}^2 \sum_{i_5} Y_{i_5, i_5}\right] = (O(P) + (P^2) \text{tr}(\Sigma^2)) \|\gamma\|^2 \quad (199)$$

The terms with $i \neq i_3$ have zero expectation.

G.5 Fifth Term in (165) (The Only Non-Negligible One)

By Lemma 8,

$$\begin{aligned} & \frac{1}{T^2} \text{tr} E[S'_{t-1} \varepsilon_t \varepsilon'_t S_{t-1} A S'_{t-1} \varepsilon_t \varepsilon'_t S_{t-1} A] \\ &= \frac{1}{T^2} \text{tr} E[S'_{t-1} \left(\Sigma_{\varepsilon, t}^{1/2} (2\tilde{B} + \text{diag}((\kappa_\varepsilon - 3) \text{diag}(\tilde{B}) + \text{tr}(\tilde{B}))) \Sigma_{\varepsilon, t}^{1/2} \right) S_{t-1} A] \end{aligned} \quad (200)$$

where

$$\tilde{B} = \Sigma_{\varepsilon, t}^{1/2} S_{t-1} A S'_{t-1} \Sigma_{\varepsilon, t}^{1/2}. \quad (201)$$

First, we show that the excess kurtosis terms are negligible. We have (since the diag operation preserves positive definite order) that

$$\begin{aligned}
& |\operatorname{tr} E[S'_{t-1} \Sigma_{\varepsilon,t}^{1/2} \operatorname{diag}(\tilde{B}) \Sigma_{\varepsilon,t}^{1/2} S_{t-1} A]| \\
& \leq \|A\|^2 \operatorname{tr} E[S'_{t-1} \Sigma_{\varepsilon,t}^{1/2} \operatorname{diag}(\Sigma_{\varepsilon,t}^{1/2} S_{t-1} S'_{t-1} \Sigma_{\varepsilon,t}^{1/2}) \Sigma_{\varepsilon,t}^{1/2} S_{t-1}] \\
& \leq \|A\|^2 \|\Psi\|^2 \operatorname{tr} E[X'_{t-1} \Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} \operatorname{diag}(\Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} X_{t-1} X'_{t-1} \Sigma_{\varepsilon,t}^{1/2}) \Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} X_{t-1}] \\
& = \|A\|^2 \|\Psi\|^2 \operatorname{tr} E[\operatorname{diag}(\Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} X_{t-1} X'_{t-1} \Sigma_{\varepsilon,t}^{1/2}) \Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} X_{t-1} X'_{t-1} \Sigma_{\varepsilon,t}^{1/2}] \\
& = \|A\|^2 \|\Psi\|^2 E[\sum_i (\Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} X_{t-1} X'_{t-1} \Sigma_{\varepsilon,t}^{1/2})_{i,i} (\Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2} X_{t-1} X'_{t-1} \Sigma_{\varepsilon,t}^{1/2})_{i,i}] \quad (202) \\
& \leq \|A\|^2 \|\Psi\|^2 E[\sum_i ((\tilde{\Sigma}^{1/2} X_{t-1} X'_{t-1} \tilde{\Sigma}^{1/2})_{i,i})^2] \\
& \leq \|A\|^2 \|\Psi\|^2 E[\operatorname{tr}((\tilde{\Sigma}^{1/2} X_{t-1} X'_{t-1} \tilde{\Sigma}^{1/2})^2)] \\
& = \|A\|^2 \|\Psi\|^2 E[\operatorname{tr}(\tilde{Y}^2)]
\end{aligned}$$

where $\tilde{\Sigma} = (\Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2}) (\Sigma_{\varepsilon,t}^{1/2} \Sigma_{\varepsilon,t}^{1/2})'$ and $\tilde{Y} = X' \tilde{\Sigma} X$. Then,

$$E[\operatorname{tr}(Y^2)] = E[\sum_{i,j} Y_{i,j}^2] = O(P) + \operatorname{tr}(\Sigma^2) O(P^2), \quad (203)$$

so this term is negligible, where we have used that $\operatorname{tr}(\tilde{\Sigma}) \leq \operatorname{tr}(\Sigma) \|\tilde{\Sigma}_\varepsilon\|$.

Recall that

$$Z_t = S'_{t-1} S_{t-1} = \Psi^{1/2} X'_{t-1} \Sigma X_{t-1} \Psi^{1/2}$$

Then, the only remaining non-negligible terms in (200) are

$$\begin{aligned}
& \frac{1}{T^2} \operatorname{tr} E[S'_{t-1} \left(\Sigma_{\varepsilon,t}^{1/2} (2\tilde{B} + \operatorname{tr}(\tilde{B})) \Sigma_{\varepsilon,t}^{1/2} \right) S_{t-1} A] \\
&= \frac{1}{T^2} 2 \operatorname{tr} E[S'_{t-1} \Sigma_{\varepsilon,t} S'_{t-1} A S_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A] \\
&+ E[\operatorname{tr}(S_{t-1} A S'_{t-1} \Sigma_{\varepsilon,t}) \operatorname{tr}(S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1} A)] \\
&= T^{-2} E[2 \operatorname{tr}(A \tilde{Z}_t A \tilde{Z}_t) + (\operatorname{tr}(A \tilde{Z}_t))^2] \\
&= T^{-2} E[2 \operatorname{tr}(A \tilde{Z}_t A \tilde{Z}_t) + (\operatorname{tr}(A \tilde{Z}_t))^2]
\end{aligned} \tag{204}$$

where we have defined

$$\tilde{Z}_t = S'_{t-1} \Sigma_{\varepsilon,t} S_{t-1}. \tag{205}$$

By Lemma 7, we have

$$\begin{aligned}
E[S'_t S_t A S'_t S_t] &= ((\operatorname{tr} \tilde{\Sigma})^2 + \operatorname{tr}(\tilde{\Sigma}^2)) \Psi A \Psi + \operatorname{tr}(\tilde{\Sigma}^2) \operatorname{tr}(\Psi A) \Psi \\
&+ \operatorname{tr}(\tilde{\Sigma} \circ \tilde{\Sigma}) \Psi^{1/2} (\kappa - 3) \operatorname{diag}(\Psi^{1/2} A \Psi^{1/2}) \Psi^{1/2}
\end{aligned} \tag{206}$$

and, hence,

$$\begin{aligned}
& T^{-2} E[2 \operatorname{tr}(A S'_{t-1} S_{t-1} A S'_{t-1} S_{t-1})] \\
&= 2T^{-2} E[\operatorname{tr}(A \left(((\operatorname{tr} \tilde{\Sigma})^2 + \operatorname{tr}(\tilde{\Sigma}^2)) \Psi A \Psi + \operatorname{tr}(\tilde{\Sigma}^2) \operatorname{tr}(\Psi A) \Psi \right. \\
&\quad \left. + \operatorname{tr}(\tilde{\Sigma} \circ \tilde{\Sigma}) \Psi^{1/2} (\kappa - 3) \operatorname{diag}(\Psi^{1/2} A \Psi^{1/2}) \Psi^{1/2} \right))] .
\end{aligned} \tag{207}$$

The matrix in the brackets is uniformly bounded and, hence, the trace is $O(P)$, so that $T^{-2}O(P)$ is negligible.

We now use Lemma 13 with $\tilde{A} = \Psi^{1/2} A \Psi^{1/2}$ to get

$$\begin{aligned}
E[(\text{tr}(A\tilde{Z}_t))^2] &= E[(\text{tr}(A\Psi^{1/2}X'_{t-1}\tilde{\Sigma}X_{t-1}\Psi^{1/2}))^2] \\
&= E[(\text{tr}(\tilde{A}X'_{t-1}\tilde{\Sigma}X_{t-1}))^2] \\
&= E[(\sum_{i_1, i_2, i_3, i_4} \tilde{A}_{i_1, i_2} X_{i_3, i_2} \tilde{\Sigma}_{i_3, i_4} X_{i_4, i_1})^2] \\
&= 2\text{tr}(\tilde{A}^2)\text{tr}(\tilde{\Sigma}^2) + \text{tr}(\tilde{A})^2\text{tr}(\tilde{\Sigma})^2 + (\mu_4 - 3)\text{tr}(\tilde{A} \circ \tilde{A})\text{tr}(\tilde{\Sigma} \circ \tilde{\Sigma}).
\end{aligned} \tag{208}$$

Since $\tilde{A}$ is bounded and $\text{tr}(\Sigma) = 1$, $\text{tr}(\tilde{\Sigma}^2) = o(1)$, we have that the only term that is non-negligible is $(T^{-1}\text{tr}(\tilde{\Sigma})\text{tr}(\Psi A))^2$. The proof of Lemma 11 is complete. $\square$

H Computing Expectations Involving Powers of Ψ

We will use the notation

$$B_T = \frac{1}{T} \sum_{t=1}^T F_t F'_t \tag{209}$$

and

$$B_{T,t} = \frac{1}{T} \sum_{\tau=1}^T F_\tau F'_\tau - T^{-1} F_t F'_t. \tag{210}$$

Let

$$\hat{\lambda}(z) = (zI + B_T)^{-1} \frac{1}{T} \sum_{t=1}^T F_t \tag{211}$$

while

$$\hat{R}_{T+1}^M(z) = \hat{\lambda}(z)' F_{t+1} = (S_t \hat{\lambda}(z))' R_{t+1}. \tag{212}$$

We will also need the following Lemmas from [KMZ](#).

Lemma 17 *We have*

$$P^{-1} \text{tr}(A_1(zI + B_T)^{-1}A_2) - P^{-1} \text{tr} E[A_1(zI + B_T)^{-1}A_2] \rightarrow 0$$

almost surely for any bounded A_1, A_2 that are independent of $F_t, t \leq T$.

Lemma 18 *Under Assumption [4](#), there exists a function $\xi(z; c)$ such that*

$$\frac{1}{T} \text{tr}((zI + B_T)^{-1}\Psi\sigma_*) \rightarrow \xi(z; c) \quad (213)$$

almost surely and

$$\frac{1}{T} F'_t(zI + B_{T,t})^{-1}F_t \rightarrow \xi(z; c), \quad (214)$$

in probability, where

$$\frac{c^{-1}\xi(z; c)}{1 + \xi(z; c)} = 1 - m(-z; c)z \quad (215)$$

Proof. First, Lemma [11](#) implies that

$$\frac{1}{T} F'_t(zI + B_{T,t})^{-1}F_t - \frac{1}{T} \text{tr}((zI + B_{T,t})^{-1}E[F_t F'_t]) \rightarrow 0.$$

in probability. Next Lemma [17](#) applied to our setting implies that for any bounded matrix Q_T independent of $B_{T,t}$ we have

$$\frac{1}{T} \text{tr}((zI + B_{T,t})^{-1}Q_T) - \frac{1}{T} E[\text{tr}((zI + B_{T,t})^{-1}Q_T)] \rightarrow 0$$

almost surely. At the same time, by Lemma 9,

$$\begin{aligned} E[F_t F_t'] &= ((\text{tr } \Sigma)^2 + \text{tr}(\Sigma^2)) \Psi \Sigma_F \Psi \\ &+ \text{tr}(\Sigma^2)(\kappa - 3) \Psi^{1/2} \text{diag}(\Psi^{1/2} \Sigma_F \Psi^{1/2}) \Psi^{1/2} + \Psi \left(\text{tr}(\Sigma \Sigma_\varepsilon) + \text{tr}(\Psi \Sigma_F) \text{tr}(\Sigma^2) \right) \end{aligned} \quad (216)$$

We have

$$\frac{1}{T} \text{tr}((zI + B_{T,t})^{-1} (\text{tr } \Sigma)^2 \Psi \Sigma_{F,t} \Psi) = O(1/T) \quad (217)$$

The same argument applies to the second term because the trace of

$$\text{tr}(\Sigma^2)(\kappa - 3) \Psi^{1/2} \text{diag}(\Psi^{1/2} \Sigma_F \Psi^{1/2}) \Psi^{1/2}$$

is also uniformly bounded. Thus, we get

$$\begin{aligned} \frac{1}{T} F_t' (zI + B_{T,t})^{-1} F_t &\sim \frac{1}{T} \text{tr}((zI + B_{T,t})^{-1} E[F_t F_t']) \\ &\sim T^{-1} \text{tr}[(zI + B_{T,t})^{-1} \Psi \sigma_*] \rightarrow \xi(z; c). \end{aligned} \quad (218)$$

Now, we have

$$\begin{aligned} 1 &= P^{-1} \text{tr } E[(zI + B_T)^{-1} (zI + B_T)] \\ &= zm(-z; c) + \frac{1}{P} \text{tr } \frac{1}{T} \sum_t E[(zI + B_T)^{-1} F_t F_t'] \\ &= zm(-z; c) + \frac{1}{P} \text{tr } E[(zI + B_T)^{-1} F_t F_t'] \end{aligned} \quad (219)$$

where we have used symmetry across t in the last step. Using the Sherman-Morrison formula, we get

$$\frac{1}{T} \text{tr } E[(zI + B_T)^{-1} F_t F_t'] = E \left[\frac{\frac{1}{T} F_t' (zI + B_{T,t})^{-1} F_t}{1 + \frac{1}{T} F_t' (zI + B_{T,t})^{-1} F_t} \right],$$

where

$$B_{T,t} = \frac{1}{T} \sum_{\tau \neq t} F_{\tau} F'_{\tau}.$$

Furthermore, since all functions involved are uniformly bounded, a standard argument implies that we can replace

$$\frac{1}{T} F'_t (zI + B_{T,t})^{-1} F_t$$

with

$$\xi(z; c)$$

by (218).⁴⁷

□

Lemma 19 *Let X_P be a sequence of positive semi-definite matrices with $\text{tr}(X_P) \leq K$. Then,*

$$\lim_{M \rightarrow \infty} \left(\frac{1}{P} \text{tr}(zI + A_P + X_P)^{-1} - \frac{1}{P} \text{tr}(zI + A_P)^{-1} \right) = 0$$

for any positive semi-definite matrices A_P .

