

Abstract of “Improving Portfolio Construction through Side Information and Quantified Uncertainty” by Anish R. Shah, Ph.D., Brown University, May 2023.

Motivated largely by the problems of estimation error in investment portfolio optimization, this dissertation develops techniques to improve covariance forecasting and portfolio construction. Covariance forecasting is improved by incorporating contemporaneous side information and by propagating uncertainty of the parameters in factor-modeled covariance to uncertainty of the covariance matrix itself. The latter makes it possible to approach several investment problems explicitly modeling uncertainty. Markowitz mean-variance portfolio optimization becomes robust optimization with an uncertainty set on portfolio utility. Uncertain risk parity considers deviation from the risk budget. Risk under uncertainty and price movement measures the stability of hedges under forecasting uncertainty and shifts in risk exposure caused by dynamic prices.

Abstract of “Improving Portfolio Construction through Side Information and Quantified Uncertainty” by Anish R. Shah, Ph.D., Brown University, May 2023.

Motivated largely by the problems of estimation error in investment portfolio optimization, this dissertation develops techniques to improve covariance forecasting and portfolio construction. Covariance forecasting is improved by incorporating contemporaneous side information and by propagating uncertainty of the parameters in factor-modeled covariance to uncertainty of the covariance matrix itself. The latter makes it possible to approach several investment problems explicitly modeling uncertainty. Markowitz mean-variance portfolio optimization becomes robust optimization with an uncertainty set on portfolio utility. Uncertain risk parity considers deviation from the risk budget. Risk under uncertainty and price movement measures the stability of hedges under forecasting uncertainty and shifts in risk exposure caused by dynamic prices.

Improving Portfolio Construction through Side Information and Quantified
Uncertainty

by

Anish R. Shah

B.Sc., Tufts University, 1991

M.Sc., UC Berkeley, 1994

M.Sc., Brown University, 2003

A dissertation submitted in partial fulfillment of the
requirements for the Degree of Doctor of Philosophy
in the Division of Applied Mathematics at Brown University

Providence, Rhode Island

May 2023

© Copyright 2023 by Anish R. Shah

This dissertation by Anish R. Shah is accepted in its present form by
the Division of Applied Mathematics as satisfying the dissertation requirement
for the degree of Doctor of Philosophy.

Date _____
Matthew Harrison, Director

Recommended to the Graduate Council

Date _____
Eli Upfal, Reader
(Computer Science)

Date _____
Petter Kolm, Reader
(NYU Courant)

Approved by the Graduate Council

Date _____
Thomas A. Lewis
Dean of the Graduate School

VITA

Anish Shah graduated in 1991 from Tufts University with a B.Sc. double majoring in mathematics and mechanical engineering, in 1994 from the University of California at Berkeley with an M.Sc. in operations research, and in 2003 from Brown University with a M.Sc. in applied mathematics. He has worked in quantitative finance for decades and frequently presents his research in the US and abroad.

ACKNOWLEDGEMENTS

I send my sincere thanks to my dissertation advisor Matt Harrison and readers Petter Kolm and Eli Upfal, to faculty in data science and applied math who made completing the PhD possible after so long away, and to the many applied math professors who made it impossible not to fall in love with such magnificently beautiful material.

CONTENTS

List of Tables	x
List of Figures	xi
1 Introduction	1
2 Literature review	3
2.1 Background – stock returns and portfolio selection	3
2.1.1 Nature of stock returns	3
2.1.2 Portfolio selection	4
2.2 Estimation error in portfolio optimization	6
2.2.1 Frequentist problems	6
2.2.2 Better point-estimates	7
2.2.3 Alternative construction methods	9
2.2.4 Loss criteria for estimators	14
2.3 Research of this dissertation	14
3 Mathematical framework	16
3.1 Model - observations, latent states, side information	16
3.1.1 Gaussian returns	18
3.2 Examples	19
3.2.1 Factor-modeled covariance	19
3.2.2 Multivariate GARCH	20
4 Short-term covariance given instantaneous side information	23
4.1 Literature review	23
4.2 Problem formulation	25
4.3 A particular parameterization	27
4.4 Practical considerations	28

4.5	Abstract	30
4.6	Key Messages	30
4.7	Introduction	30
	4.7.1 Background	31
	4.7.2 General Structure of Factor Covariance Models	32
4.8	Current Information	32
	4.8.1 Portfolio or Individual Security Variance	33
	4.8.2 Expected Cross-Sectional Volatility	34
	4.8.3 Pairwise Correlation between Securities or Portfolios	34
	4.8.4 Average Correlation within a Portfolio	34
4.9	Bayesian Framework for Incorporating Current Information	34
	4.9.1 Parameters to Accommodate State	35
	4.9.2 Forecasts and Errors	36
	4.9.3 Prior	37
4.10	Adapting to Current Market Conditions by Blending Ledoit-Wolf Co- variance Reductions	38
	4.10.1 Expressed as 1-Factor Models	39
	4.10.2 Parametrization	40
4.11	Adapting to Current Market Conditions by Linear and Quadratic Pro- gramming	41
	4.11.1 Work in a Representation with Uncorrelated Factors	41
	4.11.2 Linear Properties	41
	4.11.3 Parameterization	42
	4.11.4 Observation Noise and Prior	43
	4.11.5 MAP Estimate of Parameters	43
	4.11.6 Restrictions for Observable Factors	44
	4.11.7 Casting Prior onto Orthogonal Factors	44
	4.11.8 Reducing Parameters to Prevent Overfitting	45
4.12	Orthogonal Directions that Best Explain Securities	45
4.13	Changing the Time Scale	48
4.14	Expected Cross-Sectional Volatility	50
4.15	Average Correlation and Variance within a Portfolio	52

5	Uncertain covariance and portfolio optimization under uncertainty	55
5.1	Literature and mathematical model	56
5.2	Abstract	58
5.3	Introduction	58
5.3.1	Background	59
5.4	Probability model	60
5.5	Uncertain covariance model	62
5.6	Portfolio variance with orthogonal factors	62
5.7	Robust optimization on portfolio utility	67
5.7.1	Indifference region to reduce transaction costs	68
5.8	An example	68
5.9	Summary	72
5.10	Uncertain factor-modeled $E[\Sigma_{i,j}]$ and $\text{cov}[\Sigma_{i,j}, \Sigma_{m,n}]$	73
5.11	Incorporating structural uncertainty	76
5.12	Mathematical results	77
6	Uncertain risk parity	81
6.1	Literature review	81
6.2	Mathematical model	82
6.3	Abstract	83
6.4	Introduction	83
6.4.1	Literature review	84
6.4.2	Problem statement	85
6.5	Sensitivity to covariance forecast	86
6.5.1	As solution to minimum variance with regularization	86
6.5.2	Non-robustness in risk contribution assignment	87
6.6	Adding uncertainty	87
6.6.1	Risk parity restated with uncertainty	88
6.6.2	Multiple regimes	89
6.7	Uncertain risk decomposition	90
6.8	Factor-modeled covariance uncertainty	91
6.9	An example	92
6.10	Summary	93

6.11	Uncertain covariance between portfolios from securities	95
6.11.1	Model-free	96
6.11.2	Uncertain covariance modeled	96
6.12	Parametric uncertainty covariance of unbiased covariance estimates	98
7	Risk under uncertainty and price movement	102
7.1	Mathematical model	103
7.2	Abstract	104
7.3	Introduction	104
7.3.1	Problem with conventional method	105
7.3.2	Contribution and central idea of this paper	106
7.4	Conventional risk contributions	106
7.4.1	Variance contributions of common partitions	107
7.5	Uncertain variance contributions	109
7.5.1	Uncertain factor-modeled covariance	110
7.5.2	Mean and covariance of contributions	111
7.5.3	Computational complexity	119
7.6	Uncertain standard deviation contributions	119
7.7	Modeling weight dynamics	122
7.8	An example	124
7.8.1	Portfolio 1: Diversified long, cash benchmark	125
7.8.2	Portfolio 2: 44 stock long-short, cash benchmark	127
7.9	Summary	128
7.10	Mathematical results	130
7.11	Mean and variance of factor covariance from uncertain covariance model	132
8	Conclusion	134

LIST OF TABLES

2.1	Robust optimization sets	9
5.1	Realized portfolio std dev quantiles by method (annualized)	70
5.2	Head-to-head comparison of realized volatility	71
5.3	Median expected variance, uncertainty, and realized variance (weekly, not annualized)	71
5.4	Expected minus realized annualized standard deviation	72
5.5	Z-scores: $\frac{\text{predicted} - \text{realized variance}}{\text{uncertainty (std dev) of variance}}$	72
6.1	Experiment 1: Std Dev & Correlation	93
6.2	Experiment 2: Std Dev & Correlation	94
6.3	Experiment 1: Weight and Variance Contributions	94
6.4	Experiment 2: Weight and Variance Contributions	95
7.1	Computational complexity of reporting contributions' mean	119
7.2	Computational complexity of reporting contributions' variance	119
7.3	Diversified long—mean and SD of factor exposures contributed by sector	126
7.4	Diversified long—conventional and uncertain risk contributions by factor	127
7.5	Diversified long—ex-ante modeled and realized risk contributions by sector	127
7.6	44 stock long-short—mean and SD of factor exposures contributed by side	128
7.7	44 stock long-short—mean and SD of factor exposures contributed by sector	129
7.8	44 stock long-short—conventional and uncertain risk contributions by factor	130
7.9	44 stock long-short—ex-ante modeled and realized risk contributions by sector	130

LIST OF FIGURES

3.1 Model.	17
--------------------	----

CHAPTER 1

INTRODUCTION

Quantitative finance is a broad label for the application of mathematical, algorithmic, statistical, or inferential techniques to financial markets. The expertise involved can vary from stochastic processes, PDEs, and optimal control to optimization, econometrics, statistics, machine learning, and computer science.

Attilio Meucci, for example in [Meucci \[2011\]](#), makes a general split between what he calls ‘P’ quants and ‘Q’ quants. Repeating his ideas: ‘P’ quants are concerned with investing money, operate in discrete time and high dimension (e.g., many stocks to invest across), try to predict the future, and use tools such as statistics and optimization. ‘Q’ quants are concerned with trading, operate in continuous time and low dimension, extrapolate the present to price securities that lack market information, and use tools such as stochastic calculus.

This dissertation addresses several ‘P’ quant problems I encountered over many years in the field. The impetus for much of the work is estimation error’s confounding portfolio optimization, the source of enough colossal investment blow-ups, one involving two Nobel laureates and catastrophic loss across financial markets averted only by government intervention¹. Nassim Nicholas Taleb mentions the problem in his 2012 book *Anti-Fragile*². First thinking on it in the mid 2000s, after many failed attempts, I began (to my eye) to see clearly perhaps a decade later. The insight also led to solving a couple of related problems. Progressing on the different pieces over the years, I’ve presented to combined professional/academic audiences in Boston, New York, San Francisco, London, Amsterdam, and globally via video conference.

The four topics in the dissertation are 1–forecasting short-term covariance given instantaneous side information, 2–uncertain covariance and portfolio optimization under uncertainty, 3–uncertain risk parity portfolio construction, and 4–risk under

¹Long-Term Capital Management. See [Jorion \[2000\]](#).

²[Taleb \[2012\]](#). Taleb retweeted one of my presentations in Aug 2020.

uncertainty and price movement. While all are related, only the first didn't arise from the aforementioned estimation error problem.

The chapters are arranged as follows: Chapter 2 is a literature review. Chapter 3 describes a common mathematical framework. Each of chapters 4 – 7 covers one of the four topics. Chapter 8 is a conclusion.

CHAPTER 2

LITERATURE REVIEW

2.1 Background – stock returns and portfolio selection

2.1.1 Nature of stock returns

Let S_t be the price of a stock at time t . The stock's return R_t is its relative change in value, measured as a fraction or as the change in log of value. Ignoring dividends, $R_t = (\frac{d}{dt}S_t)/S_t$ or log return $\frac{d}{dt}\log S_t$ in continuous time and $S_t/S_{t-1} - 1$ or (less common) log return $\log S_t - \log S_{t-1}$ in discrete time. The two canonical models are geometric Brownian motion¹

$$dS_t = \mu_t S_t dt + S_t \sigma_t dW_t \quad (2.1)$$

$$\text{i.e., } d\log S_t = (\mu_t - \sigma_t^2/2) dt + \sigma_t dW_t \quad (2.2)$$

and, to allow for mean reversion, Ornstein-Uhlenbeck²

$$dS_t = -\theta (S_t - \bar{S}) dt + \sigma_t dW_t \quad (2.3)$$

Discrete analogs model returns as

$$\text{independent } R_t = \mu_t + \varepsilon_t \quad (2.4)$$

$$\text{or serially correlated via an AR(1) } R_t = \rho R_{t-1} + \varepsilon_t \quad (2.5)$$

where $\varepsilon_t \sim N(0, \sigma_t^2)$ and ignoring that returns usually can't occur below -100%. With multiple stocks, S, W, ε and discrete μ are vectors and σ, θ, ρ and continuous μ matrices.

¹See section 5.8 of [Karatzas and Shreve \[1991\]](#).

²See [Barndorff-Nielsen and Shephard \[2001\]](#) or [Meucci \[2009\]](#).

While theoretically and practically important variants exist³—for example, allowing for time dependency and jumps in the randomness—the simpler versions often yet appear out of approximation, ease of manipulation, or the limits of forecasting accuracy.

Unlike a physical system, in addition to not adhering to rules, returns have a mechanism which makes them unpredictable: the market efficiency of investor action. Buying raises a stock’s price and selling lowers it, a process that occurs until the opportunity for profiting vanishes. Predictable edges are *arbitraged away*. In particular, this means *the past does not imply the future*. However, covariance persists through structural links; e.g., oil price increases affect oil companies positively and airlines negatively, linking them regardless of whether oil prices can be predicted.

An aspect of returns that can be predicted is *risk premium*, compensation for bearing risk, akin to writing insurance and receiving premiums. Apparently successful strategies are often in truth compensated for bearing risks. When the rare event occurs, profits are disgorged. Finding an apparently profitable strategy, a skillful investor wonders, “Is there a risk I am getting paid for?”

2.1.2 Portfolio selection

Portfolio selection determines the amount of capital to invest in each asset. The canonical methods are [Markowitz \[1952\]](#) (‘mean-variance’ or ‘Markowitz’) and [Kelly \[1956\]](#) (‘Kelly betting’). [Markowitz \[1952\]](#) introduces the notion of efficient portfolios—maximum expected return at their variance and minimum variance at their expected return—and finds them by maximizing expected return minus a scalar times variance.

$$w^*(\mu, \Sigma, \lambda) \equiv \arg \max_w w^T \mu - \frac{\lambda}{2} w^T \Sigma w \quad (2.6)$$

where $\mu \equiv$ expected returns $\in \mathbb{R}^n$

$\Sigma \equiv$ covariance of returns $\in \mathbb{R}^{n \times n}$

$\lambda \equiv$ risk aversion > 0

$w \equiv$ security weights $\in \mathbb{R}^n$

³Beginning with [Mandelbrot \[1963\]](#), Mandelbrot looked at price series’ fractal properties. A recent evolution of that line is [Dandapani et al. \[2021\]](#). [Bouchaud \[2021\]](#) is a clear, short overview of the history and areas of current research.

[Kelly \[1956\]](#) allocates a fraction of the capital present each period to maximize the exponential growth rate or, to same effect, expected log of wealth.

$$\text{exp growth rate} \equiv \lim_{T \rightarrow \infty} \frac{1}{T} \log \frac{S_T}{S_0} \quad (2.7)$$

$$\begin{aligned} S_T &= S_0 \prod_{t=1}^T \frac{S_t}{S_{t-1}} \\ &= S_0 \prod_{t=1}^T \left(\sum_{i=1}^n w_i [1 + R_{i,t}] \right) \\ \mathbb{E} \left[\log \frac{S_T}{S_0} \right] &= \sum_{t=1}^T \mathbb{E} \left[\log \left(\sum_{i=1}^n w_i [1 + R_{i,t}] \right) \right] \end{aligned} \quad (2.8)$$

where $S_t \equiv$ portfolio value at t

$R_{.,t} \equiv$ security returns from $t - 1$ to t

[Breiman \[1961\]](#) and [Thorp \[1971\]](#) prove results about its optimality. Page 122 of [Markowitz \[1959\]](#) shows the two term Taylor expansion of expected exponential growth is approximately mean return - 1/2 variance of return, i.e., Markowitz with a low risk aversion of $\lambda = 1$. [Chamberlain \[1983\]](#) shows that for concave utility and elliptically distributed investment returns, expected utility of a portfolio's (scalar, random) return is a function of mean and variance. [Markowitz \[2014\]](#) is a survey of approximating concave utility with Markowitz.

In practice, Kelly betting is regarded as having too large swings, especially for investing on behalf of a client. Since Markowitz is ubiquitous, this research takes it as a given.

Several areas border the scope of this dissertation. There is literature on utility functions—what utility should be, what properties are advantageous and sensible, how real human beings behave. There is (very technical) literature on optimal control in portfolio selection. Near its head, [Merton \[1971\]](#) proves results on maximizing expected utility when utility is concave and prices are modeled as continuous time Markov stochastic processes. There is literature on alterations to Markowitz. For example, risk measures⁵ in place of variance include mean absolute deviation and semi-variance, $\mathbb{E} [\min (x - target, 0)^2]$. The motivation for the latter is that, at odds with

⁴Notation: $\equiv$ means a definition.

reason, adding variable upside can decrease a portfolio’s utility when risk is measured by variance. To deal with non-Gaussian returns, one approach⁶ is to minimize (CVaR) conditional value-at-risk, $E[x|x \leq F_x^{-1}(c)]$, subject to some minimum $E[x]$. Though Markowitz considers only the present, the problem becomes multi-period in the presence of transaction costs and market impact.⁷

2.2 Estimation error in portfolio optimization

Ideas on estimation error’s effect on optimizing Markowitz utility follow three themes: frequentist problems, better point-estimates (Bayesian, robust statistics), and alternative construction methods (Bayesian posterior distribution of utility, robust optimization).

2.2.1 Frequentist problems

Though no longer so, it was once surprising that optimizing weights across stocks, which are numerous and similar, based on unbiased estimates of a few parameters (mean and covariance) taken as true—with real differences dominated by noise—yields a portfolio that performs worse than expected and a decision rule that is easily improved.

[Jobson and Korkie \[1980\]](#) derives an approximate distribution of the mean and variance of a portfolio optimized from frequentist estimates and finds that they work poorly with small histories and would be better replaced by Bayesian estimates. [Michaud \[1989\]](#) brought the problem to a broad audience. [Broadie \[1993\]](#) simulates returns from known parameters, infers unbiased mean and covariance, and optimizes the efficient frontier on both to show that optimizing estimates goes badly. A few lines proof of the mechanism is in [Shah \[2009\]](#). [Fan et al. \[2008\]](#) analyzes properties of factor-modeled⁸ vs sample covariance and finds factor-modeling is advantageous when inverse of covariance is of interest as in Markowitz portfolio optimization.

⁵See chapter 9 of [Markowitz \[1959\]](#), [Konno and Yamazaki \[1991\]](#), [Chang et al. \[2009\]](#).

⁶See [Rockafellar and Uryasev \[2000\]](#), [Rockafellar and Uryasev \[2002\]](#).

⁷See [Gârleanu and Pedersen \[2013\]](#).

⁸The standard ways of factor-modeling covariance—time series, cross-sectional, PCA—are described in chapter 9 of [Tsay \[2010\]](#).

Although there are more papers, there is little more to say once it's clear why frequentist estimates don't work well with optimization.

2.2.2 Better point-estimates

[Fabozzi et al. \[2007\]](#) gives an overview of the issues, and [Fabozzi et al. \[2010\]](#) is a survey largely on approaches to counter estimation error issues in mean-variance portfolio construction. From the latter:

There are two standard methods extensively adopted in the literature to deal with estimation errors. The first is the robust estimation methods, such as moment estimation, which can be quite robust to distributional assumptions. The second is the Bayesian approach that is neutral to uncertainty in the sense of Knight (1921) because it assumes a single prior on the portfolio distribution ([Garlappi et al. 2007](#)).

Bayesian point-estimates

First, two techniques that are widely used in practice: [Black and Litterman \[1992\]](#) converts forecasts of μ from unbiased to Bayesian by assuming Gaussian observation noise and prior. The computation is the same as a Kalman filter with one noisy observation and no dynamics. [Ledoit and Wolf \[2003\]](#) and other papers by the authors shrink covariance toward a structured form—usually one factor (effectively the first PCA) + a diagonal matrix—via convex combination and derive results for the optimal shrinkage.

Notice that, in real life, μ and Σ are forecasted separately. Again, this is because Σ is largely from history and μ from diverse outside information. The many Bayesian papers that forecast both from history ignore the predictability-erasing mechanism of the market.

Other covariance shrinkage estimators in the style of Ledoit and Wolf are [Daniels and Kass \[2001\]](#)⁹ and [Schäfer and Strimmer \[2005\]](#). Continuing in that direction, a set of shrinkage point-estimators of covariance comes from random matrix theory. The general idea: eigenvalues of the sample correlation matrix are biased—upward for

⁹Cites a 1996 working paper by Ledoit.

large eigenvalues and downward for small—so are shrunk toward a point or distribution (Marčenko–Pastur) or discarded below a critical value with their effects diagonalized. [Plerou et al. \[2002\]](#) and [Potters et al. \[2005\]](#) are early papers, [Johnstone \[2007\]](#) a clear introduction, [Bun et al. \[2017\]](#) a survey, and [Potters and Bouchaud \[2020\]](#) a textbook. [Ledoit and Wolf \[2012\]](#), which shrinks eigenvalues toward a distribution, appears in several multivariate ARCH papers mentioned in chapter 4 below.

Robust statistics

[Scutellà and Recchia \[2013\]](#), a survey on robust optimization, also covers robust statistics¹⁰ and cites the methods of the following two paragraphs.

Robust point-estimates of mean¹¹ and covariance are used to optimize portfolios in [Cavadini et al. \[2001\]](#)¹², [Perret-Gentil and Victoria-Feser \[2004\]](#)¹³, [de Melo Mendes and Leal \[2005\]](#)¹⁴, and [Welsch and Zhou \[2007\]](#)¹⁵.

Rather than using robust estimates of parameters, several papers optimize utility defined as a robust function of weighted security returns. Note that these methods are impractical since future expected returns must be forecasted and not inferred from history. [Lauprete et al. \[2003\]](#) minimizes a robust estimate of portfolio variance, the Huber function. [DeMiguel and Nogales \[2009\]](#) explores two ways of optimizing a robust estimate of portfolio variance involving both location and scale: 1) M-estimators via the Huber function, 2) S-estimators via Tukey’s biweight function.

¹⁰See [Huber \[1981\]](#).

¹¹Suffers from the usual problem—predicting mean requires outside information.

¹²Robust estimator from [Krishnakumar and Ronchetti \[1997\]](#)

¹³Robust estimator from [Rocke \[1996\]](#)

¹⁴Robust estimator from [Leung and Ng \[2004\]](#) and a variation on minimum covariance determinant of [Rousseeuw \[1985\]](#)

¹⁵1-FAST-MCD from [Rousseeuw and Driessen \[1999\]](#), 2-Iterated Bivariate Winsorization from [Chilson et al. \[2006\]](#) with modification of [Maronna and Zamar \[2002\]](#), 3-Fast 2-D Winsorization from [Khan et al. \[2007\]](#) with modification of [Maronna and Zamar \[2002\]](#).

2.2.3 Alternative construction methods

Robust optimization

Robust optimization finds the best worst case when parameters aren't fixed but fall within a set. The optimization changes from $\arg \max_w f_\theta(w)$ to $\max_w \min_{\theta \in S} f_\theta(w)$. Though the idea is straightforward, the max-min optimization limits the technique to special configurations, thus the literature is largely an exploration of particular cases described in Table 2.1.

Table 2.1: Robust optimization sets

Reference	Set
Rustem et al. [2000]	(μ, Σ) comes from a number of scenarios, $\cup_{i=1\dots n} (\mu_i, \Sigma_i)$
Halldórsson and Tütüncü [2003]	$\mu^L \leq \mu \leq \mu^U, \Sigma^L \leq \Sigma \leq \Sigma^U$
Goldfarb and Iyengar [2003]	$\mu^L \leq \mu \leq \mu^U$ $\Sigma = EFE^T + \text{diag}(\varepsilon^2)$ where $\varepsilon_L^2 \leq \varepsilon^2 \leq \varepsilon_U^2$ $\{E : E = \bar{E} + \eta, \ \eta_{i,\cdot}\ _G \leq \delta_i\}$ norm is wrt some positive definite matrix G $\{F : F^{-1} = \bar{F}^{-1} + \eta, \max \text{eigenval of } \bar{F}^{1/2} \eta \bar{F}^{1/2} \leq \delta_F\}$ or $\{F : F = \bar{F} + \eta, \ N^{-1/2} \eta N^{-1/2}\ \leq \delta_F\}$ where norm is max eigenval or sum of squared eigenvals
Ceria and Stubbs [2006]	$\{\mu : (\mu - \bar{\mu})^T \Omega (\mu - \bar{\mu}) \leq \delta\}$
Garlappi et al. [2007]	μ is in the intersection of ellipses on subsets of securities, $\cap_{i=1\dots n} \{\mu : (\mu^{S_i} - \bar{\mu}_i)^T \Omega^{S_i} (\mu^{S_i} - \bar{\mu}_i) \leq \delta_i\}$. x^{S_i} means the S_i coordinates of x . Expected returns that come from factors can be put in this form.
Lu [2011]	defined jointly on μ and factor exposures

[Anderson and Cheng \[2016\]](#) in section 2.2.3 below solves robust optimization with uncertainty in means. [Lu \[2011\]](#) makes the interesting point that separable uncertainty sets, i.e., ones defined as products, are too conservative and operate at a much higher confidence level than reported.

Note that while confidence sets defined separately on each parameter ignore the diversification of error and see every extreme outcome realized at once, jointly defined confidence sets become informationless as the universe grows. For example, imagine Gaussian μ frequentist estimated as a vector of 0's with identity noise covariance. A confidence set is a sphere of radius h . Since a single coordinate outside h breaks the sphere, the critical h for a given confidence level increases with dimension. These problems vanish in a Bayesian setting.

Previously mentioned [Scutellà and Recchia \[2013\]](#) is a survey paper on robust optimization of portfolios, and [Bertsimas et al. \[2011\]](#) on robust optimization in general. [Hansen and Sargent \[2011\]](#) explores robust decision making/control in depth.

Bayesian posterior

Most of the ideas were introduced in early, universally cited work. [Barry \[1974\]](#) assumes (μ, Σ) is Gaussian-Inverse Wishart whose parameters are the sample estimates of μ and Σ . This corresponds¹⁶ to Bayesian inference with the prior $\pi(\mu, \Sigma) \propto |\Sigma|^{-\frac{n+1}{2}}$ and Gaussian returns. The mean of the (multivariate t) predictive distribution of returns is the sample mean, and the variance a scalar multiple of the sample variance. Thus, mean-variance utility is the same as frequentist-inferred under higher risk aversion. [Klein and Bawa \[1976\]](#) considers three priors, two of which lead to different choices than can be had from frequentist estimates. [Frost and Savarino \[1986\]](#) sets the prior's parameters empirically; the prior is Gaussian-Inverse Wishart with all securities having the same mean return, variance, and pairwise correlation.

¹⁶Derived in [Klein and Bawa \[1976\]](#)

[Jorion \[1986\]](#) starts by evaluating the statistical risk of an estimator in terms of utility under optimized weights,

$$\begin{aligned}
 U(w, \theta) &\equiv \text{(random) utility of portfolio } w \text{'s returns under parameters } \theta \\
 w^*(\theta) &\equiv \arg \max_w E[U(w, \theta)] \\
 U^*(\theta) &\equiv U(w^*[\theta], \theta) \\
 r(\theta) &\equiv \text{risk of estimator} \\
 &\equiv E \left[\frac{U^*(\theta) - U(w^*[\hat{\theta}], \theta)}{|U^*(\theta)|} \right]
 \end{aligned} \tag{2.9}$$

and finds frequentist estimates of μ inadmissible. The paper then explores Stein shrinkage estimators¹⁷, of which a particular case can be framed as an informative prior on μ (Gaussian centered at a common mean). With variance assumed proportional to the sample variance, the prior's common mean and variance concentration scalar are determined empirically.

[Avramov and Zhou \[2010\]](#) is a survey paper.

Recent work leans toward specifics. [Anderson and Cheng \[2016\]](#) models next period returns as a mixture of normals having inverse-Wishart variance. The models correspond to the length of the current regime: at period t , there are t normals, with the i 'th regarding returns from periods $i \dots t$ as coming from the same distribution. Since mean returns come from history, it is not practical. In addition to maximizing posterior utility, the paper solves a robust version of the problem ("risk-sensitive optimization" in control theory) that considers uncertainty in mean returns. [Kacperczyk and Damien \[2011\]](#) is unusual in that it considers 'distribution uncertainty', the possibility of returns' following distributions other than Gaussian. [Lai et al. \[2011\]](#), noted here only because the authors are from statistics and applied math departments, looks at the mean and variance of the predictive distribution of returns. [Bauder et al. \[2021\]](#) is likewise a straight Bayesian treatment using the mean and variance of the predictive distribution assuming the series of returns is exchangeable, the future is like the past, and the number of observations exceeds the number of securities. The latter two assumptions make it practically unusable, but published in the last two

¹⁷See [Stein \[1956\]](#)

years in a credible journal, the paper offers a current snapshot of references.

Two techniques that don't forecast returns from history: [Davis and Lleo \[2013\]](#) extends [Black and Litterman \[1992\]](#) and applies the continuous time Kalman filter to incorporate information arriving across time.¹⁸ In [Meucci \[2008\]](#) and other papers, Meucci lays out a framework, 'entropy pooling', for updating general distributions with forecast information. The starting point is a prior on the world of interest, $X \sim f_X$. Views are expressed as distributions on functionals, $g(X) \sim f_V$. A distribution f^* is sought¹⁹ that is coherent with the views and has minimum K-L divergence from the prior. This is convexly combined with the prior to express the forecaster's confidence, $f_{forecast} = cf^* + (1 - c)f_X$.