Proof. We have

$$\frac{1}{P} \text{tr}(zI + A_P + X_P)^{-1} - \frac{1}{P} \text{tr}(zI + A_P)^{-1} = \frac{1}{P} \text{tr}((zI + A_P + X_P)^{-1} - (zI + A_P)^{-1})$$

and the claim follows because

$$\frac{1}{P} \text{tr}((zI + A_P + X_P)^{-1} - (zI + A_P)^{-1}) = -\frac{1}{P} \text{tr}((zI + A_P + X_P)^{-1} X_P (zI + A_P)^{-1})$$

⁴⁷Indeed, $E[\frac{Y_T}{1+Y_T} - \frac{Z_T}{1+Z_T}] = \frac{Y_T - Z_T}{(1+Y_T)(1+Z_T)}$ for any random variables Y_T, Z_T . If $Y_T, Z_T \geq 0$ then $\frac{|Y_T - Z_T|}{(1+Y_T)(1+Z_T)} \leq 1$ and hence convergence $Y_T - Z_T \rightarrow 0$ in probability implies convergence of expectations.

and

$$\begin{aligned} \text{tr}((zI + A_P + X_P)^{-1}X_P(zI + A_P)^{-1}) &= \text{tr}(X_P(zI + A_P)^{-1}(zI + A_P + X_P)^{-1}) \\ &\leq \text{tr}(X_P)\|(zI + A_P)^{-1}(zI + A_P + X_P)^{-1}\| \leq Kz^{-2} \end{aligned} \quad (220)$$

Thus, the difference is bounded in absolute value by Kz^{-2}/M . $\square$

We will, for simplicity, always assume that $\sigma_* = 1$. By Lemma 9,

$$E[FF'] \approx \Psi\Sigma_F\Psi + \sigma_*\Psi, \quad E[F] = \Psi\nu. \quad (221)$$

Furthermore, all terms involving interactions of the component $\Psi\Sigma_F\Psi$ with ν will be negligible because Σ_F has bounded trace and, hence, so do products of Σ_F with any bounded matrices. In this case, the infeasible Markowitz portfolio is $\lambda = E[FF']^{-1}E[F] \approx (\Psi\Sigma_F\Psi + \Psi)^{-1}\nu \approx \Psi^{-1}\Psi\nu = \nu$. Recall also that, by Assumption 1, $\nu = P^{-1/2}\Sigma_\nu\hat{\nu}$, $E[\hat{\nu}] = 0$, $E[\hat{\nu}\hat{\nu}'] = I$, and, hence, we are in a position to apply the arguments in Lemma 3 and get

Lemma 20 *We have*

$$\nu' A \nu - P^{-1} \text{tr}(A\Sigma_\nu) \rightarrow 0 \quad (222)$$

in L_2 and

$$\lambda' A \lambda - P^{-1} \text{tr}(A\Sigma_\nu) \rightarrow 0 \quad (223)$$

and

$$\lambda' A \nu - P^{-1} \text{tr}(A\Sigma_\nu) \rightarrow 0 \quad (224)$$

in L_2 .

Lemma 21 *Let*

$$\psi_{*,k} = \lim P^{-1} \text{tr}(\lambda' \Psi^k \lambda) \quad (225)$$

and

$$\Gamma_{k,l,T}(z) \equiv \lambda' E[\Psi^k (zI + B_T)^{-1} \Psi^l] \lambda. \quad (226)$$

We have

$$\psi_{*,k+\ell} = z \Gamma_{k,\ell,T}(z) + \left(\psi_{*,k+1} \Gamma_{1,\ell,T}(z) + \sigma_* \Gamma_{k+1,\ell,T} \right) (1 + \xi(z; c))^{-1} + O(P^{-1/2} + \text{tr}(\Sigma^2)), \quad (227)$$

where $E[|O(P^{-1/2})|] \leq K \|\Psi\|^{k+\ell} (P^{-1/2} + \text{tr}(\Sigma^2))$ for some $K > 0$.

Proof of Lemma 21. Using the Sherman-Morrison formula and Lemma 18, we get

$$F'_t(zI + B_T)^{-1} = F'_t(zI + B_{T,t})^{-1} (1 + \hat{\xi}_t)^{-1},$$

where

$$\hat{\xi}_t = (T)^{-1} F'_t(zI + B_{T,t})^{-1} F_t \approx \xi(z; c), \quad (228)$$

where $\approx$ means that the difference between the two quantities has an L_2 -norm less than $K P^{-1/2}$ for some $K > 0$. We also have

$$\begin{aligned} E[F_t F'_t] &= ((\text{tr} \Sigma)^2 + \text{tr}(\Sigma^2)) \Psi \Sigma_F^* \Psi \\ &+ \text{tr}(\Sigma \circ \Sigma) (\kappa - 3) \Psi^{1/2} \text{diag}(\Psi^{1/2} \Sigma_F^* \Psi^{1/2}) \Psi^{1/2} + \Psi \left(\text{tr}(\Sigma \Sigma_\varepsilon) + \text{tr}(\Psi \Sigma_F^*) \text{tr}(\Sigma^2) \right) \\ &= \hat{\Sigma}_F + \Psi \Sigma_F^* \Psi + \sigma_* \Psi, \end{aligned} \quad (229)$$

where $\|\widehat{\Sigma}_F\| = o(1)$, and

$$\Sigma_F^* = \nu\nu' + \Sigma_F. \quad (230)$$

We will need the following important observations:

Lemma 22 *We have*

$$\lambda' A_P Q_P \lambda \rightarrow 0 \quad (231)$$

in probability, for any uniformly bounded Q_P (even if they correlate with λ) and any A_P with a uniformly bounded trace norm, such that A_P is independent of λ .

Proof of Lemma 22. We have

$$\begin{aligned} \lambda' A_P Q_P \lambda &= \text{tr}(\lambda \lambda' A_P Q_P) \\ &\leq \|\lambda \lambda' A_P Q_P\|_1 \leq \|Q_P\|_\infty \|\lambda \lambda' A_P\|_1 \\ &= \|Q_P\|_\infty \text{tr}((\lambda \lambda' A_P A_P' \lambda \lambda')^{1/2}) = \|Q_P\|_\infty (\lambda' A_P A_P' \lambda)^{1/2} \text{tr}((\lambda \lambda')^{1/2}) = (\lambda' A_P A_P' \lambda)^{1/2} \|\lambda\| \\ &= (\text{tr}(A_P A_P' \lambda \lambda'))^{1/2} \|\lambda\| \rightarrow (P^{-1} \text{tr}(\Sigma_\nu))^{1/2} (P^{-1} \text{tr}(A_P A_P' \Sigma_\nu))^{1/2} \\ &\leq (P^{-1} \text{tr}(\Sigma_\nu))^{1/2} \|\Sigma_\nu\|^{1/2} (P^{-1} \text{tr}(A_P A_P'))^{1/2} \rightarrow 0 \end{aligned} \quad (232)$$

The proof of Lemma 22 is complete. □

In the next calculation, we use Lemma 20 to interchangeably use λ and ν . Thus, for any A_P

with bounded trace norm, we get with $\Sigma_F^* = \Sigma_F + \nu\nu'$ that

$$\begin{aligned}
\psi_{*,k+\ell} &= P^{-1} \text{tr}(\Psi^{k+\ell} \Sigma_\nu) \approx \lambda' \Psi^{k+\ell} \lambda = \lambda' E[\Psi^k (zI + B_T)(zI + B_T)^{-1} \Psi^\ell] \lambda \\
&= z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k B_T (zI + B_T)^{-1} \Psi^\ell] \lambda \\
&\stackrel{\text{symmetry over } t}{=} z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k F_t F_t' (zI + B_T)^{-1} \Psi^\ell] \lambda \\
&\stackrel{(98)}{=} z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k F_t F_t' (zI + B_{T,t})^{-1} (1 + \hat{\xi}_t)^{-1} \Psi^\ell] \lambda \\
&\stackrel{\text{Lemma 18}}{\sim} z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k F_t F_t' (zI + B_{T,t})^{-1} \Psi^\ell] \lambda (1 + \xi(z; c))^{-1} \\
&\stackrel{(229)}{\sim} z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k (\hat{\Sigma}_F + \Psi \Sigma_F^* \Psi + \sigma_* \Psi) (zI + B_{T,t})^{-1} \Psi^\ell] \lambda (1 + \xi(z; c))^{-1} \\
&\stackrel{\|\hat{\Sigma}_F\|=o(1)}{\sim} z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k (\Psi(\Sigma_F + \nu\nu') \Psi + \sigma_* \Psi) (zI + B_T)^{-1} \Psi^\ell] \lambda (1 + \xi(z; c))^{-1} \\
&\stackrel{(231)}{\sim} z\Gamma_{k,\ell,T}(z) + \lambda' E[\Psi^k (\nu\nu' + \sigma_* \Psi) (zI + B_T)^{-1} \Psi^\ell] \lambda (1 + \xi(z; c))^{-1} \\
&\stackrel{\text{Lemma 20}}{\sim} z\Gamma_{k,\ell,T}(z) + \lambda' \Psi^{k+1} \lambda E[\nu' (zI + B_T)^{-1} \Psi^\ell] \lambda (1 + \xi(z; c))^{-1} \\
&+ \nu' E[\Psi^{k+1} \sigma_* (zI + B_T)^{-1} \Psi^\ell] \lambda (1 + \xi(z; c))^{-1} \\
&\sim z\Gamma_{k,\ell,T}(z) + \left(\psi_{*,k+1} \Gamma_{1,\ell,T}(z) + \sigma_* \Gamma_{k+1,\ell,T} \right) (1 + \xi(z; c))^{-1}
\end{aligned} \tag{233}$$

□

Lemma 23 *The limits*

$$\Gamma_{k,l}(z) = \lim_{P,T,N \rightarrow \infty, P/T \rightarrow c} \Gamma_{k,l,T}(z) \tag{234}$$

exist in probability (as they depend on the realization of ν , they are potentially random).

They can be computed as follows. Let

$$\delta(z) = -\sigma_* z^{-1}(1 + \xi(z; c))^{-1}. \quad (235)$$

Then,

$$\Gamma_{1,l}(z) = \lim \frac{z^{-1}P^{-1} \operatorname{tr}(\Psi^{1+l}(I - \Psi\delta(z))^{-1}\Sigma_\nu)}{1 - \delta(z)P^{-1} \operatorname{tr}(\Psi^2(I - \Psi\delta(z))^{-1}\Sigma_\nu)} \quad (236)$$

and

$$\Gamma_{k,\ell} = \lim \{z^{-1}P^{-1} \operatorname{tr}(\Psi^{k+\ell}(I - \Psi\delta(z))^{-1}\Sigma_\nu) - z^{-1}P^{-1} \operatorname{tr}(\Psi^{k+1}(I - \Psi\delta(z))^{-1}\Sigma_\nu)\Gamma_{1,\ell}(1 + \xi(z; c))^{-1}\} \quad (237)$$

Proof. Passing to a subsequence and using the standard diagonal argument, we can assume that the limits exist (but might depend on the subsequence). Lemma 21 implies that these limits satisfy

$$\Gamma_{k,\ell} = a_{k+1} + \delta \Gamma_{k+1,\ell}, \quad (238)$$

where

$$a_{k+1,\ell} = z^{-1}(\psi_{*,k+\ell} - \psi_{*,k+1}\Gamma_{1,\ell}(1 + \xi(z; c))^{-1}), \quad \delta(z) = -\sigma_* z^{-1}(1 + \xi(z; c))^{-1}. \quad (239)$$

Let us pick $z > \max(1, \|\Psi\|)$ sufficiently large, so that $\sigma_* z^{-1}(1 + \xi(z; c))^{-1} < 1$ and⁴⁸

$$|\delta^k \Gamma_{k,\ell}(z)| \leq z^{-k+1} \|\lambda\|^2 \|\Psi\|^{k+\ell} \rightarrow_{k \rightarrow \infty} 0. \quad (240)$$

⁴⁸This uniform exponential decay also implies that the infinite sum of the limits equals the limit of the infinite sum, as we pass to the $P \rightarrow \infty$ limit.

Then, iterating forward, we get

$$\Gamma_{k,\ell} = \sum_{\tau=0}^{\infty} a_{k+\tau+1,\ell} \delta^\tau. \quad (241)$$

Now,

$$a_{k+\tau+1,\ell} = z^{-1}(\psi_{*,k+\tau+\ell} - \psi_{*,k+\tau+1} \Gamma_{1,\ell} (1 + \xi(z; c))^{-1}), \quad \delta(z) = -\sigma_* z^{-1} (1 + \xi(z; c))^{-1}. \quad (242)$$

$$\begin{aligned} \Gamma_{1,\ell} &= \sum_{\tau=0}^{\infty} a_{\tau+2,\ell} \delta^\tau \\ &= \sum_{\tau=0}^{\infty} z^{-1} (\psi_{*,1+\tau+\ell} - \psi_{*,1+\tau+1} \Gamma_{1,\ell} (1 + \xi(z; c))^{-1}) \delta^\tau \\ &= \sum_{\tau=0}^{\infty} (z^{-1} (P^{-1} \text{tr}(\Psi^{\tau+\ell+1} \Sigma_\nu) - P^{-1} \text{tr}(\Psi^{\tau+2} \Sigma_\nu) \Gamma_{1,\ell} (1 + \xi(z; c))^{-1})) \delta^\tau \\ &= z^{-1} P^{-1} \text{tr}(\Psi^{1+\ell} (I - \Psi \delta(z))^{-1} \Sigma_\nu) - z^{-1} P^{-1} \text{tr}(\Psi^2 (I - \Psi \delta(z))^{-1} \Sigma_\nu) \Gamma_{1,\ell} (1 + \xi(z; c))^{-1}, \end{aligned} \quad (243)$$

implying that

$$\Gamma_{1,\ell} = \lim \frac{z^{-1} P^{-1} \text{tr}(\Psi^{1+\ell} (I - \Psi \delta(z))^{-1} \Sigma_\nu)}{1 - \delta(z) P^{-1} \text{tr}(\Psi^2 (I - \Psi \delta(z))^{-1} \Sigma_\nu)} \quad (244)$$

Then, the same argument implies

$$\Gamma_{k,\ell} = \lim z^{-1} P^{-1} \text{tr}(\Psi^{k+\ell} (I - \Psi \delta(z))^{-1} \Sigma_\nu) - z^{-1} P^{-1} \text{tr}(\Psi^{k+1} (I - \Psi \delta(z))^{-1} \Sigma_\nu) \Gamma_{1,\ell} (1 + \xi(z; c))^{-1} \quad (245)$$

Furthermore,

$$\delta(z) = -\sigma_* z^{-1} (1 + \xi(z; c))^{-1}. \quad (246)$$

Thus, the limits are independent of the subsequence chosen. Furthermore, a standard argument based on the Vitali Theorem (see (Kelly et al., 2024b)) implies that the limits are analytic for z with $\Re z > 0$. Thus, if the identities for the limits hold in the open subset $\Re z > \max(1, \|\Psi\|)$, they hold in the whole half-plane by the uniqueness theorem for analytic functions. $\square$

I Expectation of $\hat{R}^M$

The goal of this section is to compute the asymptotics of out-of-sample expected returns of the feasible efficient portfolio.

Proposition 6 *We have*

$$\lim_{T, P \rightarrow \infty, P/T \rightarrow c, N \rightarrow \infty} E[\hat{R}_{t+1}^M(z)] = \frac{\Gamma_{1,1}(z)}{1 + \xi(z; c)}, \quad (247)$$

where

$$\Gamma_{1,1}(z) = \lim_{T, P \rightarrow \infty, P/T \rightarrow c, N \rightarrow \infty} \nu' E[\Psi(zI + B_T)^{-1} \Psi] \nu. \quad (248)$$

Proof of Proposition 6. We start by computing

$$E[F_{t+1}] = E[S_t' R_{t+1}] = E[S_t' (S_t \tilde{F}_{t+1} + \varepsilon_{t+1})] = \text{tr}(\Sigma) \Psi \nu = \Psi \nu \quad (249)$$

by the assumed normalization of $\text{tr}(\Sigma) = 1$. Therefore, by (212), we have

$$E[\hat{R}_{t+1}^M(z)] = E[\hat{\lambda}(z)' F_{t+1}] = E\left[\frac{1}{T} \sum_t F_t'(zI + B_T)^{-1}\right] \Psi \nu \quad (250)$$

where we have used the normalization $\text{tr} \Sigma = 1$. Now, by the interchangeability of F_t across t and the Sherman-Morrison formula, we have

$$\begin{aligned} & E\left[\frac{1}{T} \sum_t F_t'(zI + B_T)^{-1}\right] \Psi \nu \\ &= E[F_t'(zI + B_T)^{-1} \Psi] \nu = E[F_t'(zI + B_{T,t})^{-1} \frac{1}{1 + (T)^{-1} F_t'(zI + B_{T,t})^{-1} F_t} \Psi] \nu, \end{aligned} \quad (251)$$

where

$$B_{T,t} = \frac{1}{T} \sum_{\tau \neq t} F_\tau F_\tau'.$$

By Lemma 18,

$$(T)^{-1} F_t'(zI + B_{T,t})^{-1} F_t \rightarrow \xi(z; c)$$

is probability and therefore

$$E[F_t'(zI + B_{T,t})^{-1} \frac{1}{1 + (T)^{-1} F_t'(zI + B_{T,t})^{-1} F_t} \Psi] \nu \sim \frac{E[F_t'(zI + B_{T,t})^{-1} \Psi \nu]}{1 + \xi(z; c)}, \quad (252)$$

(taking $1 + \xi$ out of expectation is justified by Lemma 4), whereas $E[F_t] = \Psi \nu$ implies

$$E[F_t'(zI + B_{T,t})^{-1} \Psi \nu] = \nu' E[\Psi(zI + B_{T,t})^{-1} \Psi \nu] \sim \Gamma_{1,1}(z) \quad (253)$$

by Lemma 23. Recall that, by (46),

$$\mathcal{E}(z) = E[R^{infeas}(z)] = E[F]'(zI + E[FF'])^{-1}E[F] = \frac{A(z)}{1 + A(z)}, \quad (254)$$

and

$$\mathcal{V}(z) = \frac{(zA(z))'}{(1 + A(z))^2}. \quad (255)$$

where (assuming, as above, the normalization $\sigma_* = 1$)

$$\lim_{P \rightarrow \infty} A(z) = \lim \lambda' \Psi(zI + \Psi)^{-1} \lambda = \lim P^{-1} \text{tr}(\Sigma_\nu \Psi(zI + \Psi)^{-1} \Psi) \quad (256)$$

and, similarly,

$$\lim_{P \rightarrow \infty} A'(z) = -\lim \lambda' \Psi(zI + \Psi)^{-2} \lambda = -\lim P^{-1} \text{tr}(\Sigma_\nu \Psi(zI + \Psi)^{-2} \Psi) \quad (257)$$

by the standard arguments based on the Cauchy theorem. Now, by (244),

$$\Gamma_{1,1} = \lim \frac{z^{-1} P^{-1} \text{tr}(\Psi^2(I - \Psi\delta(z))^{-1} \Sigma_\nu)}{1 - \delta(z) P^{-1} \text{tr}(\Psi^2(I - \Psi\delta(z))^{-1} \Sigma_\nu)} \quad (258)$$

with

$$\delta(z) = -z^{-1}(1 + \xi(z; c))^{-1} = -1/Z^*. \quad (259)$$

Thus,

$$\begin{aligned}
& \lim P^{-1} \text{tr}(\Psi^2(I - \Psi\delta(z))^{-1}\Sigma_\nu) \\
&= \lim P^{-1} \text{tr}(\Psi^2(I + \Psi/Z^*)^{-1}\Sigma_\nu) \\
&= \lim P^{-1} Z^* \text{tr}(\Psi^2(Z^* + \Psi)^{-1}\Sigma_\nu) \\
&= Z^* A(Z^*)
\end{aligned} \tag{260}$$

and, hence,

$$\begin{aligned}
\Gamma_{1,1} &= \lim \frac{z^{-1} P^{-1} \text{tr}(\Psi^2(I - \Psi\delta(z))^{-1}\Sigma_\nu)}{1 - \delta(z) P^{-1} \text{tr}(\Psi^2(I - \Psi\delta(z))^{-1}\Sigma_\nu)} \\
&= \frac{z^{-1} Z^* A(Z^*)}{1 + A(Z^*)}
\end{aligned} \tag{261}$$

so that

$$\lim_{T, P \rightarrow \infty, P/T \rightarrow c, N \rightarrow \infty} E[\hat{R}_{t+1}^M(z)] = \frac{\Gamma_{1,1}(z)}{1 + \xi(z; c)} = \frac{A(Z^*)}{1 + A(Z^*)} \tag{262}$$

because

$$1 + \xi = z^{-1} Z^* \tag{263}$$

The proof of Proposition 6 is complete.

□

J Proof of Theorem 3, ii.: Second Moment of $\hat{R}^M$

We start with

Lemma 24 *We have*

$$G(z; c) = \frac{d}{dz}(z\xi(z; c)) \in (0, cz^{-2}] \quad (264)$$

satisfies

$$G(z; c) = \mathcal{M}(z; Z^*(z; c)), \quad (265)$$

where

$$\mathcal{M}(z; Z) = -1 + \frac{Z}{z + c\phi(Z)Z^2}, \quad \phi(z) = \lim P^{-1} \text{tr}(E[FF'](zI + E[FF'])^{-2}). \quad (266)$$

Proof of Lemma 24. By the master equation (Theorem 4),

$$m(z; c) = \frac{1}{1 - c - czm(z; c)} m_{\sigma_*\Psi} \left(\frac{z}{1 - c - czm(z; c)} \right). \quad (267)$$

Recall that

$$Z^*(z; c) = \frac{z}{1 - c + czm(-z; c)}, \quad (268)$$

and we can thus rewrite the Master equation as

$$zm(-z; c) = Z^*(z; c)m(-Z^*(z; c)), \quad z > 0. \quad (269)$$

By the definition of the $\xi(z; c)$ function (Lemma 18),

$$\frac{c^{-1}\xi(z; c)}{1 + \xi(z; c)} = 1 - m(-z; c)z. \quad (270)$$

and hence

$$\xi(z; c) = \frac{1 - zm(-z; c)}{c^{-1} - 1 + zm(-z; c)} \quad (271)$$

and hence

$$1 + \xi(z; c) = \frac{c^{-1}}{c^{-1} - 1 + zm(-z; c)} = \frac{1}{1 - c + czm(-z; c)} \quad (272)$$

Differentiating this identity, we get

$$\xi'(z; c) = -c(zm(-z; c))'(1 + \xi(z; c))^2 \quad (273)$$

Furthermore, differentiating the identity

$$zm(-z; c) = Z^*(z; c)m(-Z^*(z; c)), \quad (274)$$

we get

$$(zm(-z; c))' = (zm(-z))'(Z^*)Z^{*'} = (zm(-z))'(Z^*)(1 + \xi(z; c) + z\xi'(z; c)) \quad (275)$$

so that

$$\xi'(z; c) = -c(zm(-z))'(Z^*)(1 + \xi(z; c) + z\xi'(z; c))(1 + \xi(z; c))^2, \quad (276)$$

implying that

$$\xi'(z; c) = \frac{-c(zm(-z))'(Z^*)(1 + \xi(z; c))^3}{1 + c(zm(-z))'(Z^*)z(1 + \xi(z; c))^2} \quad (277)$$

and hence

$$1 + \xi(z; c) + z\xi'(z; c) = \frac{1 + \xi(z; c)}{1 + c(zm(-z))'(Z^*)z(1 + \xi(z; c))^2} = \frac{Z^*(z; c)}{z + c(zm(-z))'(Z^*)Z^{*2}(z; c)} \quad (278)$$

□

Let

$$\overline{F}_t = \sum_t F_t.$$

All kurtosis terms with $\kappa - 3$ vanish asymptotically due to their vanishing trace norm. Using Lemma 9, we get⁴⁹

$$\begin{aligned} E[(\hat{R}_{t+1}^M(z))^2] &= E\left[\frac{1}{T}\overline{F}_t'(zI + B_T)^{-1}F_{t+1}F_{t+1}'(zI + B_T)^{-1}\frac{1}{T}\overline{F}_t\right] \\ &= E\left[\frac{1}{T}\overline{F}_t'(zI + B_T)^{-1}E_{t-}[F_{t+1}F_{t+1}'](zI + B_T)^{-1}\frac{1}{T}\overline{F}_t\right] \\ &\stackrel{\text{Lemma 9}}{=} E\left[\frac{1}{T}\overline{F}_t'(zI + B_T)^{-1}\left((\text{tr} \Sigma)^2 + \text{tr}(\Sigma^2))\Psi\Sigma_F\Psi + \Psi\left(\text{tr}(\Sigma\Sigma_\varepsilon) + \text{tr}(\Psi\Sigma_F)\text{tr}(\Sigma^2)\right)\right)\right. \\ &\quad \left.(zI + B_T)^{-1}\frac{1}{T}\overline{F}_t\right] \\ &\stackrel{\text{tr}(\Sigma^2)=o(1)}{\approx} E\left[\frac{1}{T}\overline{F}_t'(zI + B_T)^{-1}\left((\text{tr} \Sigma)^2\Psi\Sigma_F\Psi + \Psi\text{tr}(\Sigma\Sigma_\varepsilon)\right)\right. \\ &\quad \left.(zI + B_T)^{-1}\frac{1}{T}\overline{F}_t\right] \\ &= \frac{1}{T^2} \sum_{t_1, t_2} E[F_{t_1}(zI + B_T)^{-1}\left((\text{tr} \Sigma)^2\Psi\Sigma_F\Psi + \Psi\text{tr}(\Sigma\Sigma_\varepsilon)\right)(zI + B_T)^{-1}F_{t_2}] \\ &\sim \text{Term1} + \text{Term2}, \end{aligned} \quad (279)$$

⁴⁹ E_{t-} denotes the expectation averaging over realizations of S_t and R_{t+1} .

with

$$Term1 = \frac{1}{T} E[F'_{t_1} (zI + B_T)^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_T)^{-1} F_{t_1}] \quad (280)$$

and

$$Term2 = \frac{T(T-1)}{T^2} E[F'_{t_1} (zI + B_T)^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_T)^{-1} F_{t_2}] \quad (281)$$

for any $t_1 \neq t_2$, where we have used interchangeability of F_t to get rid of the summation over t_1 and t_1, t_2 .

J.1 *Term1* in (280)

We first deal with the first term. Using the Sherman-Morrison formula and Lemma 18, and Lemma 9, we get

$$\begin{aligned} Term1 &= \frac{1}{T} \text{tr } E \left[\left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_T)^{-1} F_{t_1} F'_{t_1} (zI + B_T)^{-1} \right] \\ &\sim \frac{1}{T} \text{tr } E \left[\left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t_1})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1})^{-1} \right] (1 + \xi(z; c))^{-2} \\ &\sim \frac{1}{T} \text{tr } E \left[\left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t})^{-1} \right. \\ &\quad \left. \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t_1})^{-1} \right] (1 + \xi(z; c))^{-2}, \end{aligned} \quad (282)$$

where we have used Lemma 5 and Lemma 2 to take the two $1 + \hat{\xi}_t$ terms out of expectation and replace them with their limit, $1 + \xi(z; c)$. We can now split this expression into several

terms. We have

$$\begin{aligned} & \frac{1}{T} \text{tr} E[(\text{tr} \Sigma)^2 \Psi \Sigma_F \Psi (zI + B_{T,t})^{-1} (\text{tr} \Sigma)^2 \Psi \Sigma_F \Psi (zI + B_{T,t})^{-1}] (1 + \xi(z; c))^{-2} \\ &= \frac{1}{T} \text{tr} E[\Psi \Sigma_F \Psi (zI + B_{T,t})^{-1} \Psi \Sigma_F \Psi (zI + B_{T,t})^{-1}] (1 + \xi(z; c))^{-2} \rightarrow 0 \end{aligned} \quad (283)$$

because

$$\text{tr}(\Sigma_F) = \text{tr}(\Sigma_{F,t}) + \|\nu\|^2 = O(1) = o(T),$$

and all other matrices involved are uniformly bounded. The second term is

$$\frac{1}{T} \text{tr} E[(\text{tr} \Sigma)^2 \Psi \Sigma_F \Psi (zI + B_{T,t})^{-1} \text{tr}(\Sigma \Sigma_\varepsilon) \Psi (zI + B_{T,t})^{-1}] / (1 + \xi(z; c))^2 = O(T^{-1}) \quad (284)$$

by the same argument. Finally, the last term is

$$(\text{tr}(\Sigma \Sigma_\varepsilon))^2 \frac{1}{T} \text{tr} E[\Psi (zI + B_{T,t})^{-1} \Psi (zI + B_{T,t})^{-1}] / (1 + \xi(z; c))^2 \quad (285)$$

and it needs to be evaluated directly.