Resampling

[Michaud \[1998\]](#), which is critiqued by [Scherer \[2002\]](#), introduces 'resampling'—repeatedly generating a fixed length sample of a normal distribution under the point-estimated parameters and optimizing using parameters inferred from the sample, finally averaging the output. Particulars aside, it is akin to a softened point-estimate of covariance: Because the maximizer of $w^T \mu - \frac{\lambda}{2} w^T \Sigma w$ is $w^* = \frac{1}{\lambda} \Sigma^{-1} \mu$ and samples $(\hat{\Sigma}, \mu + \hat{\delta})$ and $(\hat{\Sigma}, \mu - \hat{\delta})$ are equally likely, $E[\hat{w}^*] = \frac{1}{\lambda} E[\hat{\Sigma}^{-1}] \mu$.²⁰

Risk parity

Named by [Qian \[2005\]](#), risk parity is a popular portfolio construction technique that determines weights by making each investment's risk contribution (component in the direction of the net portfolio risk²¹) match the given, typically equally-weighted risk

¹⁸See [Kolm and Ritter \[2017\]](#), [Kolm and Ritter \[2020\]](#), and [Kolm et al. \[2021\]](#) for other extensions.

¹⁹Cannot be solved analytically but numerically by reweighting the probabilities on a fixed set of samples.

²⁰With linear equality constraints, sampling μ washes out for the same reason: $\arg \max_w w^T \mu - \frac{\lambda}{2} w^T \Sigma w$ s.t. $Aw = b$ is $w^* = \frac{1}{\lambda} \left[\Sigma^{-1} - \Sigma^{-1} A^T (A \Sigma^{-1} A^T)^{-1} A \Sigma^{-1} \right] \mu + \Sigma^{-1} A^T (A \Sigma^{-1} A^T)^{-1} b$,
so $E[\hat{w}^*] = \frac{1}{\lambda} E \left[\hat{\Sigma}^{-1} - \hat{\Sigma}^{-1} A^T (A \hat{\Sigma}^{-1} A^T)^{-1} A \hat{\Sigma}^{-1} \right] \mu + E \left[\hat{\Sigma}^{-1} A^T (A \hat{\Sigma}^{-1} A^T)^{-1} \right] b$.

²¹See [Menchero and Davis \[2011\]](#).

budget.²²

$$\begin{aligned}
 w^* (\Sigma, \sigma^2, f) &\equiv \text{any } w > 0 \text{ s.t. } w \odot \Sigma w = \sigma^2 f & (2.10) \\
 \text{where } \sigma^2 &\equiv \text{target variance of portfolio} \\
 f &\equiv \text{fraction of risk budgeted to each security, } f \in \mathbb{R}_+^n \\
 \odot &\equiv \text{element-by-element multiplication}
 \end{aligned}$$

It is typically done at the index or asset class level rather than on individual stocks, e.g., one security is the S&P500 index, another US corporate bonds.

Maillard et al. [2010] shows risk parity is equivalent²³ to minimizing variance with a regularization penalty.

$$w^* (\Sigma, \sigma^2, f) = \arg \min_{w > 0} \frac{1}{2} w^T \Sigma w - \sigma^2 \sum_i f_i \log w_i \quad (2.11)$$

In addition to much simpler computation, strict convexity makes obvious the existence of a unique solution in each orthant (using $\log |\cdot|$) when w isn't sign constrained.

Maillard et al. [2010] and Bai et al. [2016] explore mathematical properties and alternative formulations. Instead of equalizing risk contributions across investments, Lohre et al. [2014] equalizes them across principal component directions of the covariance matrix, Lassance et al. [2022] across independent component directions, and Roncalli and Weisang [2016] across directions specified by any factor model. Bernardi et al. [2019] looks at risk parity's 'second order risk', the ratio of realized to predicted risk, essentially the bias from estimate-optimized selection.

Hierarchical risk parity

De Prado [2016] introduces 'hierarchical risk parity' as a way to avoid covariance's low eigenvalue directions (seemingly low risk portfolios) that appear as artifacts under a large universe. The algorithm 1—constructs a notion of distance between securities based on differences in their correlations to all securities, 2—extracts a tree structure through hierarchical clustering, 3—reorders securities according to proximity (a depth

²²Controlling the scale via full investment ($1^T w = 1$) instead of target variance σ^2 is achieved easily by arbitrarily setting σ^2 and normalizing the solution, $w^*/1^T w^*$.

²³Setting the gradient to 0 yields the risk parity equation (2.10).

first search of the tree), and finally, 4–starting from a single group containing all securities equally weighted, repeatedly bisects each group into a finer scale (groups containing $1/2$ as many securities) and scales the weights on each half inversely proportional to its relative variance. Cotton [2022b] generalizes the idea.

2.2.4 Loss criteria for estimators

In equation (2.9) above, Jorion [1986] evaluates an estimator by the utility of a portfolio built from its estimates. Facing Markowitz utility as in equation (2.6), why do few papers estimate $\{\mu, \Sigma\}$ and set w^* to maximize the performance of $U\left(w^* \left[\hat{\mu}, \hat{\Sigma}\right], \mu, \Sigma\right)$, which is standard in statistics?²⁴

As previously mentioned, market efficiency dictates that forecasts for μ and Σ come from different information sets, with μ 's continually changing and expanding. Nevertheless, utility considerations appear implicitly. For example, factor-modeling and the many covariance shrinkage estimators in section 2.2.2 aim to erase low eigenvalue directions of covariance *because* the maximizer of $w^T \mu - \frac{\lambda}{2} w^T \Sigma w$ is $w^* = \frac{1}{\lambda} \Sigma^{-1} \mu$. This wouldn't be important if w were independent of the estimates. The widely popular Black and Litterman [1992] sets the prior on μ to stabilize the results of optimization. And most papers include empirical results of optimized portfolios.

Papers that explicitly tie the estimation to utility, Jorion [1986] included, typically infer both μ and Σ from history in a traditional framing that doesn't work with markets.²⁵ However interesting, they can't directly be used in practice.

2.3 Research of this dissertation

From this background arise the four problems involving covariance, portfolio construction, and uncertainty that are addressed by this dissertation:

1. Short-term covariance given instantaneous side information
2. Uncertain covariance and portfolio optimization under uncertainty

²⁴In this vein, presenting a deep learning application to trading which didn't involve different information sets, (Pelger [2023]) often repeated the importance of making the loss function match one's objectives.

²⁵Cotton [2022a] presents many arguments for market information over forecasting from history.

3. Uncertain risk parity
4. Risk under uncertainty and price movement

Additional references specific to each topic are in the corresponding chapters.

CHAPTER 3

MATHEMATICAL FRAMEWORK

An abstract mathematical framework serves as starting point for conceptualization and finding approaches to solutions.

3.1 Model - observations, latent states, side information

$Y_t \equiv$ random stock returns ($Y_t \in \mathbb{R}^n$, $t \in \mathbb{N}$)

$Z_t \equiv$ unobserved latent states ($Z_t \in \mathcal{Z}$, $t \in \mathbb{N}$)

$X_t \equiv$ random side information ($X_t \in \mathcal{X}$, $t \in \mathbb{N}$)

Of interest is the distribution given information at time T of future stock returns, $P(Y_T|X_{1:T}, Y_{1:T-1})$ or $P(Y_T|X_{1:T-1}, Y_{1:T-1})$, or latent states, $P(Z_T|X_{1:T}, Y_{1:T-1})$ or $P(Z_T|X_{1:T-1}, Y_{1:T-1})$.

Assume Y and X are independent given Z ,

$$P(Y_{1:T}, X_{1:T}|Z_{1:T}) = P(Y_{1:T}|Z_{1:T}) P(X_{1:T}|Z_{1:T}),$$

and conditionally depend on Z only at the same period,

$$P(Y_{1:T}|Z_{1:T}) = \prod_{t=1}^T P(Y_t|Z_t)$$

$$P(X_{1:T}|Z_{1:T}) = \prod_{t=1}^T P(X_t|Z_t).$$

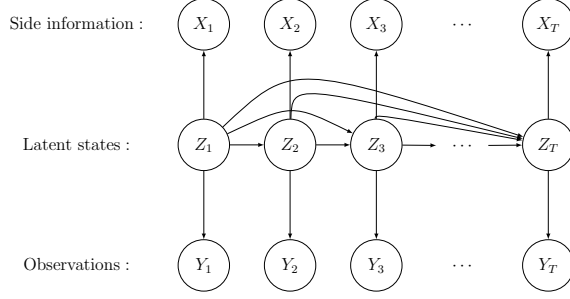


Figure 3.1: Model.

$P(Z_{1:T})$ describes the dynamics of Z .

$$\begin{aligned}
 P(Y_T | X_{1:T}, Y_{1:T-1}) &= \int P(Y_T, Z_{1:T} | X_{1:T}, Y_{1:T-1}) dZ_1 \dots dZ_T \\
 &= \int P(Y_T | X_{1:T}, Y_{1:T-1}, Z_{1:T}) P(Z_{1:T} | X_{1:T}, Y_{1:T-1}) dZ_1 \dots dZ_T \\
 &= \int P(Y_T | Z_T) P(Z_{1:T} | X_{1:T}, Y_{1:T-1}) dZ_1 \dots dZ_T \\
 &= \frac{1}{P(X_{1:T}, Y_{1:T-1})} \int P(Y_T | Z_T) P(X_{1:T}, Y_{1:T-1} | Z_{1:T}) P(Z_{1:T}) dZ_1 \dots dZ_T \\
 &= \frac{1}{P(X_{1:T}, Y_{1:T-1})} \int \left[\prod_{t=1}^T P(X_t | Z_t) P(Y_t | Z_t) \right] P(Z_{1:T}) dZ_1 \dots dZ_T
 \end{aligned} \tag{3.1}$$

Similarly,

$$P(Y_T | X_{1:T-1}, Y_{1:T-1}) = \frac{1}{P(X_{1:T-1}, Y_{1:T-1})} \int \left[\prod_{t=1}^{T-1} P(X_t | Z_t) P(Y_t | Z_t) \right] P(Y_T | Z_T) P(Z_{1:T}) dZ_1 \dots dZ_T \tag{3.2}$$

$$\begin{aligned}
 P(Z_T | X_{1:T}, Y_{1:T-1}) &= \int P(Z_{1:T} | X_{1:T}, Y_{1:T-1}) dZ_1 \dots dZ_{T-1} \\
 &= \frac{P(X_T | Z_T)}{P(X_{1:T}, Y_{1:T-1})} \int \left[\prod_{t=1}^{T-1} P(X_t | Z_t) P(Y_t | Z_t) \right] P(Z_{1:T}) dZ_1 \dots dZ_{T-1}
 \end{aligned} \tag{3.3}$$

$$P(Z_T | X_{1:T-1}, Y_{1:T-1}) = \frac{1}{P(X_{1:T-1}, Y_{1:T-1})} \int \left[\prod_{t=1}^{T-1} P(X_t | Z_t) P(Y_t | Z_t) \right] P(Z_{1:T}) dZ_1 \dots dZ_{T-1} \tag{3.4}$$

Covariance of $Y_T|X_{1:T-1}, Y_{1:T-1}$ can be written in terms of the state variables:

$$\begin{aligned}
\text{cov}[Y_T|X_{1:T-1}, Y_{1:T-1}] &= \text{E} \left[\underbrace{\text{cov}[Y_T|X_{1:T-1}, Y_{1:T-1}, Z_T]}_{\text{cov}[Y_T|Z_T]} | X_{1:T-1}, Y_{1:T-1} \right] \\
&\quad + \text{cov} \left[\underbrace{\text{E}[Y_T|X_{1:T-1}, Y_{1:T-1}, Z_T]}_{\text{E}[Y_T|Z_T]} | X_{1:T-1}, Y_{1:T-1} \right] \\
&= \text{E}[\text{cov}[Y_T|Z_T] | X_{1:T-1}, Y_{1:T-1}] \\
&\quad + \text{cov}[\text{E}[Y_T|Z_T] | X_{1:T-1}, Y_{1:T-1}]
\end{aligned} \tag{3.5}$$

3.1.1 Gaussian returns

$Y_t|Z_t \sim \mathcal{N}(\mu_t(Z_t), \Sigma_t(Z_t))$ is a canonical model of stock returns, with μ_t independent of $\{Y_s, \mu_s, \Sigma_s\}_{s < t}$ through market efficiency. μ is typically close to 0 and of smaller magnitude than the standard deviation.

3.2 Examples

The specific cases below illustrate how the model accommodates standard techniques for working with investment returns.

3.2.1 Factor-modeled covariance

This is the standard model throughout this dissertation. Next period's return covariance is point-estimated from history, $\text{cov}[Y_T|X_{1:T-1}, Y_{1:T-1}]$, by assuming returns follow a linear regression model and computing the covariance of the modeled returns with assumptions on independence.

$$F_t \equiv \text{factors} \in \mathbb{R}^k$$

$$\beta \equiv \text{exposures} \in \mathbb{R}^{n \times k}$$

$$\varepsilon_t \equiv \text{idiosyncratic effects} \in \mathbb{R}^n$$

$$\Omega \equiv \text{covariance of factors} \in \mathbb{R}^{k \times k}$$

$$\Lambda \equiv \text{covariance of idiosyncratic effects} \in \mathbb{R}^{n \times n}, \text{ a diagonal matrix}$$

$$Y_t = \beta F_t + \varepsilon_t$$

$$F_t|\Omega \stackrel{\text{iid}}{\sim} \mathcal{N}(0, \Omega), \text{ ind of } \{\beta, \varepsilon\}$$

$$\varepsilon_t|\Lambda \stackrel{\text{iid}}{\sim} \mathcal{N}(0, \Lambda), \text{ ind of } \{\beta, F\}$$

$$Z_t \equiv \{\beta, \Omega, \Lambda, F_t\}$$

$$P(Z_1, \dots, Z_T) = \pi(\beta, \Omega, \Lambda) \prod_{t=1}^T P(F_t|\Omega)$$

$$Y_t|Z_t \sim \mathcal{N}(\beta F_t, \Lambda)$$

$$Y_t|Z_t \setminus F_t \sim \mathcal{N}(0, \beta \Omega \beta^T + \Lambda)$$

$$X_t \equiv \{F_t\}$$

$$X_t|Z_t \sim \delta_{F_t}$$

As in equation (3.5) but conditioning on $Z_T \setminus F_T$,

$$\begin{aligned}
\text{cov}[Y_T | X_{1:T-1}, Y_{1:T-1}] &= \text{E} \left[\underbrace{\text{cov}[Y_T | X_{1:T-1}, Y_{1:T-1}, Z_T \setminus F_T]}_{\beta \Omega \beta^T + \Lambda} | X_{1:T-1}, Y_{1:T-1} \right] \\
&\quad + \text{cov} \left[\underbrace{\text{E}[Y_T | X_{1:T-1}, Y_{1:T-1}, Z_T \setminus F_T]}_{\beta \text{E}[F_T | Z_T \setminus F_T] = 0} | X_{1:T-1}, Y_{1:T-1} \right] \\
&= \text{E} [\beta \Omega \beta^T + \Lambda | X_{1:T-1}, Y_{1:T-1}] \tag{3.6}
\end{aligned}$$

Versions with more complicated econometrics have non-iid ε_t , e.g., following an ARMA, and β, Ω , and Λ varying with t according to some dynamics.