Lemma 25 *We have*

$$\begin{aligned} & \frac{1}{P} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}] \\ & \sim \sigma_*^2 \frac{1}{P} \text{tr} E[\Psi (zI + B_T)^{-1} \Psi (zI + B_T)^{-1}] \\ & \rightarrow \Gamma_3(z) = \left(1 - (-z^2 m'(-z; c) + 2zm(-z; c) + c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2) \right) (1 + \xi(z; c))^4 \quad (286) \\ & = c^{-1} (\xi(z; c) + z\xi'(z; c)) (1 + \xi(z; c))^2 \\ & = c^{-1} G(z; c) (1 + \xi(z; c))^2 \end{aligned}$$

The proof of Lemma 25 is non-trivial as it requires deriving remarkable algebraic identities for the expectations involving both true and estimated covariance matrices. We do it by

expanding the trivial identity

$$1 = \frac{1}{P} \text{tr} E[(zI + B_T)(zI + B_T)^{-1}(zI + B_T)(zI + B_T)^{-1}]$$

and then writing out B_T and using the Sherman-Morrison formula in a smart way multiple times.

Proof of Lemma 25. We have, by the Sherman-Morrison formula, that

$$\begin{aligned} & \frac{1}{P} \frac{1}{T} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_T)^{-1} F_{t_1} F'_{t_1} (zI + B_T)^{-1}] \\ & \sim \frac{1}{c} \frac{1}{T^2} E[F'_{t_1} (zI + B_T)^{-1} F_{t_1} F'_{t_1} (zI + B_T)^{-1} F_{t_1}] \\ & = c^{-1} E \left[\left(\frac{\frac{1}{T} F'_{t_1} (zI + B_{T,t_1})^{-1} F_{t_1}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1})^{-1} F_{t_1}} \right)^2 \right] \\ & \sim c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \end{aligned} \tag{287}$$

by Lemma 18. Now,

$$m'(-z; c) = \lim P^{-1} \text{tr} E[(zI + B_T)^{-2}] \tag{288}$$

and hence

$$\begin{aligned} 1 &= \frac{1}{P} \text{tr} E[(zI + B_T)(zI + B_T)^{-1}(zI + B_T)(zI + B_T)^{-1}] \\ &= \frac{1}{P} z^2 \text{tr} E[(zI + B_T)^{-2}] + 2z \frac{1}{P} \text{tr} E[(zI + B_T)^{-2} B_T] \\ &+ \frac{1}{P} \text{tr} E[B_T (zI + B_T)^{-1} B_T (zI + B_T)^{-1}] \\ &\sim z^2 m'(-z; c) + 2z \frac{1}{P} \text{tr} E[(zI + B_T)^{-2} (B_T + zI - zI)] \\ &+ \frac{1}{P} \frac{1}{T^2} \sum_{t_1, t_2} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_T)^{-1} F_{t_2} F'_{t_2} (zI + B_T)^{-1}] \end{aligned} \tag{289}$$

We now use symmetry over t to replace summations over t and rewrite this expression as

$$\begin{aligned}
& -z^2 m'(-z; c) + 2zm(-z; c) + \frac{1}{P} \frac{1}{T} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_T)^{-1} F_{t_1} F'_{t_1} (zI + B_T)^{-1}] \\
& + \frac{1}{P} \frac{T(T-1)}{T^2} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_T)^{-1} F_{t_2} F'_{t_2} (zI + B_T)^{-1}] \\
& \sim -z^2 m'(-z; c) + 2zm(-z; c) + c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \\
& + \frac{1}{P} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_T)^{-1} F_{t_2} F'_{t_2} (zI + B_T)^{-1}] \\
& \sim -z^2 m'(-z; c) + 2zm(-z; c) + c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \\
& + \frac{1}{P} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_{T,t_1})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_2})^{-1}] / (1 + \xi(z; c))^2 \\
& \sim -z^2 m'(-z; c) + 2zm(-z; c) + c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \\
& + \frac{1}{P} E[F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}] / (1 + \xi(z; c))^4 \\
& = -z^2 m'(-z; c) + 2zm(-z; c) + c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \\
& + \frac{1}{P} \text{tr} E[F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}] / (1 + \xi(z; c))^4
\end{aligned} \tag{290}$$

where we have defined

$$B_{T,t_1,t_2} = \frac{1}{T} \sum_{\tau \notin \{t_1, t_2\}} F_\tau F'_\tau. \tag{291}$$

We also used that

$$F'_{t_1} (zI + B_T)^{-1} \sim F'_{t_1} (zI + B_{T,t_1})^{-1} / (1 + \xi(z; c))$$

by Lemma 18 and the Sherman-Morrison formula. The replacement of four different $1 + \hat{\xi}_t$ with $(1 + \xi)$ and the emergence of $(1 + \xi)^4$ is justified by Lemma 5 and Lemma 2.

Now,

$$\begin{aligned}
& \frac{1}{P} \operatorname{tr} E[F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}] \\
&= \frac{1}{P} \operatorname{tr} E \left[\left(((\operatorname{tr} \Sigma)^2 + \operatorname{tr}(\Sigma^2)) \Psi \Sigma_F \Psi \right. \right. \\
&\quad \left. \left. + \Psi \left(\operatorname{tr}(\Sigma \Sigma_\varepsilon) + \operatorname{tr}(\Sigma_F \Psi) \operatorname{tr}(\Sigma^2) \right) \right) (zI + B_{T,t_1,t_2})^{-1} \left(((\operatorname{tr} \Sigma)^2 + \operatorname{tr}(\Sigma^2)) \Psi \Sigma_F \Psi \right. \right. \\
&\quad \left. \left. + \Psi \left(\operatorname{tr}(\Sigma \Sigma_\varepsilon) + \operatorname{tr}(\Sigma_F \Psi) \operatorname{tr}(\Sigma^2) \right) \right) (zI + B_{T,t_1,t_2})^{-1} \right] \tag{292}
\end{aligned}$$

which coincides with the expression in (282). By the derivations in formulas (283) and (284), we get

$$\begin{aligned}
& \frac{1}{P} \operatorname{tr} E[F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}] \\
&\sim \sigma_*^2 \frac{1}{P} \operatorname{tr} E[\Psi (zI + B_T)^{-1} \Psi (zI + B_T)^{-1}], \tag{293}
\end{aligned}$$

and hence

$$\begin{aligned}
1 &= -z^2 m'(-z; c) + 2zm(-z; c) + c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \\
&+ \sigma_*^2 \frac{1}{P} \operatorname{tr} E[\Psi (zI + B_T)^{-1} \Psi (zI + B_T)^{-1}] / (1 + \xi(z; c))^4 \tag{294}
\end{aligned}$$

Finally,

$$\frac{\xi(z; c)}{1 + \xi(z; c)} = c(1 - zm(-z; c)) \tag{295}$$

and therefore, we can perform the following (surprising) algebraic manipulations:

$$\begin{aligned}
& (1 + z^2 m'(-z; c) - 2zm(-z; c) - c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2)(1 + \xi(z; c))^4 \\
&= \left(\frac{d}{dz} (z(1 - zm(-z; c))) - c^{-1} \left(\frac{\xi(z; c)}{1 + \xi(z; c)} \right)^2 \right) (1 + \xi(z; c))^4 \\
&= c^{-1} \left(\frac{d}{dz} \left(\frac{z\xi(z; c)}{1 + \xi(z; c)} \right) (1 + \xi(z; c))^2 - (\xi(z; c))^2 \right) (1 + \xi(z; c))^2 \\
&= c^{-1} \left(\frac{d}{dz} \left(z - \frac{z}{1 + \xi(z; c)} \right) (1 + \xi(z; c))^2 - (\xi(z; c))^2 \right) (1 + \xi(z; c))^2 \\
&= c^{-1} \left(\left(1 - \frac{1}{1 + \xi(z; c)} + \frac{z\xi'(z; c)}{(1 + \xi(z; c))^2} \right) (1 + \xi(z; c))^2 - (\xi(z; c))^2 \right) (1 + \xi(z; c))^2 \\
&= c^{-1} (\xi(z; c) + z\xi'(z; c)) (1 + \xi(z; c))^2.
\end{aligned} \tag{296}$$

The proof of Lemma 25 is complete. □

We conclude that the first term from (279) characterized in (282) satisfies

$$\begin{aligned}
Term1 &= \frac{1}{T} E[F'_{t_1} (zI + B_T)^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_T)^{-1} F_{t_1}] \\
&\sim (1 + \xi(z; c))^{-2} c \Gamma_3(z)
\end{aligned} \tag{297}$$

because $1/T \sim c/P$.

J.2 Term2 in (281)

We now proceed with the second term (281). By the Sherman-Morrison formula and Lemma 18,

$$\begin{aligned}
& E[F'_{t_1}(zI + B_T)^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_T)^{-1} F_{t_2}] \\
& \sim E[F'_{t_1}(zI + B_{T,t_1})^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t_2})^{-1} F_{t_2}] / (1 + \xi(z; c))^2 \\
& \sim E[F'_{t_1} \left((zI + B_{T,t_1,t_2})^{-1} - \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} F_{t_2}} \right) \\
& \quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \left((zI + B_{T,t_1,t_2})^{-1} \right. \\
& \quad \left. - \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}} \right) F_{t_2}] / (1 + \xi(z; c))^2 \\
& = \text{Term1} + \text{Term2} + \text{Term3}
\end{aligned} \tag{298}$$

where

$$\begin{aligned}
\text{Term1} &= E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
& \quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t_1,t_2})^{-1} F_{t_2}] / (1 + \xi(z; c))^2 \\
\text{Term2} &= -2E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
& \quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \\
& \quad \times \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}} F_{t_2}] / (1 + \xi(z; c))^2 \\
\text{Term3} &= E[F'_{t_1} \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} F_{t_2}} \\
& \quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}} F_{t_2}] / (1 + \xi(z; c))^2
\end{aligned} \tag{299}$$

We now analyze each term separately.

J.3 Term1 in (299)

We will need the following lemma.

Lemma 26 *We have*

$$F(A) = \lambda' E[(zI + B_T)^{-1} A (zI + B_T)^{-1}] \lambda \rightarrow 0 \quad (300)$$

for any A with uniformly bounded trace norm, with A independent of λ .

Proof of Lemma 26. This lemma is non-trivial because B_T is not independent of λ and, hence, simple arguments above based on the observation that $\nu = P^{-1} \Sigma_\nu^{1/2} \hat{\nu}$ and, hence, annihilates any bounded-trace matrix, are not directly applicable. Indeed, if A is independent of ν , then (under the normalization $\sigma_* = 1$),

$$\lambda' A \lambda \approx \nu' A \nu \approx P^{-1} \text{tr}(A \Sigma_\nu) \approx 0 \quad (301)$$

if A is independent on λ . By contrast, if, for example, $A = \lambda \lambda'$, we have $\text{tr}(A) = \|\lambda\|^2 \approx P^{-1} \text{tr}(\Sigma_\nu)$ and, hence, is bounded. Yet,

$$\lambda' A \lambda = \|\lambda\|^4 \approx (P^{-1} \text{tr}(\Sigma_\nu))^2$$

does not converge to zero. Note also that “extracting” ν from F_t is tricky because $F_t - \nu$ is not independent of ν . We now proceed with our proof.

We know from Lemma 22 that $\lambda' E[A(zI + B_T)^{-1}] \lambda \rightarrow 0$. Furthermore,

$$\begin{aligned}
\lambda' E[A(zI + B_T)^{-1}] \lambda &= \lambda' E[(zI + B_T)^{-1}(zI + B_T)A(zI + B_T)^{-1}] \lambda \\
&\stackrel{\text{symmetry}}{=} z \lambda' E[(zI + B_T)^{-1}A(zI + B_T)^{-1}] \lambda + \frac{1}{T} \lambda' E[(zI + B_T)^{-1}F_t F_t' A(zI + B_T)^{-1}] \lambda \\
&= z \lambda' E[(zI + B_T)^{-1}A(zI + B_T)^{-1}] \lambda \\
&+ E\left[\left((zI + B_{T,t})^{-1} - \frac{\frac{1}{T}(zI + B_{T,t})^{-1}F_t F_t'(zI + B_{T,t})^{-1}}{1 + \frac{1}{T}F_t'(zI + B_{T,t})^{-1}F_t}\right)F_t F_t' A(zI + B_T)^{-1} \lambda\right] \\
&\approx z \lambda' E[(zI + B_T)^{-1}A(zI + B_T)^{-1}] \lambda + (1 + \xi(z; c))^{-1} \lambda' E[(zI + B_{T,t})^{-1}F_t F_t' A \\
&\times \left((zI + B_{T,t})^{-1} - \frac{\frac{1}{T}(zI + B_{T,t})^{-1}F_t F_t'(zI + B_{T,t})^{-1}}{1 + \xi(z; c)}\right) \lambda] \\
&= z \lambda' E[(zI + B_T)^{-1}A(zI + B_T)^{-1}] \lambda \\
&+ (1 + \xi(z; c))^{-1} \lambda' E[(zI + B_{T,t})^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) A(zI + B_{T,t})^{-1}] \lambda \\
&- (1 + \xi(z; c))^{-2} \lambda' E[(zI + B_{T,t})^{-1}F_t F_t' A \frac{1}{T}(zI + B_{T,t})^{-1}F_t F_t'(zI + B_{T,t})^{-1}] \lambda \\
&\approx z \lambda' E[(zI + B_T)^{-1}A(zI + B_T)^{-1}] \lambda \\
&+ (1 + \xi(z; c))^{-1} \lambda' E[(zI + B_{T,t})^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) A(zI + B_{T,t})^{-1}] \lambda \\
&- Q(z)(1 + \xi(z; c))^{-2} \lambda' E[(zI + B_{T,t})^{-1}F_t F_t'(zI + B_{T,t})^{-1}] \lambda
\end{aligned} \tag{302}$$

where

$$Q(z) = F_t' A \frac{1}{T} (zI + B_{T,t})^{-1} F_t \sim T^{-1} \text{tr } E[\Psi A(zI + B_{T,t})^{-1}] \rightarrow 0 \tag{303}$$

by Lemma 3 because $\|A\|_1 = o(P)$ by assumption, and

$$\begin{aligned}
&\lambda' E[(zI + B_{T,t})^{-1}F_t F_t'(zI + B_{T,t})^{-1}] \lambda \\
&= \lambda' E[(zI + B_{T,t})^{-1} \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t})^{-1}] \lambda = O(1).
\end{aligned} \tag{304}$$