3.2.2 Multivariate GARCH

Multivariate GARCH(p,q) models time series with heteroskedasticity. The general form:

$$\begin{aligned}
Y_t | \mu_t, \Sigma_t &\sim \mathcal{N}(\mu_t, \Sigma_t) \\
\Sigma_t &= f(Y_{t-p} Y_{t-p}^T, \dots, Y_{t-1} Y_{t-1}^T) + g(\Sigma_{t-q}, \dots, \Sigma_{t-1})
\end{aligned}$$

μ_t is usually fixed ($\mu_t = \mu$) or zero. Different assumptions on f and g lead to numerous acronym-named variants¹. It is also possible to incorporate side information. Among the many versions of GARCH, here are two to show how it fits within the model above.

¹Silvennoinen and Teräsvirta [2009] is an overview.

Vanilla Multivariate GARCH(1,1)

To limit the number of parameters, structure needs to be imposed. One form is GARCH(1,1) in the ‘BEKK’ model of [Engle and Kroner \[1995\]](#):

$$Y_t | \mu, \Sigma_t \sim \mathcal{N}(\mu, \Sigma_t)$$

$$\Sigma_t = aa^T + \sum_{i=1}^K b_i (Y_{t-1} - \mu)(Y_{t-1} - \mu)^T b_i^T + \sum_{i=1}^K c_i \Sigma_{t-1} c_i^T$$

$$\text{where } a, b_i, c_i \in \mathbb{R}^{n \times n}, \quad i = 1 \dots K$$

$$Z_t = \{\mu, a, b, c, \Sigma_t, Y_t\}$$

$$P(Z_1, \dots, Z_T) = \pi(\mu, a, b, c, \Sigma_1) \prod_{t=2}^T P(\Sigma_t | \Sigma_{t-1}, Y_{t-1}, a, b, c) \prod_{t=1}^T P(Y_t | \mu, \Sigma_t)$$

$$Y_t | Z_t \sim \delta_{Y_t}$$

$$X_t \equiv \{ \}$$

Factor model with GARCH

GARCH can be applied within a factor model, for example, GARCH(1,1) in [Engle et al. \[1990\]](#) and [Vrontos et al. \[2003\]](#). From [Vrontos et al. \[2003\]](#):

$$\begin{aligned} F_t &\equiv \text{unobserved factors} \in \mathbb{R}^k \\ \beta &\equiv \text{exposures} \in \mathbb{R}^{n \times k} \\ \Theta_t &\equiv \text{dynamic factor variances} \in \mathbb{R}^k \\ a, b, c &\equiv \text{parameters} \in \mathbb{R}^k \end{aligned}$$

$$\begin{aligned} Y_t &= \mu + \beta F_t \\ F_t &= \text{diag}[\Theta_t]^{\frac{1}{2}} \times \text{iid } \mathcal{N}[0, I_k] \\ \Theta_{i,t} &= a_i + b_i F_{i,t-1}^2 + c_i \Theta_{i,t-1} \\ F_t | \Theta_t &\sim \mathcal{N}(0, \text{diag}[\Theta_t]) \end{aligned}$$

$$\begin{aligned} Z_t &\equiv \{\mu, \beta, a, b, c, \Theta_t, F_t\} \\ P(Z_1, \dots, Z_T) &= \pi(\mu, \beta, a, b, c, \Theta_1) P(F_1 | \Theta_1) \\ &\quad \prod_{t=2}^T P(\Theta_t | \Theta_{t-1}, F_{t-1}, a, b, c) P(F_t | \Theta_t) \end{aligned}$$

$$\begin{aligned} Y_t | Z_t \setminus F_t &\sim \mathcal{N}[\mu, \beta \text{diag}[\Theta_t] \beta^T] \\ X_t &\equiv \{ \} \end{aligned}$$

$$\begin{aligned} \text{cov}[Y_T | Y_{1:T-1}] &= \text{E}[\text{cov}[Y_T | \mu, \beta, \Theta_T, Y_{1:T-1}] | Y_{1:T-1}] \\ &\quad + \text{cov} \left[\underbrace{\text{E}[Y_T | \mu, \beta, \Theta_T, Y_{1:T-1}]}_{\mu} | Y_{1:T-1} \right] \\ &= \text{E}[\beta \text{diag}[\Theta_T] \beta^T | Y_{1:T-1}] + \text{cov}[\mu | Y_{1:T-1}] \end{aligned}$$

CHAPTER 4

SHORT-TERM COVARIANCE GIVEN INSTANTANEOUS SIDE INFORMATION

Covariance of stock returns varies over time. Volatility tends to increase abruptly and fall slowly.¹ For example, a chart of the S&P 500 option-implied volatility VIX looks different forward (impenetrable vertical faces, gradual drops) than backward (gradual increases, cliffs). Correlations between securities also shift. However unpredictable, these changes are apparent in sources of side information almost immediately, long before behavior can be inferred from observing the process alone.

This chapter describes a Bayesian framework for incorporating side information to forecast return covariance over a given horizon (e.g., one or two weeks) and frequency (e.g., daily or weekly returns). One particular parameterization is computationally fast, solved by linear or quadratic programming depending on distribution assumptions.

4.1 Literature review

Originating from [Engle \[1982\]](#) and [Bollerslev \[1986\]](#), ARCH/GARCH models are a way to capture dynamic volatility in time series. The general idea—updating variance at a point toward recent error between fit and data—has many variants. [Engle \[2004\]](#) is a survey, and [Hansen and Lunde \[2005\]](#) a comparison. Multivariate versions include [Bollerslev \[1990\]](#) and the DCC model of [Engle \[2002\]](#); [Bauwens et al. \[2006\]](#) is a survey. [Engle et al. \[2019\]](#) improves DCC by adding the random matrix theory shrinkage of [Ledoit and Wolf \[2012\]](#). [Engle et al. \[1990\]](#) and section 6.2 of [Bollerslev et al. \[1994\]](#)

¹Section 1.2 of [Bollerslev et al. \[1994\]](#) describes some of the structure present in stock returns.

apply ARCH to factor-modeled covariance. Though often used, like many econometric time series techniques, ARCH type models miss relevant information.

Range-based estimators infer volatility from a shorter, i.e., more recent, history. [Parkinson \[1980\]](#) models $\log(\text{price})$ as a Brownian motion and, from the distribution of its high-low², works out the ratio between powers of volatility and expected powers of $\log(\text{intraday high price/low})$ yielding an estimator with the accuracy of the unbiased time series estimator in $1/5$ the observations. The idea can be applied at higher frequency for shorter lag and more localization. Other range estimators, adding the information of price at market open and close, include [Garman and Klass \[1980\]](#), [Beckers \[1983\]](#), [Rogers and Satchell \[1991\]](#), [Kunitomo \[1992\]](#), and [Yang and Zhang \[2000\]](#). [Chou et al. \[2010\]](#) is a good overview. Correlation between securities can also be inferred, as in [Rogers and Zhou \[2008\]](#) and [Chou et al. \[2009\]](#). ARCH style models that update volatility with range information are [Chou \[2005\]](#), [Brandt and Jones \[2006\]](#), [Hansen et al. \[2012\]](#), and [De Nard et al. \[2021\]](#) which is [Engle et al. \[2019\]](#) with range information added.

Implied volatility, from market-traded options and indices, e.g., VIX, is available instantaneously and contains the forward-looking, high quality information of market participants. Cross-sectional volatility, the dispersion of returns across securities, reveals correlation and can be computed over any interval. There are also continually updated predictive signals, for example, stock message board postings in [Antweiler and Frank \[2004\]](#) and [Das and Chen \[2007\]](#), news in [Mitra et al. \[2009\]](#) and [Atkins et al. \[2018\]](#), and internet search queries in [Dimpfl and Jank \[2016\]](#) in the ever-growing list of alternative data. Two examples of ARCH style models that update volatility with news signals are [Kalev et al. \[2004\]](#) and [Robertson et al. \[2007\]](#).

The fragmented information from these current information sources relates to particular aspects of covariance, e.g., the variance of a certain stock or the spread across constituents of an index. Forming a forecast covariance matrix requires grafting onto structure. One way is to gather structure from history, then tune levels to the information. Made public in 1998, [diBartolomeo and Warrick \[2005\]](#) updates an orthogonal-factor model of covariance,

$$\Sigma = E \text{diag}(\Lambda) E^T + \text{diag}(\varepsilon^2), \quad (4.1)$$

²From [Feller \[1951\]](#).

with option-implied-volatility-forecasted security variances $\hat{v}$ by applying scalars θ to the factor variances,

$$\Sigma_\theta = E \operatorname{diag}(\Lambda) \operatorname{diag}(\theta) E^T + \operatorname{diag}(\varepsilon^2), \quad (4.2)$$

and setting their value through constrained L2 regression,

$$\theta^* = \arg \min_{\theta \geq 0} \sum_i \left(\sum_k E_{ik}^2 \Lambda_k \theta_k + \varepsilon_i^2 - \hat{v}_i \right)^2. \quad (4.3)$$

[Mitra et al. \[2009\]](#) does the same adding signals from quantified news. [Shah \[2007\]](#) formalizes the idea into a Bayesian framework for fitting any factor model to forecasts of functions of covariance and derives a method to scale frequency to correct mismatch in underlying structure, e.g., monthly to daily returns. [Shah \[2014\]](#), the basis of this chapter, is a generalization that has a computationally fast parameterization for certain types of forecasts. [Miranyan \[2012\]](#) introduces scalars into a diagonalized factor model and sets their value by minimizing Frobenius distance to a target full covariance matrix having correlation from history and standard deviations from individual stock volatility forecasts.

4.2 Problem formulation

The goal is a forecast of return covariance over some frequency and horizon given history and side information. Two requirements are 1-modeling the present since covariance changes with market conditions and 2-fast computation since the horizon of interest can be short.

Within the mathematical model of chapter 3 (stock returns Y , side information X , latent states Z), assume the covariance takes a known parametric form,

$$\operatorname{cov}[Y_t|S_t] = f(S_t) \text{ where } S_t \text{ is a subset of } Z_t$$

One example is the factor model of 3.2.1,

$$\text{cov}[Y_t|S_t] = \beta \Omega \beta^T + \Lambda \quad (4.4)$$

$$S_t \equiv \{\beta, \Omega, \Lambda\}$$

$$Z_t \equiv S_t \cup \{F_t\}$$

$$\text{where } Y_t \equiv \beta F_t + \varepsilon_t$$

$$F_t \sim \text{iid } \mathcal{N}[0, \Omega]$$

$$\varepsilon_t \sim \text{iid } \mathcal{N}[0, \Lambda]$$

Side information contains noisy estimates of functions of covariance built from contemporaneous and recent data,

$$\{G_t, g_t\}_t \text{ where } G_t = g_t(\text{cov}[Y_t|S_t]) + \eta_t \quad (4.5)$$

Two practical examples of side information:

1. Current estimates of return variance for individual securities, $\text{cov}[Y_t]_{i,i}$, or portfolios, $w^T \text{cov}[Y_t] w$. One source is range-based volatility estimation: assuming a model for security movement, e.g., geometric Brownian motion, variance is inferred from high and low prices over an interval. Another source is implied volatility from continuously traded options. A third is forecasting via any regression, including machine learning³, fed recent data.
2. Cross-sectional variance of the returns of a portfolio's constituents over an interval, $\sum_i w_i (r_i - w^T r)^2$. Its expected value $\approx \text{interval length} \times w^T \text{diag}(\text{cov}[Y_t]) - w^T \text{cov}[Y_t] w$. Paired with individual security variance estimates, these capture aspects of correlation.

The covariance forecast is calibrated to current market conditions by fitting parameters to side information and history as in equation (3.4) (repeated here),

$$P(Z_T|X_{1:T-1}, Y_{1:T-1}) = \frac{1}{P(X_{1:T-1}, Y_{1:T-1})} \int \left[\prod_{t=1}^{T-1} P(X_t|Z_t) P(Y_t|Z_t) \right] P(Z_{1:T}) dZ_1 \dots dZ_{T-1}$$

Through both $P(Z_{1:T})$ and changes in side information, slow-moving parameters are determined more by history, and fast-moving parameters by current side information.

³e.g., Zhang et al. [2023]

4.3 A particular parameterization

The fast computation parameterization of the paper starts from a factor model as in equation (4.4), diagonalizes the covariance between factors, and inserts time-varying variances. The factor covariance matrix Ω is replaced by $u \text{diag}[\Theta_t] u^T$, and the idiosyncratic variance Λ by $\gamma_t \Lambda$.

$$\text{cov}[Y_t|S_t] = \beta u \text{diag}[\Theta_t] u^T \beta^T + \gamma_t \Lambda \quad (4.6)$$

$$S_t \equiv \{\beta, u, \Lambda, \Theta_t, \gamma_t\}$$

$$Z_t \equiv S_t \cup \{F_t\}$$

$$\text{where } Y_t \equiv \beta F_t + \varepsilon_t$$

$$F_t \equiv u \text{diag}[\Theta_t]^{\frac{1}{2}} \times \text{iid } \mathcal{N}[0, I_k]$$

$$\varepsilon_t \equiv \sqrt{\gamma_t} \times \text{iid } \mathcal{N}[0, \Lambda]$$

$$\Theta_t, \gamma_t \equiv \text{dynamic variances} \in \mathbb{R}^k, \mathbb{R}$$

$$u, \Lambda \equiv \text{structure} \in \mathbb{R}^{k \times k}, \mathbb{R}^{n \times n}, \Lambda \text{ is diagonal}$$

$$\text{So, } Y_t|Z_t \sim \mathcal{N}(\beta F_t, \gamma_t \Lambda) \quad (4.7)$$

$$Y_t|S_t \sim \mathcal{N}(0, \beta u \text{diag}[\Theta_t] u^T \beta^T + \gamma_t \Lambda) \quad (4.8)$$

Add side information:

$$X_t \equiv \{F_t, G_t, g_t(\cdot)\}$$

$$\text{where } G_t = g_t(\text{cov}[Y_t|S_t]) + \eta_t \in \mathbb{R}^{m_t} \text{ as in equation (4.5)}$$

$$= g_t(\beta u \text{diag}[\Theta_t] u^T \beta^T + \gamma_t \Lambda) + \eta_t$$

$$= g_t(Z_t) + \eta_t, \quad \eta_t \sim P_{\eta,t} \text{ ind of everything else}$$

$$X_t|Z_t \sim \delta_{F_t} P_{\eta,t}(G_t - g_t(Z_t)) \quad (4.9)$$

and a distribution on Z :

$$Z_t \equiv \{\beta, u, \Lambda, \Theta_t, \gamma_t, F_t\}$$

$$P(Z_1, \dots, Z_T) = \delta_{\beta, u, \Lambda} \times \pi(\beta, u, \Lambda, \Theta_0, \gamma_0) \times (\text{dynamics on } \Theta, \gamma) \times \prod_{t=1}^T P(F_t|u, \Theta_t) \quad (4.10)$$

For dynamics, Θ_t, γ_t might be Markov $\sim \phi(\Theta_{t-1}, \gamma_{t-1})$ where ϕ is some non-negative distribution, e.g., log-normal with known parameters. Then,

$$P(Z_{1:T}) = \delta_{\beta, u, \Lambda} \times \pi(\beta, u, \Lambda, \Theta_0, \gamma_0) \times \prod_{t=1}^T \phi(\Theta_{t-1}, \gamma_{t-1}) \times \prod_{t=1}^T P(F_t | u, \Theta_t)$$

Conditioning on $Z_T \setminus F_T$ as in equation (3.6),

$$\begin{aligned} \text{cov}[Y_T | X_{1:T-1}, Y_{1:T-1}] &= \text{E} \left[\underbrace{\text{cov}[Y_T | X_{1:T-1}, Y_{1:T-1}, Z_T \setminus F_T]}_{\beta u \text{diag}[\Theta_T] u^T \beta^T + \gamma_T \Lambda} | X_{1:T-1}, Y_{1:T-1} \right] \\ &\quad + \text{cov} \left[\underbrace{\text{E}[Y_T | X_{1:T-1}, Y_{1:T-1}, Z_T \setminus F_T]}_{\beta \text{E}[F_T | Z_T \setminus F_T] = 0} | X_{1:T-1}, Y_{1:T-1} \right] \\ &= \text{E} [\beta u \text{diag}[\Theta_T] u^T \beta^T + \gamma_T \Lambda | X_{1:T-1}, Y_{1:T-1}] \end{aligned}$$

The distribution of parameters $P(Z_T | X_{1:T-1}, Y_{1:T-1})$ can be computed via equation (3.4) from $P(Y_t | Z_t)$, $P(X_t | Z_t)$, and $P(Z_{1:T})$ (equations 4.7, 4.9, and 4.10). Sampling numerically yields an approximate expectation and likely configurations. However, practical considerations intrude.

4.4 Practical considerations

1. Data. Some numbers are missing, wrong, or peculiar and unindicative of the future⁴. Additionally, the universe of securities changes — stocks disappear, merge, or suddenly appear, e.g., after an initial public offering or with the creation of the euro in 1999. Thus, the matrix of relevant historic returns $Y_{1:T}$ is incomplete and dirty. In practice, steps are taken to trim outliers and proxy from similar securities. Admittedly, this could be addressed through an informative prior on (β, Λ) with missing entries of $Y_{1:T}$ unobserved.
2. Complicated econometric modeling. For example, β is often inferred using outside data, e.g., from SEC financial filings. While making β, Λ time-dependent and rolling the outside data into a prior for each period could accommodate within the formal framework, other issues such as markets trading over different time intervals in Asia and the US are trickier. Because estimation specifics

can be both involved and continually revised, there is an advantage to modularity.

3. Reporting. Firms across the investment industry use a common set of models of established structure to communicate, e.g., for comparison, auditing, and showing portfolio risks without revealing positions. A model built from scratch is less valuable than a standard adapted to current conditions. For example, with the latter, one could see how much risk different well-known factors contribute over several horizons, e.g., 2 days, 2 weeks, and 1 month.
4. Computational speed. Models need to be recomputed quickly under short horizons or regime change, e.g., at the inception of the coronavirus pandemic. Separately, it is likely that, even with full computation, material change occurs in few areas.

For these reasons, the approach taken is to forecast covariance given current information by altering an existing estimate, Σ_t , using only the most recent observation of current information, X_t , rather than the full history, $\{X_{1:t}, Y_{1:t}\}$. Assuming the effects of history and dynamics to be encapsulated in the prior simplifies the problem to a single period. Though a distribution is available, only the MAP estimate of S_t is used in a point-estimate of covariance.

⁴A difference between a market and a physical system

from [Shah \[2014\]](#) **Short-Term Risk and Adapting Covariance to Current Market Conditions**

4.5 Abstract

Covariance models of stock returns appear throughout the investment process, e.g., forecasting portfolio risk, hedging, constructing mean-variance optimal portfolios, and algorithmic trading. Typically built from historic time-series, they estimate the past but—because markets and regimes continually change—not the present and future. There are non-time-series, instantaneous ways to infer or predict volatility, e.g., option implied volatility, intra-period trading range, machine learning on alternative data. This paper describes a conceptual framework and algorithm for incorporating these inferences into any linear factor covariance model so that it represents one’s beliefs about future behavior instead of echoing the past.

4.6 Key Messages

1. Factor covariance models are built from historic series, necessary to infer structure.
2. Thus, their information is stale and ignores much of what’s known and relevant.
3. Expressing (often non-time series) current information as quantities relating to covariance, one can inject it into a factor model to forecast the current and future.

4.7 Introduction

Covariance models of stock returns appear throughout the investment process, e.g., forecasting portfolio risk, hedging, constructing mean-variance optimal portfolios, and algorithmic trading. They are built from historical data, with more complex models

requiring more data thus reaching farther back in time. Richness in structure is impossible without data.

But the world continually changes. What happened years ago or recently in a different regime poorly represents the present and future. Moreover, the data entering these models is just a fraction of what’s relevant. Many immediate information sources—ignored due to econometric legacy—can be used to infer or predict volatility, e.g., the cross-sectional spread of returns, option implied volatility, intra-period high/low prices, volume, bond spreads, short interest, news/social media and other alternative data. What follows is a framework and algorithm for incorporating this information.

4.7.1 Background

The literature on several related threads is rich and well-developed. [Engle \[1982\]](#) introduced ARCH models to forecast volatility of heteroscedastic time-series from history. Countless variants have since been investigated. [Parkinson \[1980\]](#) describes one of many estimators to infer volatility of stock returns from intra-period high/low prices or other intra-period statistics. [Antweiler and Frank \[2004\]](#) forecasted volatility from stock message board postings. Today, alternative data sources are usually the best way to observe or predict essentially every quantity of interest in the world.

On this paper’s particular thread, [diBartolomeo and Warrick \[2005\]](#) adjusted an orthogonal factor covariance model to match forward-looking option-implied volatility. [Mitra et al. \[2009\]](#) applied the technique to incorporate an aspect of quantified news. [Shah \[2007\]](#) presented a Bayesian framework to adjust models whose factors are correlated while considering information accuracy. [Miranyan \[2012\]](#) used stock level volatility information to adjust a diagonalized factor model – first creating a target by inflating or deflating rows and columns of modeled covariance to match individual stock volatilities, then scaling the model’s factor variances to yield covariance as close (measured by Frobenius distance) to the target as possible.

This paper contributes 1—a general framework to incorporate, considering accuracy, forecasts of any quantity that can be calculated by a covariance matrix; 2—a computationally fast algorithm to incorporate stock, portfolio, and cross-sectional volatility forecasts with consideration of accuracy; 3—a model and formula to scale

covariance to longer/shorter horizons under serial correlation; and 4—a diagonalization of factor covariance that maximizes explained variance of securities linked to factors by exposures.

4.7.2 General Structure of Factor Covariance Models

Because there are too many securities to accurately infer all the pairwise entries, models are built on a skeleton of common factors, e.g., market, industry, country, and growth/value spread. Covariance between securities happens through links to these factors. Movement uncorrelated with the factors is treated as uncorrelated across securities and called stock-specific.

The factor model view of covariance:

$$\begin{aligned}\mathbf{C} &= (n \times n) \text{ modeled covariance of stock returns} \\ &= \mathbf{E}\mathbf{F}\mathbf{E}^T + \text{diag}(\boldsymbol{\varepsilon}^2)\end{aligned}\tag{4.11}$$

where $n = \#$ of securities

$m = \#$ of factors

$\mathbf{E} = (n \times m)$ security exposures to factors

$\mathbf{F} = (m \times m)$ covariance between factors

$\boldsymbol{\varepsilon}^2 = (n \times 1)$ uncorrelated stock specific variance

The individual pieces are typically inferred from time-series reaching back 3-5 years. Some might come from stock prices and accounting or econometric data updated perhaps quarterly.

4.8 Current Information

Current information here is communicated through *forecasts of quantities related to variance or covariance* over one's horizon. Forecasts can be made exclusively from instantly observable data, e.g., VIX, or involve history, e.g., the volume of news today and every day over the past year. They can be the result of black box machine learning on alternative data.

A quantity forecasted might be the variance of TSLA over the next 3 months or the standard deviation of the S&P500 over the next day. Quantities can be relative—the variance of TSLA over the next 3 months divided by its variance over the past 5 years.

Insufficient to construct the entire covariance, *current information is projected onto the structure of a base covariance model*. More information yields projections that are more descriptive and, over time, dynamic. Less, and the projection resembles the base model. Communicating information in quantities calculated by a covariance matrix makes the process possible. Examples:

4.8.1 Portfolio or Individual Security Variance

$$\begin{aligned}\sigma_v^2 &= \text{variance of portfolio } \mathbf{v} \\ &= \mathbf{v}^T \mathbf{C} \mathbf{v}\end{aligned}\tag{4.12}$$

where $\mathbf{v} = (n \times 1)$ portfolio weights

Useful cases:

1. variance stripped of market (or other) effects⁵
 $\mathbf{v} = \text{portfolio weights} - \beta \times \text{market weights}$
2. tracking error
 $\mathbf{v} = \text{portfolio weights} - \text{benchmark weights}$
3. The difference in variance between two portfolios, e.g. a stock's variance above the market

⁵Imagine different sources are relevant for predicting market and non-market behavior, e.g., VIX for market, company news for non-market. In a forecast of combined market and non-market behavior, the (very large) market effect would dominate the error. Without separating the two, information in the non-market part would get lost.

4.8.2 Expected Cross-Sectional Volatility

The cross-sectional expected squared difference from the average return, weighted by portfolio value

$$\begin{aligned}\sigma_{v,cs}^2 &= \text{expected cross-sectional variance}^6 \text{ of portfolio } \mathbf{v} \\ &\approx \mathbf{v}^T \boldsymbol{\omega}^2 - \mathbf{v}^T \mathbf{C} \mathbf{v}\end{aligned}\tag{4.13}$$

$$\begin{aligned}\text{where } \boldsymbol{\omega}^2 &= (n \times 1) \text{ variance of stock returns} \\ &= \text{diag}(\mathbf{C})\end{aligned}\tag{4.14}$$

4.8.3 Pairwise Correlation between Securities or Portfolios

$$\begin{aligned}\rho_{u,v} &= \text{correlation between portfolios } \mathbf{u} \text{ and } \mathbf{v} \\ &= \frac{\mathbf{u}^T \mathbf{C} \mathbf{v}}{\sqrt{(\mathbf{u}^T \mathbf{C} \mathbf{u})(\mathbf{v}^T \mathbf{C} \mathbf{v})}}\end{aligned}\tag{4.15}$$

4.8.4 Average Correlation within a Portfolio

$$\begin{aligned}\bar{\rho}_v &= \text{average correlation}^7 \text{ within portfolio } \mathbf{v} \\ &= \frac{\mathbf{v}^T \mathbf{C} \mathbf{v} - (\mathbf{v}^2)^T \boldsymbol{\omega}^2}{(\mathbf{v}^T \boldsymbol{\omega})^2 - (\mathbf{v}^2)^T \boldsymbol{\omega}^2}\end{aligned}\tag{4.16}$$

4.9 Bayesian Framework for Incorporating Current Information

Presented in [Shah \[2007\]](#), the framework is as follows: Into the model (1), introduce parameters $\boldsymbol{\theta}$ which will serve to characterize the state of the market. State here

⁶Derivation in [4.14](#)

⁷Derivation in [4.15](#)

resembles regime but can be richer and more granular. Let $\pi(\boldsymbol{\theta})$ be $\boldsymbol{\theta}$'s prior distribution. Current information forecasts, $\hat{\mathbf{g}}$, are noisy estimates of quantities, $\mathbf{g}$, calculated by the model: $\mathbf{g}$ = functions of the covariance matrix under $\boldsymbol{\theta}$.

Written in terms of $\boldsymbol{\theta}$, $\mathbf{g} = \mathbf{f}(\boldsymbol{\theta})$.

Estimates $\hat{\mathbf{g}} = \mathbf{g} + \boldsymbol{\delta}$.

Errors, $\boldsymbol{\delta}$, follow some distribution, $q(\boldsymbol{\delta})$.

These give the distribution of $\boldsymbol{\theta}$ given $\hat{\mathbf{g}}$:

$$\begin{aligned}
 p(\boldsymbol{\theta}|\hat{\mathbf{g}}) &= p(\hat{\mathbf{g}}|\boldsymbol{\theta}) \pi(\boldsymbol{\theta}) / p(\hat{\mathbf{g}}) \\
 &= q(\hat{\mathbf{g}} - \mathbf{g}) \pi(\boldsymbol{\theta}) / p(\hat{\mathbf{g}}) \\
 &= q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}]) \pi(\boldsymbol{\theta}) / p(\hat{\mathbf{g}}) \\
 &\propto q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}]) \pi(\boldsymbol{\theta})
 \end{aligned} \tag{4.17}$$

Now the question can be posed, “How does the world appear considering the forecasts?” The expected value is $E[\boldsymbol{\theta}]$ taken under (4.17). The most probable configuration—the Bayesian analog of frequentist statistics’ maximum likelihood—is the maximum a posteriori (MAP) estimate:

$$\begin{aligned}
 \boldsymbol{\theta}^* &= \arg \max_{\boldsymbol{\theta}} p(\boldsymbol{\theta}|\hat{\mathbf{g}}) \\
 &= \arg \max_{\boldsymbol{\theta}} \log [p(\boldsymbol{\theta}|\hat{\mathbf{g}})] \\
 &= \arg \max_{\boldsymbol{\theta}} \log [q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}])] + \log [\pi(\boldsymbol{\theta})]
 \end{aligned} \tag{4.18}$$

The first term fits toward observations; the second stabilizes and seeks coherence with prior beliefs. Three questions: 1-How are the parameters introduced? 2-What is the observation error distribution? 3-What is a reasonable prior? Consider an example:

4.9.1 Parameters to Accommodate State

Recall the base model (4.11), $\mathbf{C} = \mathbf{E}\mathbf{F}\mathbf{E}^T + \text{diag}(\boldsymbol{\varepsilon}^2)$. Among the many possibilities, for this example, place a scalar on the standard deviation of each factor and on stock-specific volatility.

$$\begin{aligned}\mathbf{C}(\boldsymbol{\theta}) &= (n \times n) \text{ covariance} \\ &= \mathbf{E} [\text{diag}(\boldsymbol{\theta}_{1..m}) \mathbf{F} \text{diag}(\boldsymbol{\theta}_{1..m})] \mathbf{E}^T + \boldsymbol{\theta}_{m+1}^2 \text{diag}(\boldsymbol{\epsilon}^2)\end{aligned}\quad (4.19)$$

where $\boldsymbol{\theta} = (m + 1 \times 1)$

$\boldsymbol{\theta}_{1..m}$ = scales factor std deviations

$\boldsymbol{\theta}_{m+1}$ = scales security specific std deviation

as before, $\boldsymbol{\omega}^2(\boldsymbol{\theta}) = (n \times 1)$ variance of stock returns

$$= \text{diag}[\mathbf{C}(\boldsymbol{\theta})] \quad (4.20)$$

Now the model can adapt to state.

4.9.2 Forecasts and Errors

Ideally, forecasts well outnumber parameters. Redundancy connects optimized parameters to reality. However, even absent data, the prior makes the problem well-posed.

Suppose the information consists of the volatility of two indices, the relative cross-sectional volatility of one, and the volatility of certain individual stocks.