Thus, we get

$$o(1) \approx zF(A) + (1 + \xi(z; c))^{-1} F((\Psi \Sigma_F \Psi + \Psi)A), \quad (305)$$

where $o(1)$ is uniform, and the same iterative argument as in the proof of Lemma 23 give a power series representation for $F((\Psi \Sigma_F \Psi + \Psi)^k A)$ for all k , and the same uniform boundedness argument implies that $F(A) = 0$. The proof of Lemma 26 is complete. $\square$

Now, $E[F_t] = \Psi \nu$ and therefore

$$\begin{aligned} (1 + \xi(z; c))^2 \text{Term1} &= E[F'_{t_1}(zI + B_{T, t_1, t_2})^{-1} \\ &\quad \left((\text{tr} \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T, t_1, t_2})^{-1} F_{t_2}] \\ &\sim (\text{tr}(\Sigma))^2 \nu' E[\Psi(zI + B_{T, t_1, t_2})^{-1} \\ &\quad \left((\text{tr} \Sigma)^2 \Psi(\Sigma_F + \nu \nu') \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T, t_1, t_2})^{-1}] \Psi \nu \\ &= (\text{tr}(\Sigma))^2 \nu' \Psi E[(zI + B_{T, t_1, t_2})^{-1} (\text{tr} \Sigma)^2 \Psi \Sigma_F \Psi (zI + B_{T, t_1, t_2})^{-1}] \Psi \nu \\ &\quad + (\text{tr}(\Sigma))^2 \nu' \Psi E[(zI + B_{T, t_1, t_2})^{-1} (\text{tr} \Sigma)^2 \Psi \nu \nu' \Psi (zI + B_{T, t_1, t_2})^{-1}] \Psi \nu \\ &\quad + (\text{tr}(\Sigma))^2 \nu' \Psi E[(zI + B_{T, t_1, t_2})^{-1} (\text{tr} \Sigma \Sigma_\varepsilon) \Psi (zI + B_{T, t_1, t_2})^{-1}] \Psi \nu \\ &= \hat{\Gamma}_{4, T}(z), \end{aligned} \quad (306)$$

where $\hat{\Gamma}_4$ is defined in the following lemma.

Lemma 27 *We have*

$$\hat{\Gamma}_{4, T}(z) \approx \Gamma_4(z) = \frac{\Gamma_{1, 1}(z) + z \Gamma'_{1, 1}(z)}{(1 + \xi(z; c))^{-2}} \quad (307)$$

Proof. We have by the symmetry across t and the Sherman-Morrison formula and Lemma

18 that

$$\begin{aligned}
\Gamma_{1,1}(z) &\sim \lambda' E[\Psi(zI + B_T)^{-1} \Psi] \lambda = \lambda' E[\Psi(zI + B_T)^{-1} (zI + B_T) (zI + B_T)^{-1} \Psi] \lambda \\
&= z \lambda' E[\Psi(zI + B_T)^{-1} (zI + B_T)^{-1} \Psi] \lambda + \lambda' E[\Psi(zI + B_T)^{-1} B_T (zI + B_T)^{-1} \Psi] \lambda \\
&= -z \Gamma'_{1,1,T}(z) + \lambda' E[\Psi(zI + B_T)^{-1} \frac{1}{T} \sum_t F_t F_t' (zI + B_T)^{-1} \Psi] \lambda \\
&= -z \Gamma'_{1,1,T}(z) + \lambda' E[\Psi(zI + B_T)^{-1} F_t F_t' (zI + B_T)^{-1} \Psi] \lambda \\
&\sim -z \Gamma'_{1,1,T}(z) + \lambda' E[\Psi(zI + B_{T,t})^{-1} F_t F_t' (zI + B_{T,t})^{-1} \Psi] \lambda (1 + \xi(z; c))^{-2} \\
&= -z \Gamma'_{1,1,T}(z) \\
&+ \lambda' E[\Psi(zI + B_{T,t})^{-1} \left(((\text{tr } \Sigma)^2 + \text{tr}(\Sigma^2)) \Psi \Sigma_F^* \Psi \right. \\
&+ \left. \Psi \left(\text{tr}(\Sigma \Sigma_\varepsilon) + \text{tr}(\Sigma_F \Psi) \text{tr}(\Sigma^2) \right) \right) (zI + B_{T,t})^{-1} \Psi] \lambda (1 + \xi(z; c))^{-2} \\
&\sim -z \Gamma'_{1,1,T}(z) + \hat{\Gamma}_{4,T}(z) (1 + \xi(z; c))^{-2}.
\end{aligned} \tag{308}$$

The claim follows now because, by Lemma 23, $\Gamma_{1,1,T}(z) \rightarrow \Gamma_{1,1}(z)$ implies $\Gamma'_{1,1,T}(z) \rightarrow \Gamma'_{1,1}(z)$ by standard properties of analytic functions (the Vitali Theorem; (Kelly et al., 2024b)). The proof of Lemma 27 is complete. $\square$

J.4 Term2 in (299)

The next term in (299) is (note the factor of 2 as it appears two times):

$$\begin{aligned}
Term2 &= -2E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \\
&\quad \times \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}} F_{t_2}] / (1 + \xi(z; c))^2 \\
&= -2E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \\
&\quad \times \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} \Psi \nu}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}} \text{tr}(\Sigma) / (1 + \xi(z; c))^2 \\
&\sim -2E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \\
&\quad \times \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} \Psi \nu}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}}] / (1 + \xi(z; c))^2 \\
&= -2(1 + \xi(z; c))^{-2} E[X_T Y_T],
\end{aligned} \tag{309}$$

where we have used that

$$E[F_{t_2}] = \Psi \nu, \tag{310}$$

and where

$$\begin{aligned}
Y_T &= F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \Psi \nu \\
X_T &= F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \\
&\quad \times \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}}
\end{aligned} \tag{311}$$

We will need the following technical lemma whose proof follows directly from the Cauchy-Schwarz inequality.

Lemma 28 *If $X_T \rightarrow X$ in probability and is uniformly bounded and $E[Y_T^2]$ is uniformly bounded. Then,*

$$E[(X_T - X)Y_T] \rightarrow 0$$

Then, we will need

Lemma 29 *We have*

$$E[(Y_T)^2]$$

is uniformly bounded whereas

$$E[Y_T] = E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1}\lambda] \rightarrow \Gamma_{1,1}(z) \quad (312)$$

in L_2 .

Proof. Recall that

$$\lambda' \Psi^k(zI + B_T)^{-1} \Psi^\ell \lambda \rightarrow \Gamma_{k,l}(z) \quad (313)$$

by Lemma 23.

We have

$$\begin{aligned}
& E[(F'_{t_1}(zI + B_{T,t_1,t_2})^{-1}\lambda)^2] \\
&= E[F'_{t_1}(zI + B_{T,t_1,t_2})^{-1}\lambda\lambda'(zI + B_{T,t_1,t_2})^{-1}F_{t_1}] \\
&= \text{tr } E[(zI + B_{T,t_1,t_2})^{-1}\lambda\lambda'(zI + B_{T,t_1,t_2})^{-1}F_{t_1}F'_{t_1}] \\
&\sim \text{tr } E[(zI + B_{T,t_1,t_2})^{-1}\lambda\lambda'(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right)] \approx P^{-1} \text{tr } E[(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t_1,t_2})^{-1} \Sigma_\nu] \leq Kz^{-2}
\end{aligned} \tag{314}$$

for some $K > 0$. The proof of Lemma 29 is complete. $\square$

Our goal is to show that $X_T - E[X_T] \rightarrow 0$ in L_2 , so that

$$E[X_T Y_T] \sim E[X_T] E[Y_T] \sim E[X_T] \Gamma_{1,1}(z). \tag{315}$$

Lemma 30 *Let*

$$\begin{aligned}
X_T &= F'_{t_1}(zI + B_{T,t_1,t_2})^{-1} \\
&\quad \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \\
&\quad \times \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1}F_{t_1}}{1 + \frac{1}{T}F'_{t_1}(zI + B_{T,t_1,t_2})^{-1}F_{t_1}}
\end{aligned} \tag{316}$$

Then,

$$E[X_T] \rightarrow \frac{c\Gamma_3(z)}{1 + \xi(z; c)} = G(z; c)(1 + \xi(z; c)), \tag{317}$$

where $\Gamma_3(z)$ is defined in Lemma 25, and $X_T - E[X_T] \rightarrow 0$ in L_2 .

Proof of Lemma 30. We know from the proof of Lemma 11 that

$$\frac{1}{T} F'_t A F_t - \frac{1}{T} \text{tr}(A \Psi) \rightarrow 0$$

in L_2 for any bounded A independent of F_t . Thus, with $A = (zI + B_{T,t_1,t_2})^{-1} \left((\text{tr} \Sigma)^2 \Psi \Sigma_F^* \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) (zI + B_{T,t_1,t_2})^{-1}$, we have that

$$X_T \sim \frac{1}{T} \text{tr}(A \Psi), \quad (318)$$

under the assumed normalization $\sigma_* = \text{tr}(\Sigma \Sigma_\varepsilon) = 1$, and boundedness of the trace of Σ_F^* implies that its contribution is negligible, so that

$$\frac{1}{T} \text{tr}(A \Psi) \sim \frac{1}{T} \text{tr}((zI + B_{T,t_1,t_2})^{-1} \Psi (zI + B_{T,t_1,t_2})^{-1} \Psi). \quad (319)$$

A small modification of the argument in the proof of Lemma 17 (see (Kelly et al., 2024b)) implies that

$$\frac{1}{T} \text{tr}((zI + B_{T,t_1,t_2})^{-1} \Psi (zI + B_{T,t_1,t_2})^{-1} \Psi) - \frac{1}{T} \text{tr} E[(zI + B_{T,t_1,t_2})^{-1} \Psi (zI + B_{T,t_1,t_2})^{-1} \Psi] \rightarrow 0 \quad (320)$$

in L_p for any $p \geq 2$, at the rate at least $T^{-1/2}$. Thus,

$$X_T \rightarrow \frac{c\Gamma_3(z)}{1 + \xi(z; c)}$$

by Lemma 25. □

Since

$$E[Y_T] \rightarrow \Gamma_{1,1}(z)$$

by Lemma 29, and Lemma 28 and formula (309) imply that

$$Term2 \sim -2 \frac{c\Gamma_3(z)\Gamma_{1,1}(z)}{(1 + \xi(z; c))^3}. \quad (321)$$

J.5 Term3 in (299)

Finally, we now deal with Term3 in (299).

Lemma 31 *Term3 in (299) converges to zero.*

Proof of Lemma 31. We have

$$\begin{aligned} Term3 &= E \left[F'_{t_1} \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_2} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} F_{t_2}} \right. \\ &\quad \left. \left((\text{tr } \Sigma)^2 \Psi \Sigma_F \Psi + \Psi \text{tr}(\Sigma \Sigma_\varepsilon) \right) \frac{\frac{1}{T}(zI + B_{T,t_1,t_2})^{-1} F_{t_1} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1}}{1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1}} F_{t_2} \right] / (1 + \xi(z; c))^2 \\ &= E[X_T Y_T] / (1 + \xi(z; c))^2, \end{aligned} \quad (322)$$

where we have defined

$$X_T = \frac{\left(\frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} \right)^2}{\left(1 + \frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_1} \right) \left(1 + \frac{1}{T} F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} \right)}$$

and

$$Y_T = F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} \left(\Psi \Sigma_F \Psi + \sigma_* \Psi \right) (zI + B_{T,t_1,t_2})^{-1} F_{t_1}.$$

The first observation is that X_T is uniformly bounded by the Cauchy-Schwarz inequality and has a $O(1/T)$ L_2 -norm by Lemma 32. Since the first component of Y_T ,

$$F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} \Psi \Sigma_F \Psi (zI + B_{T,t_1,t_2})^{-1} F_{t_1}.$$

has a $o(T)$ L_2 -norm, we get that this part is negligible by Lemma 28.

Lemma 32 *We have that*

$$E[(F'_{t_1} A F_{t_2})^2] = O(\|A\|_1 \|A\|_\infty).$$

for any A . Thus,

$$\left(\frac{1}{T} F'_{t_1} (zI + B_{T,t_1,t_2})^{-1} F_{t_2} \right)^2$$

converges to zero in L_1 , while

$$F'_{t_2} (zI + B_{T,t_1,t_2})^{-1} \Psi \Sigma_F \Psi (zI + B_{T,t_1,t_2})^{-1} F_{t_1}$$

has a uniformly bounded L_2 -norm because $\text{tr}(\Sigma_F) = o(T)$.

Proof. We have

$$\begin{aligned} E[(F'_{t_1} A F_{t_2})^2] &= E[F'_{t_1} A F_{t_2} F'_{t_2} A F_{t_1}] \\ &= \text{tr } E[A F_{t_2} F'_{t_2} A F_{t_1} F'_{t_1}] \\ &\sim \text{tr } E\left[A \left(\Psi \Sigma_F \Psi + \sigma_* \Psi \right) \right. \\ &\quad \left. \times A \left(\Psi \Sigma_F \Psi + \sigma_* \Psi \right) \right] \end{aligned} \tag{323}$$

by Lemma 9. The proof of Lemma 32 is complete. $\square$

The proof then follows from the following result that we prove below.

Lemma 33 *We have*

$$E[(F'_{t_1} A F_{t_2})^4] = O(P^2)$$

for any uniformly bounded A and any $t_1 \neq t_2$.

Indeed, Lemma 33 implies that

$$E[X_T^2] \leq T^{-4} E[(F'_{t_1}(zI + B_{T,t_1,t_2})^{-1}F_{t_2})^4] = O(P^2/T^4)$$

while Lemma 32 implies that

$$E[Y_T^2] = O(P).$$

Thus,

$$|E[X_T Y_T]|^2 \leq E[X_T^2] E[Y_T^2] = O(P^2/T^4) O(P) \rightarrow 0$$

and the claim follows.

Proof of Lemma 33. The proof is similar to that of Lemma 11. We have

$$E[(F'_{t_1} A F_{t_2})^4] = E[(F'_{t_1} A F_{t_2} F'_{t_2} A F_{t_1})^2] \quad (324)$$

We can first evaluate the expectation over F_{t_1} using the bounds from the proof of Lemma 11. Note that the matrix $A F_{t_2} F'_{t_2} A$ has rank one and, hence, $\|A F_{t_2} F'_{t_2} A\| = F'_{t_2} A^2 F_{t_2}$. This gives us the bound

$$E[(F'_{t_1} A F_{t_2})^4] \leq K E[\text{tr}(A F_{t_2} F'_{t_2} A)^2] = K E[(F'_{t_2} A^2 F_{t_2})^2] = O(\text{tr}(A^2))^2 = O(P^2). \quad (325)$$

□

Thus, (322) converges to zero.

The proof of Lemma 31 is complete. □

Summarizing, we get from (309) and (306), (321), that

$$Term2 = (1 + \xi(z; c))^{-2} \hat{\Gamma}_4(z) - 2 \frac{c\Gamma_3(z)\Gamma_{1,1}(z)}{(1 + \xi(z; c))^3}, \quad (326)$$

and (279) implies

$$\begin{aligned} E[(\hat{R}_{t+1}^M(z))^2] &\stackrel{(279)}{\sim} Term1 + Term2 \\ &\stackrel{(297)}{\sim} (1 + \xi(z; c))^{-2} c\Gamma_3(z) + Term2 \\ &\stackrel{(326)}{\sim} (1 + \xi(z; c))^{-2} c\Gamma_3(z) + (1 + \xi(z; c))^{-2} \hat{\Gamma}_4(z) - 2 \frac{c\Gamma_3(z)\Gamma_{1,1}(z)}{(1 + \xi(z; c))^3} \\ &\approx G(z; c) + (1 + \xi)^{-2} \Gamma_4(z) - 2G(z; c)\Gamma_{1,1}(z)(1 + \xi)^{-1} \\ &= G(z; c) + (1 + \xi)^{-2} \Gamma_4(z) - 2G(z; c)\mathcal{E}(Z^*) \end{aligned} \quad (327)$$

where, by Lemma 27, we have

$$\Gamma_4(z) = \frac{\Gamma_{1,1}(z) + z\Gamma'_{1,1}(z)}{(1 + \xi(z; c))^{-2}}. \quad (328)$$

Now, we know from the above that

$$E[(\hat{R}_{t+1}^M(z))] = \mathcal{E}(Z^*) = \Gamma_{1,1}(z)/(1 + \xi(z; c)), \quad (329)$$

with $Z^* = z(1 + \xi(z; c))$. Thus,

$$\Gamma_{1,1}(z) = (1 + \xi(z; c))\mathcal{E}(Z^*) \quad (330)$$

and, hence,

$$\begin{aligned}
\Gamma'_{1,1}(z) &= \xi' \mathcal{E}(Z^*) + (1 + \xi) \mathcal{E}'(Z^*) (1 + \xi + z \xi') \\
\Rightarrow \Gamma_{1,1} + z \Gamma'_{1,1} &= (1 + \xi(z; c)) \mathcal{E}(Z^*) + z \xi' \mathcal{E}(Z^*) + z(1 + \xi) \mathcal{E}'(Z^*) (1 + \xi + z \xi') \\
&= (1 + G(z; c)) \mathcal{E}(Z^*) + Z^* \mathcal{E}'(Z^*) (1 + G(z; c)) \\
&= (1 + G(z; c)) (\mathcal{E}(Z^*) + Z^* \mathcal{E}'(Z^*))
\end{aligned} \tag{331}$$

By (51), we have

$$E[(R^{infeas})^2] = \mathcal{V}(z) + \mathcal{E}(z)^2 = \mathcal{E}(z) + z \mathcal{E}'(z) = (d/dz)(z \mathcal{E}(z)), \tag{332}$$

and we get the required identity:

$$\begin{aligned}
E[(\hat{R}_{t+1}^M(z))^2] &\rightarrow G(z; c) + (1 + G(z; c)) (\mathcal{V}(Z^*) + \mathcal{E}(Z^*)^2) - 2G(z; c) \mathcal{E}(Z^*) \\
&= \mathcal{V}(Z^*) + \mathcal{E}(Z^*)^2 + G(z; c) (1 + \mathcal{V}(Z^*) + \mathcal{E}(Z^*)^2 - 2\mathcal{E}(Z^*))
\end{aligned} \tag{333}$$

and, hence,

$$\text{Var}[(\hat{R}_{t+1}^M(z))] \rightarrow \mathcal{V}(Z^*) + G(z; c) (1 + \mathcal{V}(Z^*) + \mathcal{E}(Z^*)^2 - 2\mathcal{E}(Z^*)) \tag{334}$$

K Proof of Theorem 3, iv.: Pricing Errors

Our analysis relies on the quantities

$$\bar{F}_{OS} = E_{OS}[F] \in \mathbb{R}^P, \quad B_{OS} = E_{OS}[FF'] \in \mathbb{R}^{P \times P} \tag{335}$$

where $E_{OS}[X] = \frac{1}{T_{OS}} \sum_{t \in (T, T+T_{OS}]} X_t$ denotes an out-of-sample time series average. The pricing error properties of the SDF are particularly tractable to derive when the test assets

are the P factors F_t that underly the SDF. The out-of-sample pricing error vector is

$$\mathcal{E}_{OS}(z; P; T) = \frac{1}{T_{OS}} \sum_{t \in (T, T+T_{OS}]} F_t \hat{M}_t(z; P; T) \in \mathbb{R}^P. \quad (336)$$

Finally, following Hansen and Jagannathan (1997), we define the out-of-sample HJD as⁵⁰

$$\mathcal{D}_{OS}^{HJ}(z; P; T) = \mathcal{E}_{OS}(z; P; T)' B_{OS}^+ \mathcal{E}_{OS}(z; P; T), \quad (337)$$

where B_{OS}^+ is the Moore-Penrose quasi-inverse of the potentially degenerate matrix B_{OS} .

Recall that Theorem 3, iv, states that

- i. $\lim E[\hat{R}_{T+1}^M(z; q; c)] = \mathcal{E}(Z^*(z; q; c); q)$ where
 $Z^*(z; q; c) = z(1 + \xi(z; q; c)) \in (z, z + c)$ and $\xi(z; q; c) = \frac{c(1 - m(-z; c; q)z)}{1 - c(1 - m(-z; c; q)z)},$
- ii. $\lim \text{Var}[\hat{R}_{T+1}^M(z; q; c)] = \mathcal{V}(Z^*(z; q; c); q) + G(z; q; c) \mathcal{R}(Z^*(z; q; c); q)$ where
 $G(z; q; c) = (z\xi(z; q; c))' \in (0, cz^{-2}], \quad \mathcal{R}(z; q) \equiv (1 - \mathcal{E}(z; q))^2 + \mathcal{V}(z; q)$
- iii. $\lim \frac{\text{Var}[\hat{R}_{T+1}^M(z; q; c)]}{E[\hat{R}_{T+1}^M(z; q; c)]^2} = (1 + G(z; q; c)) \frac{1}{\mathcal{S}^2(Z^*; q)} + G(z; q; c) \left(\frac{1 - \mathcal{E}(Z^*; q)}{\mathcal{E}(Z^*; q)} \right)^2,$
- iv. $\lim E[\mathcal{D}_{OS}^{HJ}(z; q; P; T) - \bar{F}_{OS}' B_{OS}^+ \bar{F}_{OS}] = -1 + (1 + G(z; q)) \mathcal{R}(Z^*(z; q; c); q).$

As above, it suffices to deal with the case $q = 1$. Our proof is based on the following algebraic identity:

Proposition 7 *Let $B_{OS} = E_{OS}[FF']$, $\bar{F}_{OS} = E_{OS}[F]$. We have*

$$\mathcal{D}_{OS}^{HJ}(z; P; T) - \bar{F}_{OS}' B_{OS}^+ \bar{F}_{OS} = -2E_{OS}[\hat{R}_t^M(z; P; T)] + E_{OS}[(\hat{R}_t^M(z; P; T))^2] \quad (338)$$

⁵⁰At first glance, it may not be obvious whether we should define the HJD weighting matrix as the in-sample or out-of-sample second moment of factors. Upon further inspection, we find that the out-of-sample second moment is preferable because it allows us to establish a direct correspondence between $\mathcal{D}_{OS}^{HJ}(z; P; T)$ and the out-of-sample SDF Sharpe ratio.

Proof of Proposition 7. Let $F \in \mathbb{R}^{T_{OS} \times P}$ be the matrix of out-of-sample factor returns. We have $M = 1 - F\pi$, where π is our model for the efficient portfolio. Let also $\mathcal{I} = \mathbf{1}/T_{OS} \in \mathbb{R}^{T_{OS}}$. Then, $\bar{F}_{OS} = F\mathcal{I}$ and

$$\frac{1}{T_{OS}} \sum_{t \in (T, T+T_{OS}]} F_t \hat{M}_t(z; P; T) = \frac{1}{T_{OS}} \sum_{t \in (T, T+T_{OS}]} F_t (1 - F'_t \pi) = \bar{F}_{OS} - B_{OS} \pi \quad (339)$$

and, hence,

$$\begin{aligned} \mathcal{D}_{OS}^{HJ}(z; P; T) &= \mathcal{E}_{OS}(z; P; T)' B_{OS}^+ \mathcal{E}_{OS}(z; P; T) \\ &= (\bar{F}_{OS} - B_{OS} \pi)' B_{OS}^+ (\bar{F}_{OS} - B_{OS} \pi) \\ &= \bar{F}_{OS}' B_{OS}^+ \bar{F}_{OS} - 2 \bar{F}_{OS}' B_{OS}^+ B_{OS} \pi + \pi' B_{OS} \pi. \end{aligned} \quad (340)$$

Then,

$$E_{OS}[\hat{R}_t^M(z; P; T)] = \bar{F}_{OS}' \pi, \quad (341)$$

while

$$E_{OS}[(\hat{R}_t^M(z; P; T))^2] = \frac{1}{T_{OS}} \sum_t (F'_t \pi)^2 = \pi' B_{OS} \pi. \quad (342)$$

Finally,

$$\bar{F}_{OS}' B_{OS}^+ B_{OS} \pi = \bar{F}_{OS}' \pi \quad (343)$$

because $\bar{F}_{OS}$ is in the image of B_{OS} .

□

Now,

$$E[E_{OS}[\hat{R}_t^M(z; P; T)]] = E[\hat{R}_t^M(z; P; T)], \quad E[E_{OS}[(\hat{R}_t^M(z; P; T))^2]] = E[(\hat{R}_t^M(z; P; T))^2]$$

and the claim follows.

L Empirical Robustness

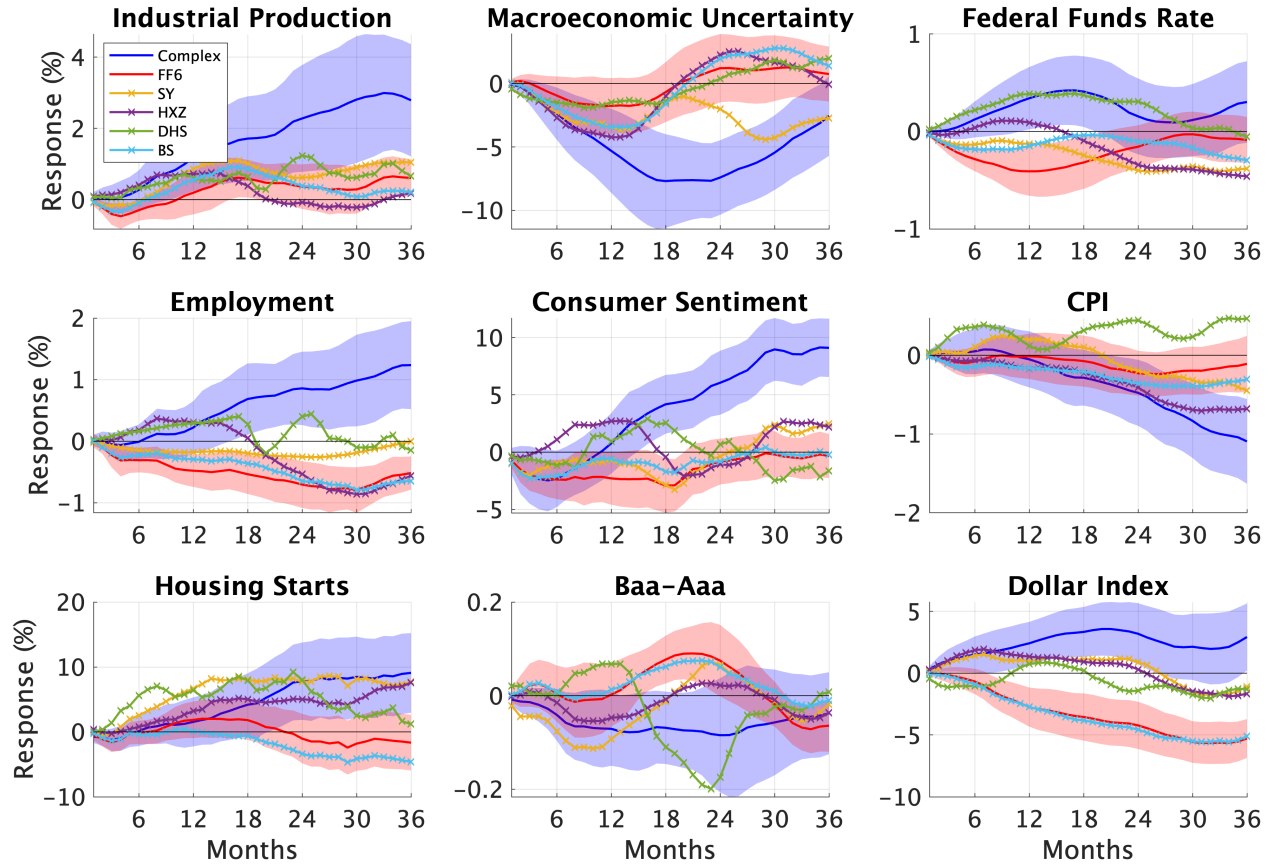


Figure IA.1: Macroeconomic Impulse Response Functions—All SDFs

Note. This figure extends the analysis of Figure 5 to include all benchmark SDFs.