- $\hat{g}_1$ = std dev of S&P500 from VIX implied vol index
- $\hat{g}_2$ = std dev of R2000 inferred from intraday high/low
- $\hat{g}_3$ = relative level (3 day avg / 3 year avg) of daily
cross-sectional std dev of R2000 returns
- $\hat{g}_4$ = std dev of MSFT from option implied vol
- $\vdots$

These correspond in the model to:

$$\begin{aligned}
f_1(\boldsymbol{\theta}) &= \sqrt{\mathbf{v}_{SP500}^T \mathbf{C}(\boldsymbol{\theta}) \mathbf{v}_{SP500}} \\
f_2(\boldsymbol{\theta}) &= \sqrt{(\mathbf{v}_{R2000}^T \mathbf{C}(\boldsymbol{\theta}) \mathbf{v}_{R2000})} \\
f_3(\boldsymbol{\theta}) &= \frac{\sqrt{(\mathbf{v}_{R2000}^T \boldsymbol{\omega}^2(\boldsymbol{\theta}) - \mathbf{v}_{R2000}^T \mathbf{C}(\boldsymbol{\theta}) \mathbf{v}_{R2000})}}{\sqrt{(\mathbf{v}_{R2000}^T \boldsymbol{\omega}^2 - \mathbf{v}_{R2000}^T \mathbf{C} \mathbf{v}_{R2000})}} \\
f_4(\boldsymbol{\theta}) &= \sqrt{\boldsymbol{\omega}_{MSFT}^2(\boldsymbol{\theta})} \\
&\vdots
\end{aligned}$$

Errors can be estimated since quantities are observed repeatedly over time or across a peer group. Depending on the method of inference, econometrically derived estimates also exist.

Suppose in this example they are distributed independent Laplace, which has the advantage of making the parameter selection robust to outliers.

$$q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}]) = c \prod_j e^{-\sqrt{2}|\hat{g}_j - f_j(\boldsymbol{\theta})|/\sigma_j} \quad (4.21)$$

$$\log[q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}])] = c_2 - \sqrt{2} \sum_j |\hat{g}_j - f_j(\boldsymbol{\theta})|/\sigma_j \quad (4.22)$$

Note that when information sources have correlated errors—e.g., errors in S&P 500 and R2000 volatility forecasts result largely from changes in market volatility—ignoring the dependency gives undue influence to not-independent observations. The correlation can be explicitly considered in the error distribution or ad hoc by inflating error estimates.

4.9.3 Prior

A model's factors are generally observable, and their statistics paired with understanding about market behavior—e.g., how certain factor volatilities naturally correlate—can be the basis of informative priors. For the example consider a simple prior, uncorrelated Gaussians around a common center that is a Gaussian centered at 1.

$$\pi(\boldsymbol{\theta}) = c \underbrace{\prod_k e^{-\frac{(\theta_k - \theta_0)^2}{2\lambda_k^2}}}_{\text{parameters}} \underbrace{e^{-\frac{(\theta_0 - 1)^2}{2\lambda_0^2}}}_{\text{center}} \quad (4.23)$$

$$\log[\pi(\boldsymbol{\theta})] = c_2 - \sum_k \frac{(\theta_k - \theta_0)^2}{2\lambda_k^2} - \frac{(\theta_0 - 1)^2}{2\lambda_0^2} \quad (4.24)$$

The spread of the center, λ_0 , can come from the standard deviation over time of (the market's volatility out the horizon, say 2 weeks / its volatility as in the model) Spreads $\lambda_1 \dots \lambda_k$ can come from the same on the factors or simply take the value λ_0 .

Assembling the pieces yields the most likely configuration:

$$\arg \max_{\boldsymbol{\theta}, \theta_0} \left\{ -\sqrt{2} \sum_j |\hat{g}_j - f_j(\boldsymbol{\theta})| / \sigma_j - \sum_k \frac{(\theta_k - \theta_0)^2}{2\lambda_k^2} - \frac{(\theta_0 - 1)^2}{2\lambda_0^2} \right\} \quad (4.25)$$

Entering the maximizing values into the parametrized model of equation (4.19) yields a covariance model that represents current information.

4.10 Adapting to Current Market Conditions by Blending Ledoit-Wolf Covariance Reductions

The framework doesn't need a particular parameterization of covariance.

To soften outliers in portfolio optimization, [Ledoit and Wolf \[2004\]](#) blend the estimated covariance matrix with highly structured reductions⁸:

1. single index – covariance happens only through a market factor
2. constant correlation – all securities have the same pairwise correlation

A variant is:

3. constant covariance – all securities have the same pairwise correlation and variance

Beyond low dimensional approximations, these can be regarded as market deformations, e.g., during a panic, securities become more alike, and the market heads toward constant covariance.

⁸Derived in [4.15](#)

4.10.1 Expressed as 1-Factor Models

One factor covariance is of the form $\mathbf{e}\mathbf{e}^T\sigma^2 + \text{diag}(\boldsymbol{\varepsilon}^2)$.

where $\mathbf{e} = (n \times 1)$ exposures to factor

$\boldsymbol{\varepsilon}^2 = (n \times 1)$ uncorrelated stock specific variance

$\sigma^2 =$ variance of factor

let $\mathbf{m} = (n \times 1)$ market portfolio weights

$\sigma_m^2 =$ variance of market portfolio

$$= \mathbf{m}^T \mathbf{C} \mathbf{m}$$

$\bar{\rho} =$ average correlation

$\bar{\sigma}^2 =$ average variance

The Ledoit-Wolf forms can be expressed as 1-factor models:

1. Single Index

$$\sigma_\beta^2 = \sigma_m^2$$

$$\mathbf{e}_\beta = \mathbf{C}\mathbf{m}/\sigma_m^2$$

$$\boldsymbol{\varepsilon}_\beta^2 = \boldsymbol{\omega}^2 - \sigma_m^2 \mathbf{e}_\beta^2$$

$$\mathbf{C}_\beta = \mathbf{e}_\beta \mathbf{e}_\beta^T \sigma_m^2 + \text{diag}(\boldsymbol{\varepsilon}_\beta^2)$$

2. Constant Correlation

$$\sigma_{\bar{\rho}}^2 = \bar{\rho}$$

$$\mathbf{e}_{\bar{\rho}} = \boldsymbol{\omega}$$

$$\boldsymbol{\varepsilon}_{\bar{\rho}}^2 = \boldsymbol{\omega}^2(1 - \bar{\rho})$$

$$\mathbf{C}_{\bar{\rho}} = \boldsymbol{\omega} \boldsymbol{\omega}^T \bar{\rho} + \text{diag}(\boldsymbol{\omega}^2)(1 - \bar{\rho})$$

3. Constant Covariance

$$\begin{aligned}
\sigma_{\bar{\rho}, \bar{\sigma}}^2 &= \bar{\rho} \\
\mathbf{e}_{\bar{\rho}, \bar{\sigma}}^2 &= \mathbf{1} \bar{\sigma} \\
\boldsymbol{\varepsilon}_{\bar{\rho}, \bar{\sigma}}^2 &= \mathbf{1} \bar{\sigma}^2 (1 - \bar{\rho}) \\
\mathbf{C}_{\bar{\rho}, \bar{\sigma}} &= \mathbf{1} \mathbf{1}^T \bar{\sigma}^2 \bar{\rho} + \mathbf{I} \bar{\sigma}^2 (1 - \bar{\rho})
\end{aligned}$$

4.10.2 Parametrization

Consider a convex combination of the three reductions $\mathbf{C}_\beta$, $\mathbf{C}_{\bar{\rho}}$, $\mathbf{C}_{\bar{\rho}, \bar{\sigma}}$ and the original $\mathbf{C}$:

$$\begin{aligned}
\mathbf{C}(\boldsymbol{\theta}, \bar{\rho}, \bar{\sigma}) &= \theta_1 \mathbf{C}_\beta + \theta_2 \mathbf{C}_{\bar{\rho}} + \theta_3 \mathbf{C}_{\bar{\rho}, \bar{\sigma}} + \theta_4 \mathbf{C} \\
&= \theta_1 [\mathbf{e}_\beta \mathbf{e}_\beta^T \sigma_m^2 + \text{diag}(\boldsymbol{\varepsilon}_\beta^2)] \\
&\quad + \theta_3 [\boldsymbol{\omega} \boldsymbol{\omega}^T \bar{\rho} + \text{diag}(\boldsymbol{\omega}^2)(1 - \bar{\rho})] \\
&\quad + \theta_3 [\mathbf{1} \mathbf{1}^T \bar{\sigma}^2 \bar{\rho} + \mathbf{I} \bar{\sigma}^2 (1 - \bar{\rho})] \\
&\quad + \theta_4 [\mathbf{E} \mathbf{F} \mathbf{E}^T + \text{diag}(\boldsymbol{\varepsilon}^2)]
\end{aligned}$$

where $\boldsymbol{\theta}, \bar{\sigma} \geq 0$

$$0 \leq \bar{\rho} \leq 1$$

$$\sum_k \theta_k = 1$$

The most likely configuration given the prior and forecasts is:

$$\arg \max_{\boldsymbol{\theta}, \bar{\rho}, \bar{\sigma}} \log [q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}, \bar{\rho}, \bar{\sigma}])] + \log [\pi(\boldsymbol{\theta}, \bar{\rho}, \bar{\sigma})] \quad (4.26)$$

With just five free parameters—one θ doesn't count since it's determined by the others—sufficient number of observations exist that they likely dominate the prior. Without a prior, i.e. with a non-informative prior, the optimization is on just the error distribution:

$$\arg \max_{\boldsymbol{\theta}, \bar{\rho}, \bar{\sigma}} \log [q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}, \bar{\rho}, \bar{\sigma}])] \quad (4.27)$$

The solution yields a current information covariance model within the parametrized form.

4.11 Adapting to Current Market Conditions by Linear and Quadratic Programming

Using a particular set of forecast quantities and assuming Laplace or Gaussian prior and error distributions, an adapting parameterization can be solved by linear or quadratic programming.

4.11.1 Work in a Representation with Uncorrelated Factors

Assume the covariance model has orthogonal factors:

$$\mathbf{C} = \mathbf{E} \operatorname{diag}(\boldsymbol{\eta}^2) \mathbf{E}^T + \operatorname{diag}(\boldsymbol{\varepsilon}^2)$$

where $\boldsymbol{\eta}^2 = (m \times 1)$ variance of the factors

$$\begin{aligned} \boldsymbol{\omega}^2 &= (n \times 1) \text{ variance of stock returns} \\ &= \mathbf{E}^2 \boldsymbol{\eta}^2 + \boldsymbol{\varepsilon}^2 \end{aligned}$$

When that isn't the case, make it so by projecting the factors onto an orthogonal basis. One way is multiplying exposures by the square root of the covariance matrix:

$$\begin{aligned} \mathbf{F} &= \mathbf{Q} \boldsymbol{\Lambda} \mathbf{Q}^T \text{ eigendecomposed} & (4.28) \\ \mathbf{F}^{\frac{1}{2}} &= \mathbf{Q} \boldsymbol{\Lambda}^{\frac{1}{2}} \mathbf{Q}^T \\ \mathbf{M} &= \mathbf{F}^{\frac{1}{2}} \\ \mathbf{M}^T \mathbf{M} &= \mathbf{F} \\ \mathbf{E} \mathbf{F} \mathbf{E}^T &= (\mathbf{E} \mathbf{M}^T)(\mathbf{E} \mathbf{M}^T)^T \\ \mathbf{E}_{new} &\leftarrow \mathbf{E} \mathbf{M}^T \\ \mathbf{F}_{new} &\leftarrow \mathbf{I} \end{aligned}$$

Another way is described in [section 4.12](#).

4.11.2 Linear Properties

Forecasting properties that can be written as linear combinations of the parameters under optimization yields an exploitable mathematical structure. Instead of requiring

a general purpose nonlinear optimization algorithm, with appropriate prior and error distributions, the optimization is solved by linear or quadratic programming.

Individual security, portfolio, and expected cross-sectional variance are linear functions of the model variances, $\boldsymbol{\eta}^2$ and $\boldsymbol{\varepsilon}^2$.

1. Portfolio Variance

$$\begin{aligned}\sigma_v^2 &= \mathbf{v}^T \mathbf{C} \mathbf{v} \\ &= \underbrace{(\mathbf{v}^T \mathbf{E})^2}_{\mathbf{a}^T} \boldsymbol{\eta}^2 + \underbrace{(\mathbf{v}^2)^T}_{\mathbf{b}^T} \boldsymbol{\varepsilon}^2 \\ &= \mathbf{a}^T \boldsymbol{\eta}^2 + \mathbf{b}^T \boldsymbol{\varepsilon}^2\end{aligned}$$

2. Portfolio-weighted Expected Cross-Sectional Variance

$$\begin{aligned}\sigma_{v,cs}^2 &= \mathbf{v}^T \boldsymbol{\omega}^2 - \mathbf{v}^T \mathbf{C} \mathbf{v} \\ &= \mathbf{v}^T [\mathbf{E}^2 \boldsymbol{\eta}^2 + \boldsymbol{\varepsilon}^2] - [(\mathbf{v}^T \mathbf{E})^2 \boldsymbol{\eta}^2 + (\mathbf{v}^2)^T \boldsymbol{\varepsilon}^2] \\ &= \underbrace{[\mathbf{v}^T \mathbf{E}^2 - (\mathbf{v}^T \mathbf{E})^2]}_{\mathbf{a}^T} \boldsymbol{\eta}^2 + \underbrace{[\mathbf{v} - \mathbf{v}^2]^T}_{\mathbf{b}^T} \boldsymbol{\varepsilon}^2 \\ &= \mathbf{a}^T \boldsymbol{\eta}^2 + \mathbf{b}^T \boldsymbol{\varepsilon}^2\end{aligned}$$

On these types of forecasts, the connection between forecasts and model predictions is $\mathbf{g} = \mathbf{A}\boldsymbol{\eta}^2 + \mathbf{B}\boldsymbol{\varepsilon}^2$.

4.11.3 Parameterization

Let $\boldsymbol{\theta} \geq 0$ scale factor variances and $\gamma \geq 0$ stock-specific variances.

$$\begin{aligned}\boldsymbol{\eta}^2(\boldsymbol{\theta}) &= \text{diag}(\boldsymbol{\eta}^2) \boldsymbol{\theta} \\ \boldsymbol{\varepsilon}^2(\gamma) &= \gamma \boldsymbol{\varepsilon}^2\end{aligned}$$

Then $\mathbf{g} = \mathbf{f}(\boldsymbol{\theta}, \gamma)$

$$\begin{aligned}&= \underbrace{\mathbf{A} \text{diag}(\boldsymbol{\eta}^2)}_{\tilde{\mathbf{A}}} \boldsymbol{\theta} + \underbrace{\mathbf{B}\boldsymbol{\varepsilon}^2}_{\tilde{\mathbf{b}}} \gamma \\ &= \tilde{\mathbf{A}}\boldsymbol{\theta} + \tilde{\mathbf{b}}\gamma \text{ which is linear in } \boldsymbol{\theta} \text{ and } \gamma\end{aligned}$$

Note: if the model begin with correlated factors, adjustments would transfer back via

$$\mathbf{F}(\boldsymbol{\theta}) = \mathbf{M}^T \text{diag}[\boldsymbol{\eta}^2(\boldsymbol{\theta})] \mathbf{M} \quad (4.29)$$

4.11.4 Observation Noise and Prior

Suppose both the noise and $\boldsymbol{\theta}$'s prior are distributed independent Laplace⁹, with $\boldsymbol{\theta}$ centered at θ_0 . θ_0 and γ have non-informative priors.