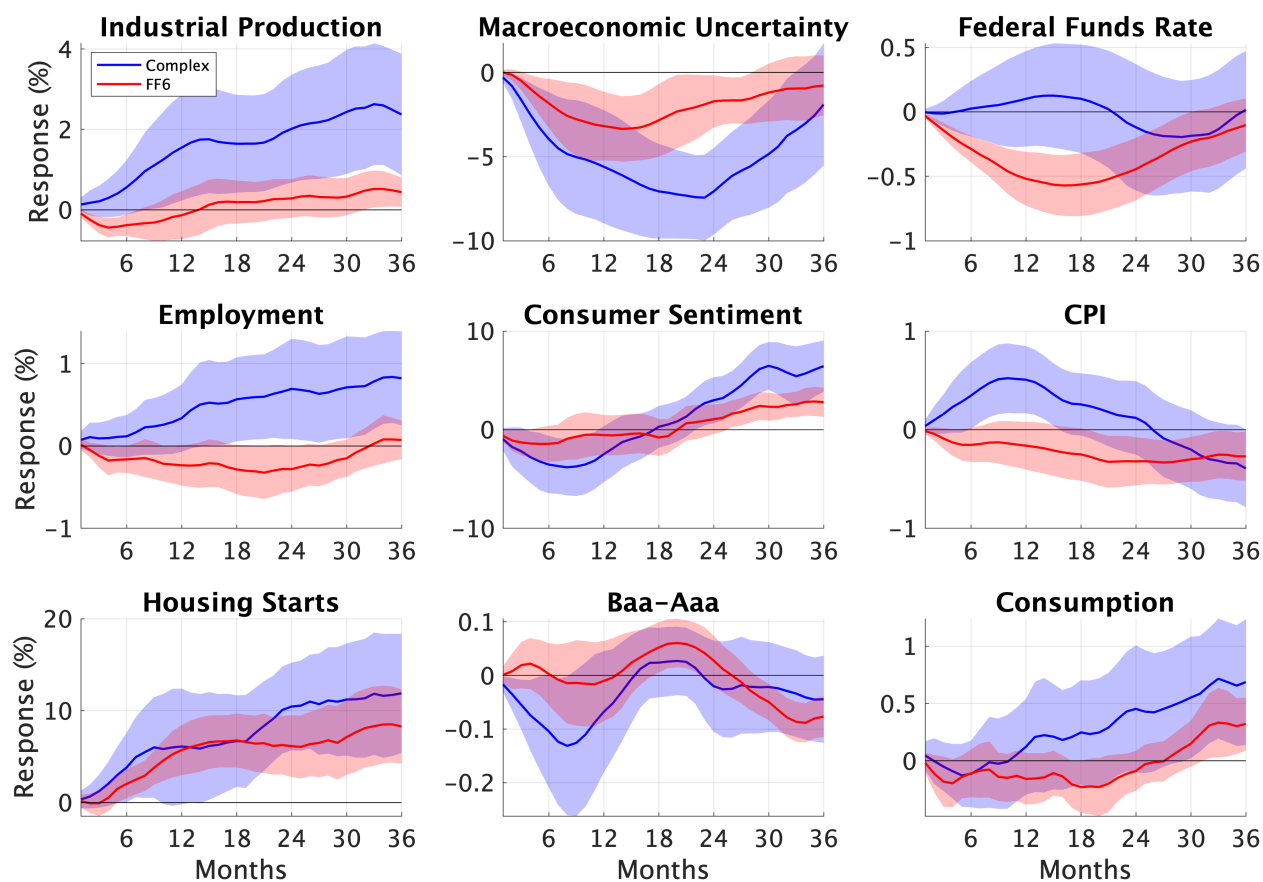


Figure IA.2: Macroeconomic Impulse Response Functions–Consumption

Note. This figure extends the analysis of Figure 5 to but replaces the last variable (Dollar Index) with consumption (Fred-MD variable “DPCERA3M086SBEA”).

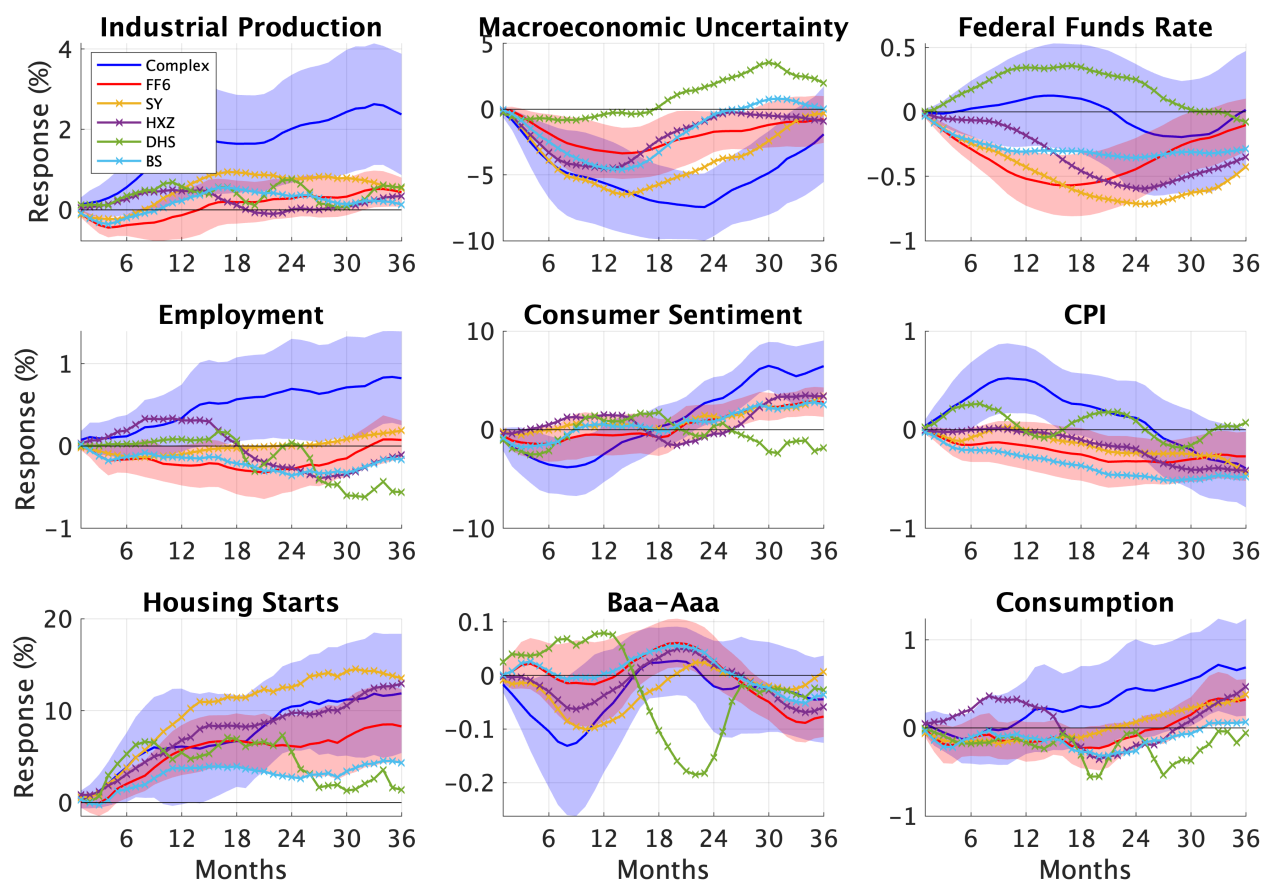


Figure IA.3: Macroeconomic Impulse Response Functions—Consumption and All SDFs

Note. This figure extends the analysis of Figure 5 to but replaces the last variable (Dollar Index) with consumption (Fred-MD variable “DPCERA3M086SBEA”).

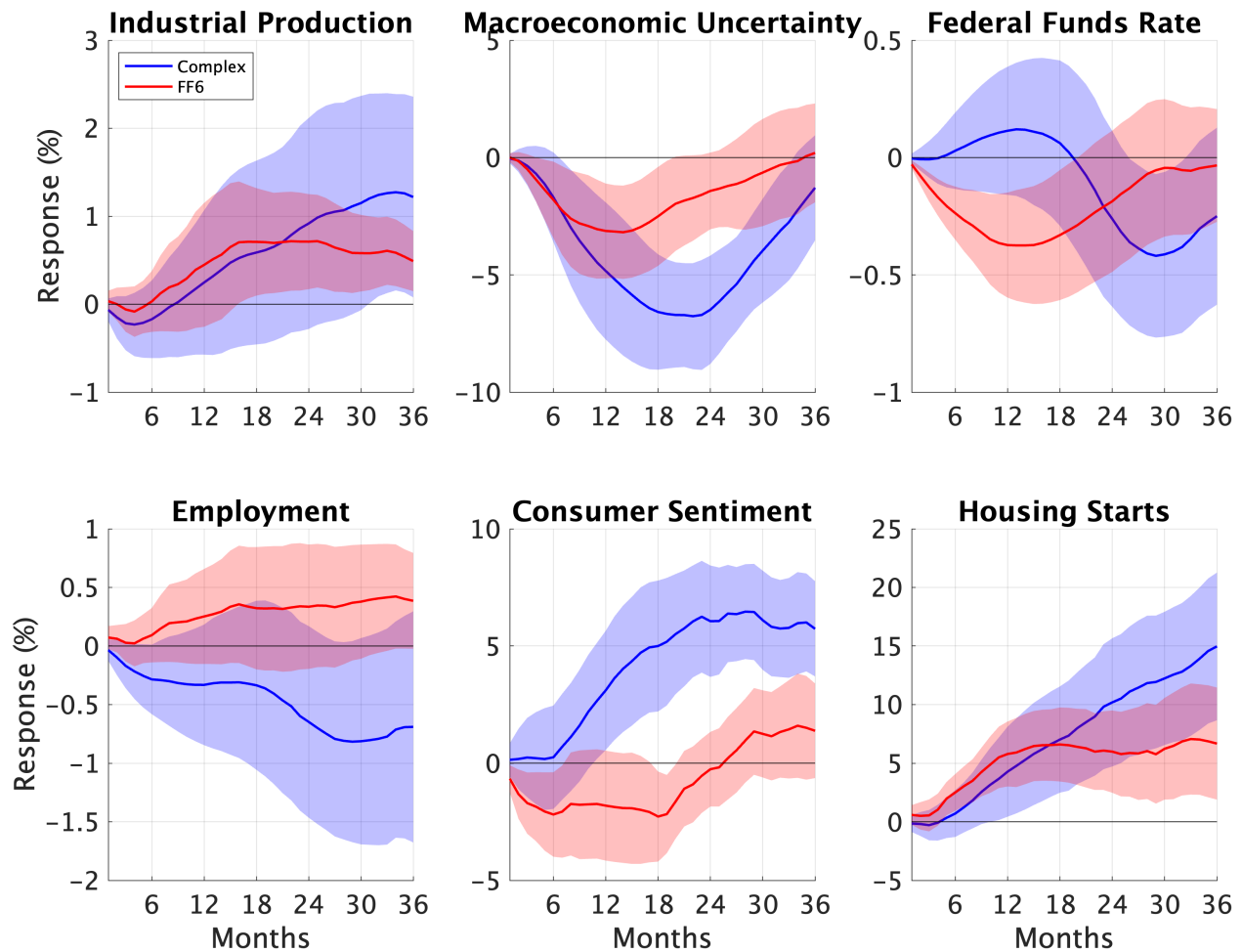


Figure IA.4: Macroeconomic Impulse Response Functions—Alternative Macro Variable Set (Medium)

Note. This figure repeats the analysis of Figure 5 using a smaller set of six macroeconomic variables in the y vector.

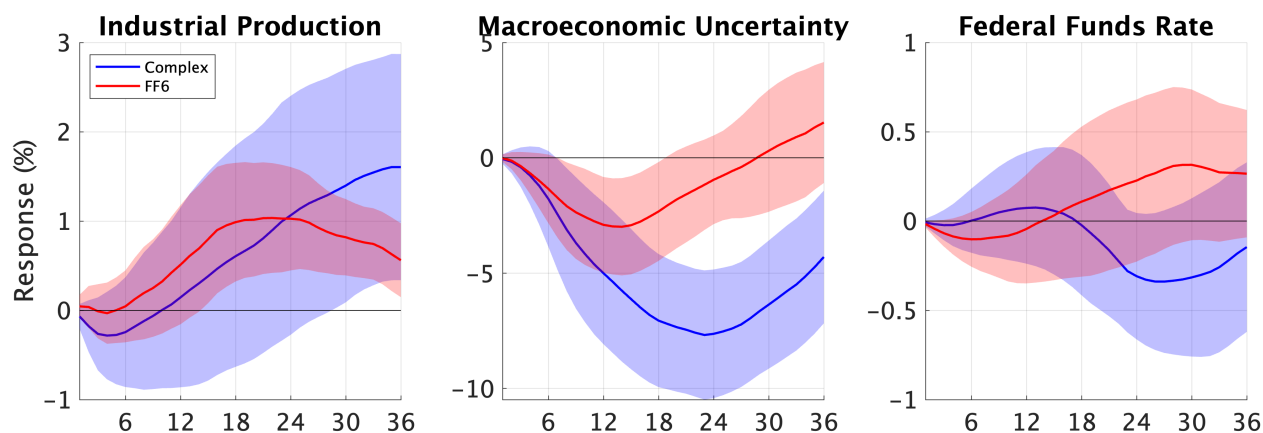


Figure IA.5: Macroeconomic Impulse Response Functions—Alternative Macro Variable Set (Small)

Note. This figure repeats the analysis of Figure 5 using a smaller set of three macroeconomic variables in the y vector.