Noise:

$$q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}, \gamma]) = c \prod_j e^{-\sqrt{2}|\hat{g}_j - f_j(\boldsymbol{\theta}, \gamma)|/\sigma_j} \quad (4.30)$$

$$\begin{aligned} \log[q(\hat{\mathbf{g}} - \mathbf{f}[\boldsymbol{\theta}, \gamma])] &= c_2 - \sqrt{2} \sum_j |\hat{g}_j - f_j(\boldsymbol{\theta}, \gamma)|/\sigma_j \\ &= c_2 - \left(\sqrt{2}/\boldsymbol{\sigma}\right)^T |\hat{\mathbf{g}} - \mathbf{f}(\boldsymbol{\theta}, \gamma)|^{10} \\ &= c_2 - \left(\sqrt{2}/\boldsymbol{\sigma}\right)^T |\hat{\mathbf{g}} - \tilde{\mathbf{A}}\boldsymbol{\theta} - \tilde{\mathbf{b}}\gamma| \end{aligned} \quad (4.31)$$

Prior:

$$\pi(\boldsymbol{\theta}) = c \prod_k e^{-\sqrt{2}|\theta_k - \theta_0|/\lambda_k} \quad (4.32)$$

$$\log[\pi(\boldsymbol{\theta})] = c_2 - \left(\sqrt{2}/\boldsymbol{\lambda}\right)^T |\hat{\boldsymbol{\theta}} - \mathbf{1}\theta_0| \quad (4.33)$$

4.11.5 MAP Estimate of Parameters

Under these distributions and forecasts, the most likely configuration is

$$\arg \max_{\boldsymbol{\theta}, \theta_0, \gamma} \log[q(\cdot)] + \log[\pi(\cdot)] \quad (4.34)$$

$$= \arg \min_{\boldsymbol{\theta}, \theta_0, \gamma} (1/\boldsymbol{\sigma})^T |\hat{\mathbf{g}} - \tilde{\mathbf{A}}\boldsymbol{\theta} - \tilde{\mathbf{b}}\gamma| + (1/\boldsymbol{\lambda})^T |\hat{\boldsymbol{\theta}} - \mathbf{1}\theta_0|$$

$$\text{where } \boldsymbol{\theta}, \theta_0, \gamma \geq 0 \quad (4.35)$$

⁹but restricted so variance numbers are positive: $\boldsymbol{\theta} \geq 0$, and a variance forecast's error $\leq$ the forecast. These will be imposed as constraints and not properly considered in the distributions.

¹⁰ $|\cdot|$ here and throughout means absolute value taken component by component.

which is solved by linear programming¹¹.

Following the same steps with Gaussian errors and priors leads to an optimization problem solved by quadratic programming.

4.11.6 Restrictions for Observable Factors

Factors, before orthogonalization, are often observable and have sources of information about the future—e.g., oil price and the option implied volatility of oil—that should be represented in the fitted model.

$$\begin{aligned}
 \mathbf{F}(\boldsymbol{\theta})_{i,i} &= \text{adjusted variance of factor } i \\
 &= (\mathbf{M}^T \text{diag} [\boldsymbol{\eta}^2(\boldsymbol{\theta})] \mathbf{M})_{i,i} \\
 &= \sum_k M_{i,k}^2 \eta^2(\boldsymbol{\theta})_k \\
 &= \sum_k \underbrace{M_{i,k}^2 \eta_k^2}_{X_{i,k}} \theta_k \\
 &= \mathbf{X}_i \boldsymbol{\theta}
 \end{aligned} \tag{4.36}$$

Mapping to a linear combination of parameters, a factor's prediction can be incorporated like any forecast or as an optimization constraint.

4.11.7 Casting Prior onto Orthogonal Factors

Prior beliefs are best made on factors pre-orthogonalization. Beyond the first few, factors from eigendecomposition are hard to map to identifiable phenomena and change over time.

To avoid using reusing symbols, let $\boldsymbol{\psi}$ = the variance multipliers on the original factors, and $\boldsymbol{\theta}$ the same on the orthogonal factors.

$$\begin{aligned}
 \psi_i &= \text{scaled/base variance of original factor } i \\
 &= \frac{\mathbf{X}_i \boldsymbol{\theta}}{\mathbf{X}_i \mathbf{1}} \\
 &= \frac{\mathbf{X}_i}{\underbrace{\mathbf{X}_i \mathbf{1}}_u} \boldsymbol{\theta}
 \end{aligned}$$

¹¹Apply standard trick: $\min |x|$ becomes $\min \delta$ where $\delta \geq x, \delta \geq -x$

The map is linear, $\psi(\theta) = U\theta$. Change of variables gives the prior on the orthogonal factors.

$$\begin{aligned} \pi(\psi) &= \text{prior on original variance scalars } \psi \\ \text{prior on } \theta &= \pi(U\theta) |U| \\ &\propto \pi(U\theta) \end{aligned} \tag{4.37}$$

4.11.8 Reducing Parameters to Prevent Overfitting

More data relative to the number of parameters provides stability, makes the optimizing configuration less sensitive to individual data points.

Suppose $m' < m$ of the m factor variances were individually scaled and the remaining $(m - m')$ scaled as one. What factors should be among the m' ? Ideally, the set yields the highest posterior likelihood under its maximizer. However, finding that set is a combinatorial problem where each function evaluation requires solving a linear program.

Consider a heuristic. If forecasts are on portfolio or cross-sectional variance (not on the difference in portfolio variances), entries in $\tilde{\mathbf{A}}$ and $\tilde{\mathbf{b}}$ are non-negative, and rows sum to the forecasted quantities under the base model. The i 'th column of $\tilde{\mathbf{A}}$'s holds the i 'th factor's contributions. $\tilde{\mathbf{b}}$ holds the contributions from stock-specific variance. The heuristic is to sum the columns of $\tilde{\mathbf{A}}$ and take the highest m' . To place more weight on accurate forecasts, weight rows by the inverse of the forecast's error standard deviation.

In the opposite direction, for more granularity, stock-specific variance scalars can vary by group, for example, by a company's economic sector.

4.12 Orthogonal Directions that Best Explain Securities

Orthogonalizing via eigendecomposing factor covariance, as in equation (4.28), doesn't consider connection to securities. The most important first orthogonal factor could be a direction of high factor variance but zero explanatory power.

What is truly sought is the entirely different goal of directions that best explain securities (as specified by a set of securities and weights).

This can be tailored to one's particular universe. For example, investment style or sector focus might concentrate a portfolio's risk along certain directions. Viewing under a tailored basis could better reveal risks and improve the incorporation of current information.

The orthogonal decomposition that best explains the $\mathbf{w}$ -weighted variance of securities is $\mathbf{F} = \mathbf{N}\mathbf{N}^T$ where

$\mathbf{w} = (n \times 1)$ weight of securities in the fit

$\mathbf{E}_{i,\cdot} = (1 \times m)$ i 'th row of $\mathbf{E}$

$\mathbf{H} = (m \times m)$ $\mathbf{w}$ -weighted second moments of $\mathbf{E}$

$$= \sum_i w_i \mathbf{E}_{i,\cdot}^T \mathbf{E}_{i,\cdot}$$

$\mathbf{Q}, \mathbf{\Lambda} = (m \times m)$ from eigendecomposition

$$\mathbf{Q}\mathbf{\Lambda}\mathbf{Q}^T = \mathbf{F}$$

$\tilde{\mathbf{Q}} = (m \times m)$ from eigendecomposition

$$\tilde{\mathbf{Q}}\tilde{\mathbf{\Lambda}}\tilde{\mathbf{Q}}^T = \mathbf{\Lambda}^{\frac{1}{2}}\mathbf{Q}^T\mathbf{H}\mathbf{Q}\mathbf{\Lambda}^{\frac{1}{2}}$$

$$\mathbf{N} = (m \times m) \mathbf{Q}\mathbf{\Lambda}^{\frac{1}{2}}\tilde{\mathbf{Q}}$$

Proof.

Let $\mathbf{d} = (k \times 1)$ a combination (direction) of factors

$\boldsymbol{\beta} = (n \times 1)$ stock exposures to $\mathbf{d}$

$$= \text{cov}[\text{stock returns}, \mathbf{d} \text{ returns}] / \text{var}[\mathbf{d} \text{ returns}]$$

$$= \frac{\mathbf{E}\mathbf{F}\mathbf{d}}{\mathbf{d}^T\mathbf{F}\mathbf{d}}$$

The variance of stock i explained by $\mathbf{d}$ is $\beta_i^2 \text{var}[\mathbf{d} \text{ returns}] = \frac{(\mathbf{E}_{i,\cdot}\mathbf{F}\mathbf{d})^2}{\mathbf{d}^T\mathbf{F}\mathbf{d}}$. The goal is to

find $\mathbf{d}$ that maximizes the weighted explained variance.

$$\begin{aligned}
\text{wt explained var} &= \sum_i w_i \frac{(\mathbf{E}_{i,\cdot} \mathbf{F} \mathbf{d})^2}{\mathbf{d}^T \mathbf{F} \mathbf{d}} \\
&= \sum_i w_i \frac{(\mathbf{d}^T \mathbf{F} \mathbf{E}_{i,\cdot}^T) (\mathbf{E}_{i,\cdot} \mathbf{F} \mathbf{d})}{\mathbf{d}^T \mathbf{F} \mathbf{d}} \\
&= \frac{\mathbf{d}^T \mathbf{F} (\sum_i w_i \mathbf{E}_{i,\cdot}^T \mathbf{E}_{i,\cdot}) \mathbf{F} \mathbf{d}}{\mathbf{d}^T \mathbf{F} \mathbf{d}} \\
&= \frac{\mathbf{d}^T \mathbf{F} \mathbf{H} \mathbf{F} \mathbf{d}}{\mathbf{d}^T \mathbf{F} \mathbf{d}}
\end{aligned} \tag{4.38}$$

$$\text{Let } \mathbf{v} = \Lambda^{\frac{1}{2}} \mathbf{Q}^T \mathbf{d} \tag{4.39}$$

$$\text{Then } \mathbf{d} = \mathbf{Q} \Lambda^{-\frac{1}{2}} \mathbf{v} \tag{4.40}$$

$$\begin{aligned}
\text{and } \mathbf{v}_1^T \mathbf{v}_2 &= \left(\mathbf{d}_1^T \mathbf{Q} \Lambda^{\frac{1}{2}} \right) \left(\Lambda^{\frac{1}{2}} \mathbf{Q}^T \mathbf{d}_2 \right) \\
&= \mathbf{d}_1^T \mathbf{Q} \Lambda \mathbf{Q}^T \mathbf{d}_2 \\
&= \mathbf{d}_1^T \mathbf{F} \mathbf{d}_2
\end{aligned} \tag{4.41}$$

$$= \text{cov} [\mathbf{d}_1, \mathbf{d}_2] \tag{4.42}$$

$$\begin{aligned}
\text{wt explained var} &= \frac{\mathbf{d}^T \mathbf{F} \mathbf{H} \mathbf{F} \mathbf{d}}{\mathbf{d}^T \mathbf{F} \mathbf{d}} \\
&= \frac{(\mathbf{v}^T \Lambda^{-\frac{1}{2}} \mathbf{Q}^T) \mathbf{F} \mathbf{H} \mathbf{F} (\mathbf{Q} \Lambda^{-\frac{1}{2}} \mathbf{v})}{(\mathbf{v}^T \Lambda^{-\frac{1}{2}} \mathbf{Q}^T) \mathbf{F} (\mathbf{Q} \Lambda^{-\frac{1}{2}} \mathbf{v})} \\
&= \frac{\mathbf{v}^T \left(\Lambda^{\frac{1}{2}} \mathbf{Q}^T \mathbf{H} \mathbf{Q} \Lambda^{\frac{1}{2}} \right) \mathbf{v}}{\mathbf{v}^T \mathbf{v}}
\end{aligned} \tag{4.43}$$

Eigendecomposing the center matrix, $\Lambda^{\frac{1}{2}} \mathbf{Q}^T \mathbf{H} \mathbf{Q} \Lambda^{\frac{1}{2}}$, into $\tilde{\mathbf{Q}} \tilde{\Lambda} \tilde{\mathbf{Q}}^T$ gives $\mathbf{v}$ directions that maximize weighted explained variance. They are the columns of $\tilde{\mathbf{Q}}$.

Written in terms of $\mathbf{d}$, via (4.40), they are the columns of $\mathbf{D} \equiv \mathbf{Q} \Lambda^{-\frac{1}{2}} \tilde{\mathbf{Q}}$.

Because vectors from the eigendecomposition are orthonormal and $\text{cov} [\mathbf{d}_1, \mathbf{d}_2] = \mathbf{v}_1^T \mathbf{v}_2$, $\text{cov} [\mathbf{D}] = \mathbf{I}$.

$$\begin{aligned}
\text{Let } \tilde{\mathbf{E}} &= (n \times m) \text{ exposures to factors in the new basis} \\
&= \mathbf{E}\mathbf{F}\mathbf{D}
\end{aligned} \tag{4.44}$$

$$\begin{aligned}
\mathbf{R} &= (m \times m) \text{ exposure of existing factors to it} \\
&= \mathbf{I}\mathbf{F}\mathbf{D} \\
&= \mathbf{F}\mathbf{D}
\end{aligned} \tag{4.45}$$

The covariance decomposition is

$$\begin{aligned}
\mathbf{F} &= \text{cov} [\mathbf{R}] \\
&= (\mathbf{F}\mathbf{D}) \text{cov} [\mathbf{D}] (\mathbf{F}\mathbf{D})^T \\
&= \left(\mathbf{Q}\mathbf{\Lambda}^{\frac{1}{2}}\tilde{\mathbf{Q}} \right) \mathbf{I} \left(\mathbf{Q}\mathbf{\Lambda}^{\frac{1}{2}}\tilde{\mathbf{Q}} \right)^T
\end{aligned} \tag{4.46}$$

□

Note: One could also find the solution via the first m entries of the eigendecomposition of the $n \times n$ matrix $\mathbf{E}\mathbf{F}\mathbf{E}^T$ (typically $n \gg m$) with computational complexity $\mathcal{O}(n^3)$ instead of $\mathcal{O}(m^3 + nm^2)$.

4.13 Changing the Time Scale

Return frequency of the base and adapted model need not agree. The base model might look at monthly returns with a horizon of 6 months or a year while current information typically applies to short horizons (2 days, 1 week) and return frequencies (1 day, 1 week). Adapting addresses the horizon, but time scale differences remain.