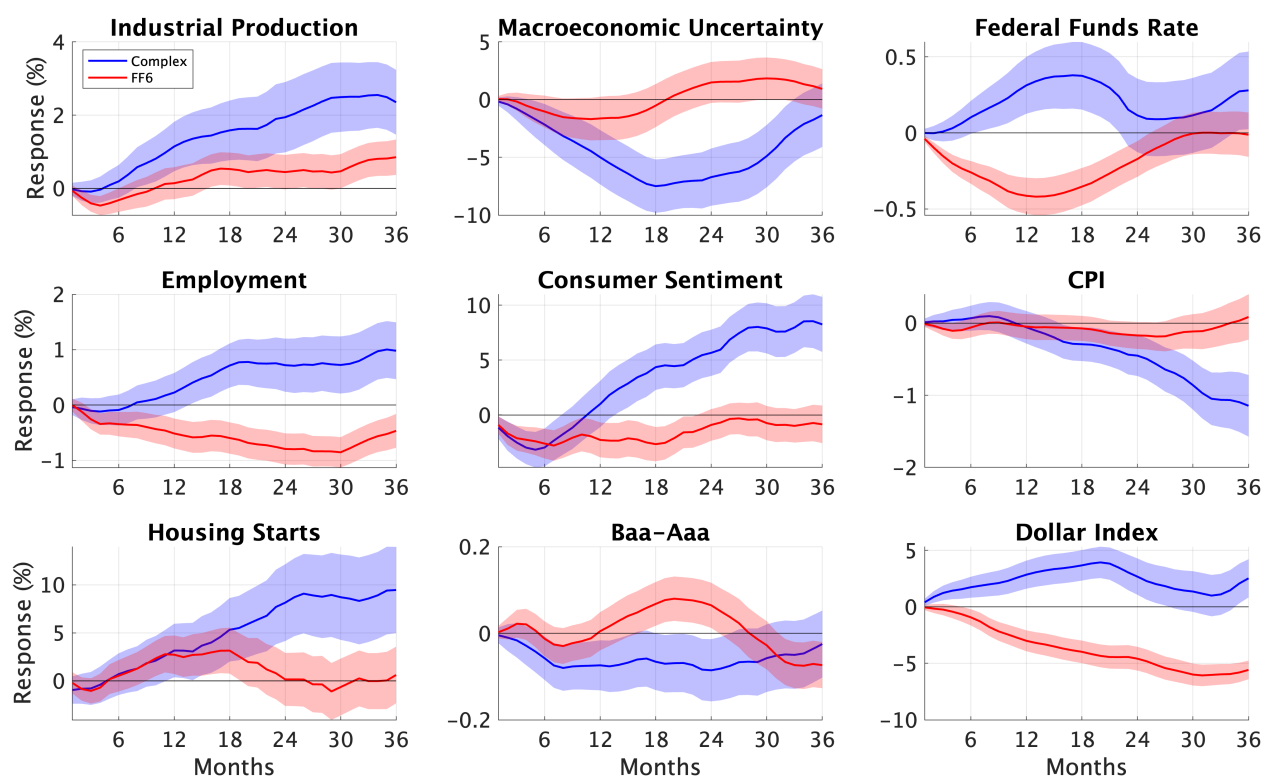


Figure IA.6: Macroeconomic Impulse Response Functions—Additional Macro Variable Lags

Note. This figure extends the analysis of Figure 5 to include six lags of macroeconomic variables (as opposed to three in the main analysis).

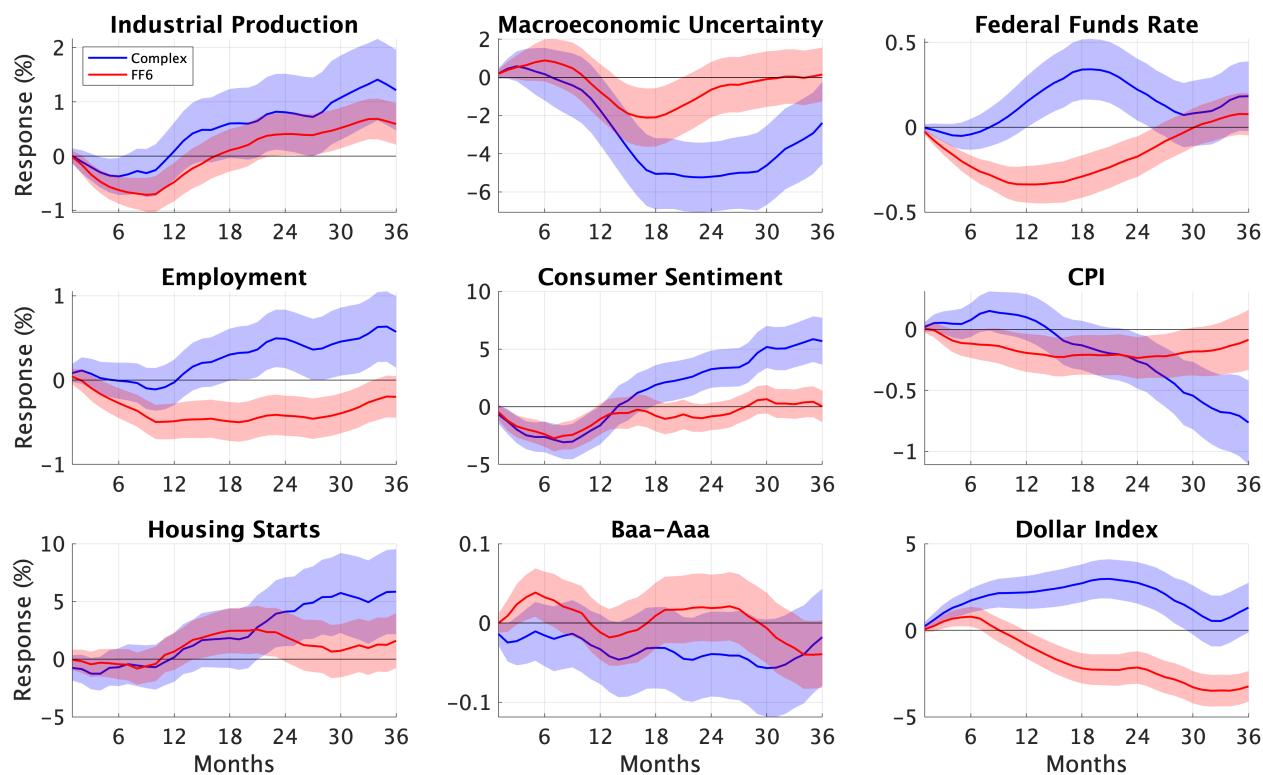


Figure IA.7: Macroeconomic Impulse Response Functions—Six-month SDF Returns

Note. This figure repeats the analysis of Figure 5 using six-month SDF returns (as opposed to 12 months in the main analysis).

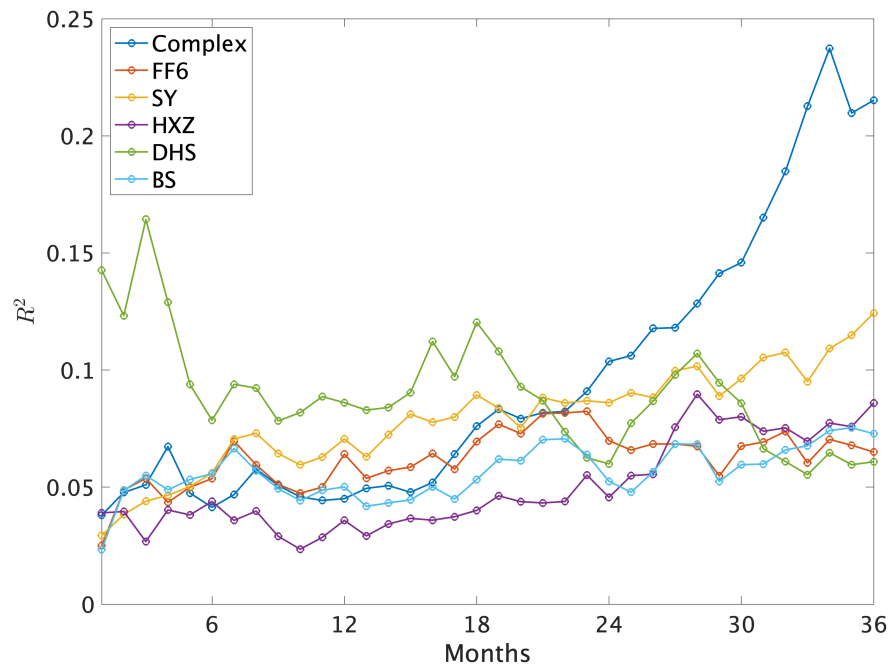


Figure IA.8: SDF Returns and Future Macroeconomic Outcomes—All SDFs

Note. This figure reports the R^2 from a regression of SDF returns on future macroeconomic outcomes per equation (15).

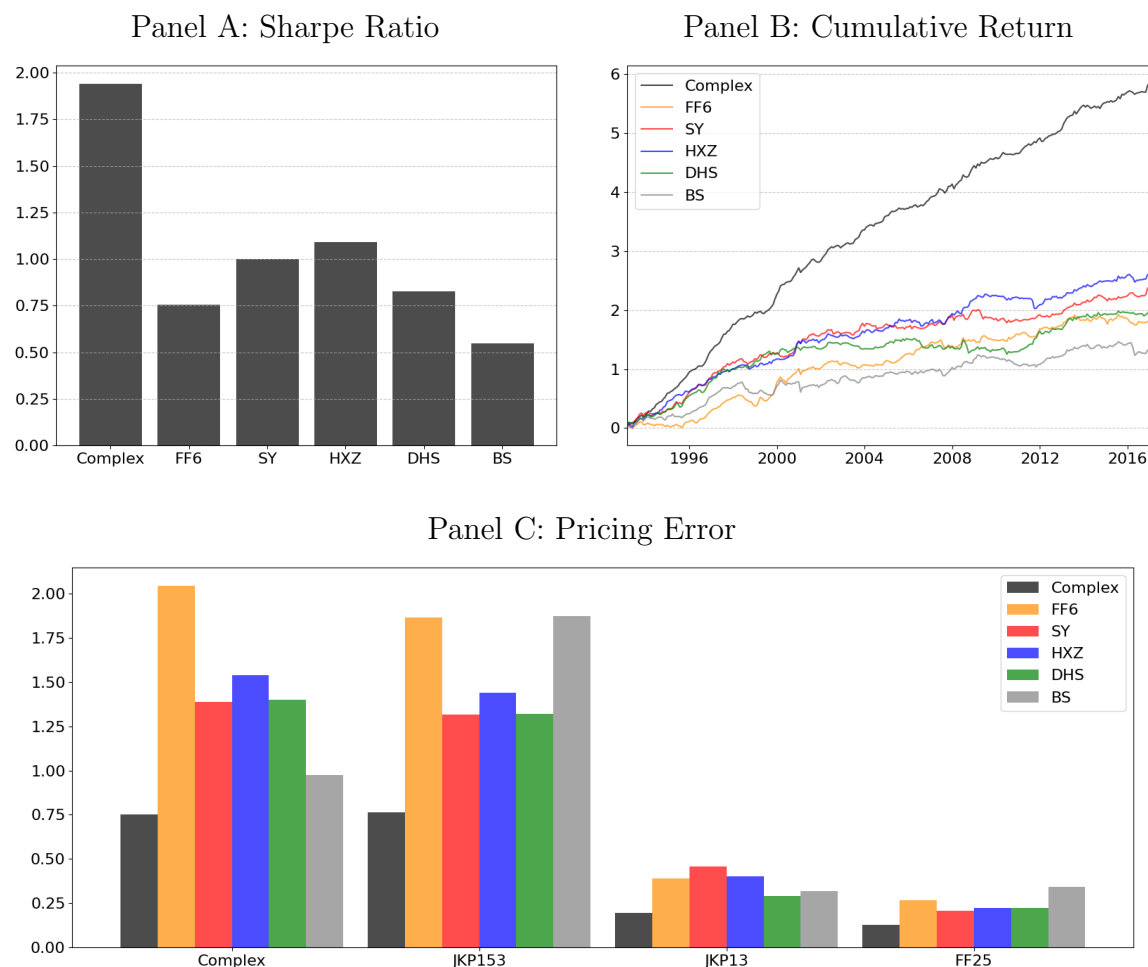


Figure IA.9: Performance Comparison of Complex and Benchmark Factor Models for $T = 12$

Note. This figure replicates Figure 3, replacing the original $T = 360$ month rolling window with $T = 12$ months for the estimation of both the complex model and the simple benchmark models. All other aspects of the analysis remain unchanged.

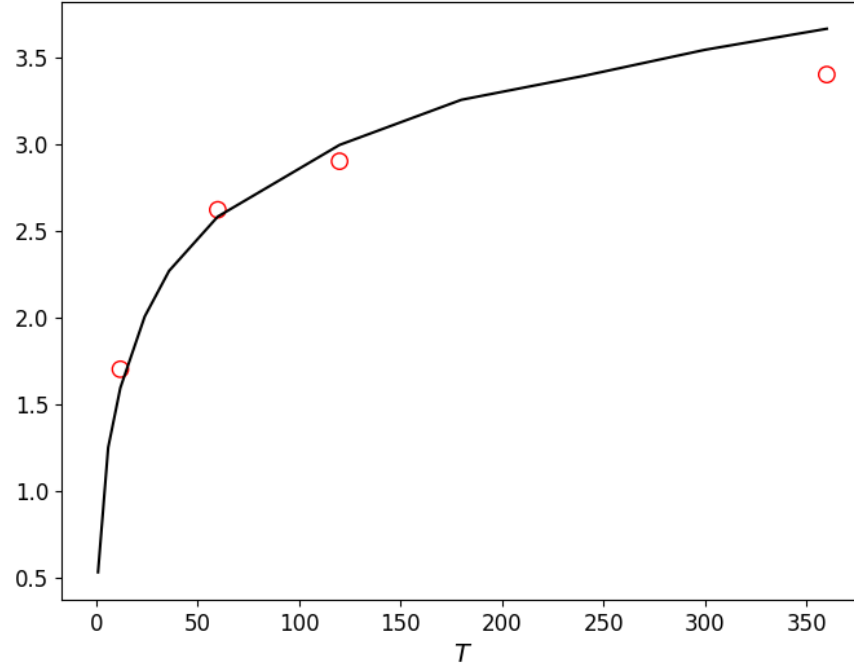


Figure IA.10: Complex SDF Performance for Different T (Data Versus Simulation)

Note. This figure compares the empirical Sharpe ratio of the complex SDF (reported in Section 3) to that of the complex SDF estimated on simulated data. We simulate from the AIPT model calibration described in the theoretical analysis of Section 4.7.1 and estimate the complex SDF using training samples of $T = 1$ to $T = 360$. The solid line reports out-of-sample Sharpe ratios in simulated data for each T . Red circles show the actual empirical out-of-sample Sharpe ratio in training windows of $T = 12, 60, 120$, and 360 (corresponding to the empirical results in Figures 2 and 11). Because our model calibration uses $c = 10$, we report empirical results for $c = 10$.