Though both sets of numbers can be reported in annualized terms¹², negative serial correlation in daily returns inflates variance over short intervals relative to long. Not accounting for the serial correlation underpredicts risk. The variance multiplier to correct this is approximately

$$\frac{1 - \rho^2 - 2\rho(1 - \rho^{T_1})/T_1}{1 - \rho^2 - 2\rho(1 - \rho^{T_2})/T_2} (1 + \mu)^{2(T_1 - T_2)} \tag{4.47}$$

¹²annualized variance = periods in yr $\times$ 1 period variance

where T_1 = short horizon, e.g., 5 days

T_2 = long horizon, e.g., 60 days

ρ = daily serial correlation in returns

μ = mean daily return

Proof.

r_t = return on day t

$r_{0 \rightarrow T}$ = cumulative return from time 0 to time T

$$= f(r_1, \dots, r_T) = \prod_{t=1}^T (1 + r_t) - 1$$

Linearized around its mean,

$$\begin{aligned} &\approx f|_{r_1, \dots, r_T = \mu} + \nabla f|_{r_1, \dots, r_T = \mu} (\mathbf{r} - \boldsymbol{\mu}) \\ &= c + (1 + \mu)^{T-1} \sum_{t=1}^T r_t - \mu \end{aligned}$$

Taking the variance of both sides,

$$\text{var}[r_{0 \rightarrow T}] \approx (1 + \mu)^{2T-2} \sum_{1 \leq j, k \leq T} \text{cov}[r_j, r_k]$$

Suppose 1 day returns follow an AR(1)

$$\begin{aligned} r_{t+1} &= c + \rho r_t + \varepsilon_{t+1} \\ \text{var}[r_t] &= \sigma^2 \end{aligned}$$

Then

$$\begin{aligned}
\text{cov}[r_j, r_k] &= \sigma^2 \rho^{|j-k|} \\
\text{var}[r_{0 \rightarrow T}] &\approx (1 + \mu)^{2T-2} \sum_{1 \leq j, k \leq T} \text{cov}[r_j, r_k] \\
&= (1 + \mu)^{2T-2} \sum_{1 \leq j, k \leq T} \sigma^2 \rho^{|j-k|} \\
&= (1 + \mu)^{2T-2} \sigma^2 T \left[\frac{1 - \rho^2 - 2\rho(1 - \rho^T)/T}{(1 - \rho)^2} \right]
\end{aligned}$$

$\frac{\text{var}[r_{0 \rightarrow T_1}]/T_1}{\text{var}[r_{0 \rightarrow T_2}]/T_2}$ = the ratio of annualized variances

$$\approx \frac{1 - \rho^2 - 2\rho(1 - \rho^{T_1})/T_1}{1 - \rho^2 - 2\rho(1 - \rho^{T_2})/T_2} (1 + \mu)^{2(T_1 - T_2)} \quad (4.48)$$

□

4.14 Expected Cross-Sectional Volatility

Portfolio weighted cross-sectional expected squared difference from portfolio return equals portfolio weighted average variance minus the variance of the portfolio plus (negligible) squared difference in means.

$$\begin{aligned}
\sigma_{v,cs}^2 &= \text{expected cross-sectional variance of portfolio } \mathbf{v} \\
&\approx \mathbf{v}^T \boldsymbol{\omega}^2 - \mathbf{v}^T \mathbf{C} \mathbf{v} \quad (4.49)
\end{aligned}$$

Proof.

$$\begin{aligned}
\mathbf{r} &= (n \times 1) \text{ security returns} \\
\boldsymbol{\mu} &= \mathbb{E}[\mathbf{r}] \\
\sum_i v_i &= 1 \\
\sigma_{v,cs}^2 &= \mathbb{E} \left[\sum_i v_i (r_i - \mathbf{v}^T \mathbf{r})^2 \right] \\
&= \mathbb{E} \left[\sum_i v_i ([r_i - \mu_i] - \mathbf{v}^T [\mathbf{r} - \boldsymbol{\mu}] + [\mu_i - \mathbf{v}^T \boldsymbol{\mu}])^2 \right] \\
&= \sum_i v_i \mathbb{E} [(r_i - \mu_i)^2] + \sum_i v_i \mathbb{E} [(\mathbf{v}^T [\mathbf{r} - \boldsymbol{\mu}])^2] \\
&\quad + \sum_i v_i (\mu_i - \mathbf{v}^T \boldsymbol{\mu})^2 \\
&\quad - 2 \mathbb{E} \left[\mathbf{v}^T (\mathbf{r} - \boldsymbol{\mu}) \underbrace{\sum_i v_i (r_i - \mu_i)}_{\mathbf{v}^T (\mathbf{r} - \boldsymbol{\mu})} \right] \\
&\quad + 2 \sum_i v_i \underbrace{\mathbb{E} [r_i - \mu_i]}_0 (\mu_i - \mathbf{v}^T \boldsymbol{\mu}) \\
&\quad - 2 \sum_i v_i \underbrace{\mathbb{E} [\mathbf{v}^T (\mathbf{r} - \boldsymbol{\mu})]}_0 (\mu_i - \mathbf{v}^T \boldsymbol{\mu}) \\
&= \underbrace{\sum_i v_i \mathbb{E} [(r_i - \mu_i)^2]}_{\mathbf{v}^T \boldsymbol{\omega}^2} - \underbrace{\mathbb{E} [(\mathbf{v}^T [\mathbf{r} - \boldsymbol{\mu}])^2]}_{\mathbf{v}^T \mathbf{C} \mathbf{v}} \\
&\quad + \underbrace{\sum_i v_i (\mu_i - \mathbf{v}^T \boldsymbol{\mu})^2}_{\text{negligible}^{13}} \\
&\approx \mathbf{v}^T \boldsymbol{\omega}^2 - \mathbf{v}^T \mathbf{C} \mathbf{v}
\end{aligned}$$

¹³Both variance and mean return scale with time. Because the mean return term is squared, it vanishes over small time. For example, say the annual variance is $(25\%)^2$ and the annual return is 10%. Over one trading day, the variance is $(25\%)^2/250 = 2.5\%^2$, and the squared mean return is $(10\%/250)^2 = 0.0016\%^2$.

□

4.15 Average Correlation and Variance within a Portfolio

$$\begin{aligned}
 \bar{\rho}_v &= \text{average correlation within portfolio } \mathbf{v} \\
 &= \frac{\mathbf{v}^T \mathbf{C} \mathbf{v} - (\mathbf{v}^2)^T \boldsymbol{\omega}^2}{(\mathbf{v}^T \boldsymbol{\omega})^2 - (\mathbf{v}^2)^T \boldsymbol{\omega}^2}
 \end{aligned} \tag{4.50}$$

$$\begin{aligned}
 \bar{\sigma}_v^2 &= \text{average variance under average correlation} \\
 &= \frac{\mathbf{v}^T \mathbf{C} \mathbf{v}}{\bar{\rho}_v - (1 - \bar{\rho}_v) \mathbf{v}^T \mathbf{v}}
 \end{aligned} \tag{4.51}$$

Proof.

$\bar{\rho}$:

$$\begin{aligned}
\mathbf{v}^T \mathbf{C} \mathbf{v} &= \text{var}_v \\
&= \sum_i v_i^2 \text{var}_i + \sum_{j \neq i} v_i v_j \text{cov}_{i,j} \\
&= \sum_i v_i^2 \omega_i^2 + \sum_{j \neq i} v_i v_j \omega_i \omega_j \rho_{i,j} \\
&= \sum_i v_i^2 \omega_i^2 + \sum_{j \neq i} v_i v_j \omega_i \omega_j \bar{\rho} \\
&= \sum_i v_i^2 \omega_i^2 + \sum_{j \neq i} v_i v_j \omega_i \omega_j \bar{\rho} \\
&\quad + \sum_{j=i} v_i v_j \omega_i \omega_j \bar{\rho} - \sum_{j=i} v_i v_j \omega_i \omega_j \bar{\rho} \\
&= \sum_i v_i^2 \omega_i^2 + \bar{\rho} \left(\sum_i v_i \omega_i \right) \left(\sum_j v_j \omega_j \right) \\
&\quad - \bar{\rho} \sum_i v_i^2 \omega_i^2 \\
&= (\mathbf{v}^2)^T \boldsymbol{\omega}^2 + \bar{\rho} (\mathbf{v}^T \boldsymbol{\omega})^2 - \bar{\rho} (\mathbf{v}^2)^T \boldsymbol{\omega}^2
\end{aligned}$$

$$\begin{aligned}
\bar{\rho} &= \frac{\mathbf{v}^T \mathbf{C} \mathbf{v} - (\mathbf{v}^2)^T \boldsymbol{\omega}^2}{(\mathbf{v}^T \boldsymbol{\omega})^2 - (\mathbf{v}^2)^T \boldsymbol{\omega}^2} \\
&= \frac{\text{var} - \text{uncorrelated var}}{\text{perfectly correlated var} - \text{uncorrelated var}}
\end{aligned}$$

$\bar{\sigma}$:

$$\begin{aligned}
 \mathbf{v}^T \mathbf{C} \mathbf{v} &= (\mathbf{v}^2)^T \boldsymbol{\omega}^2 + \bar{\rho} (\mathbf{v}^T \boldsymbol{\omega})^2 - \bar{\rho} (\mathbf{v}^2)^T \boldsymbol{\omega}^2 \\
 &= (\mathbf{v}^2)^T [\bar{\sigma} \mathbf{1}]^2 + \bar{\rho} (\mathbf{v}^T [\bar{\sigma} \mathbf{1}])^2 - \bar{\rho} (\mathbf{v}^2)^T [\bar{\sigma} \mathbf{1}]^2 \\
 &= \bar{\sigma}^2 \left[\underbrace{(\mathbf{v}^2)^T \mathbf{1}}_{\mathbf{v}^T \mathbf{v}} + \bar{\rho} \underbrace{(\mathbf{v}^T \mathbf{1})^2}_1 - \bar{\rho} (\mathbf{v}^2)^T \mathbf{1} \right] \\
 &= \bar{\sigma}^2 [\bar{\rho} + (1 - \bar{\rho}) \mathbf{v}^T \mathbf{v}]
 \end{aligned}$$

$$\bar{\sigma}^2 = \frac{\mathbf{v}^T \mathbf{C} \mathbf{v}}{\bar{\rho} + (1 - \bar{\rho}) \mathbf{v}^T \mathbf{v}}$$

□

CHAPTER 5

UNCERTAIN COVARIANCE AND PORTFOLIO OPTIMIZATION UNDER UNCERTAINTY

The work in this chapter grew out of the problem of estimation error in portfolio optimization. It can also be applied elsewhere in quantitative finance, e.g., chapters 6 and 7, and beyond finance.

Because optimized selection of frequentist estimates finds outliers in observation noise, especially as dimension grows, forecasts need to represent a predictive (according to one's beliefs) distribution where each number with its associated uncertainty is believed and articulated.

Standard Bayesian techniques to infer covariance of time series are unwieldy, perhaps ill-equipped to deal with stock returns' high dimension and non-stationarity, and hard to visualize. Meanwhile, modeling covariance through factors,

$$\Sigma_t \equiv \text{cov}[Y_t | \beta, \Omega, \Lambda] = \beta \Omega \beta^T + \Lambda,$$

is ubiquitous in the investment world, and versions with parameters occupying a range rather than point-estimated appear in the robust optimization literature¹.

To users of factor models, the covariance matrix is of interest, but so are the values of each parameter; indeed, a slew of diagnostic reports are generated from them. For example, the portfolio averaged exposures, $w^T \beta$, reveal the bets, intentional and unintentional, taken by the manager. Incorporating uncertainty while retaining the factor model structure yields information not available through alternative Bayesian methods in addition to reducing computation.

Parameters change from point-estimated to having a mean and covariance, making

¹e.g., Goldfarb and Iyengar [2003].

Σ_t random. Formulae are derived for $E[\Sigma_t]$ and $\text{cov}[\Sigma_t]$ under assumptions on independence and distributional shape. Note that, because of convexity, expected portfolio variance exceeds the point-estimate assembled from parameter means², $w^T E[\Sigma_t] w \geq w^T \hat{\Sigma}_t w$.

The second part of the research uses the result to improve portfolio construction. Given Markowitz utility as in equation (2.6),

$$U_\lambda(w) = w^T \mu - \frac{\lambda}{2} w^T \Sigma w$$

$$\text{So, } \text{var}[U(w)] = \text{var}[w^T \mu] + \frac{\lambda^2}{4} \text{var}[w^T \Sigma w] - \lambda \text{cov}[w^T \mu, w^T \Sigma w] \quad (5.1)$$

$$\approx \text{var}[w^T \mu] + \frac{\lambda^2}{4} \text{var}[w^T \Sigma w] \quad (5.2)$$

The first term comes from the forecaster, the second from the uncertain covariance model. Maximizing expected utility with a penalty on uncertainty,

$$w^* \equiv \arg \max_w E[U(w)] - \kappa \times \text{stdev}[U(w)], \quad (5.3)$$

is effectively robust optimization with the uncertainty set defined on utility and is a way to mitigate the estimation error problems encountered in portfolio optimization.

5.1 Literature and mathematical model

Section 2.2.1 ‘Frequentist problems’ in the literature review of chapter 2 outlines the motivation for this work. The alternative construction methods covered in section 2.2.3 describe creating a predictive distribution of covariance (under ‘Bayesian posterior’) and adding uncertainty to the parameters of a factor model (under ‘Robust optimization’). This work bridges the two.

For efficiency in the repetitive computations of optimization, the paper that is the basis of this chapter, [Shah \[2015\]](#), uses a diagonalized factor model similar to that of equation (4.6),

$$\text{cov}[Y_t | \beta, \Theta, \gamma, \Lambda] \equiv \beta \text{diag}[\Theta] \beta^T + \gamma \Lambda, \quad (5.4)$$

² $\hat{\Sigma}_t = E[\beta] E[\Omega] E[\beta]^T + E[\Lambda]$

with different variable names,

$$\begin{aligned}\Sigma_t &\equiv \text{cov}[Y_t|Z_t] \\ &\equiv E_t \text{diag}(\theta_{t,1:K}) E_t^T + \theta_{t,K+1} \text{diag}(\varepsilon_t^2)\end{aligned}\tag{5.5}$$

$$\text{where } Z_t \equiv \{E_t, \theta_t, \varepsilon_t^2\}\tag{5.6}$$

In the mathematical framework of chapter 3, we are interested in the mean and variance of $\Sigma_t|X_t$ where the side information $X_t = \{\text{mean and variance of } Z_t\}$ comes from a forecaster. X_t encapsulates all the information of $Y_{1:t-1}, X_{1:t-1}, Z_{1:t-1}$, i.e., forecasts can't be improved. The act of forecasting is a separate earlier procedure whose methods are left to the forecaster and free to evolve.

Assumptions on distributional shape (Gaussian E) and independence (E, θ, ε^2 jointly independent, the elements of ε^2 independent) add the information to make computation possible.

from [Shah \[2015\]](#) Uncertain Covariance Models and Uncertainty-Penalized Portfolio Optimization

5.2 Abstract

Covariance appears throughout investment management, e.g., risk reporting, portfolio optimization, risk parity, smart beta, algorithmic trading, and hedging. It is often represented via multi-factor model. The imposed structure—comovement through sensitivity to common factors, diagonal residuals—offers advantages such as less data for estimation, easy proxying, and removing transient behavior and low eigenvalue directions. Nevertheless, parameter values are inferred and imperfect. Though it's common to ignore the error and proceed from point estimates, the spread of possible values is considered in Bayesian techniques and robust optimization. The research in this paper connects to both: Forecasts of mean and variance of the factor model parameters propagate to forecast the mean and variance of the covariance matrix. Using the result, optimization is performed to maximize portfolio utility a given number of standard deviations below the mean.

5.3 Introduction

Covariance between security returns, whose behavior is non-stationary and dimension high relative to the length of history, is often represented via factor model. Connections between securities happen through *exposure* to common factors, e.g., market,

industry, principal components. Movement uncorrelated with factors, called *idiosyncratic*, is assumed uncorrelated across securities. The standard form:

$$\begin{aligned}\Sigma &= (N \times N) \text{ factor-modeled covariance of stock returns} \\ &= \mathbf{E}\mathbf{F}\mathbf{E}^T + \text{diag}[\epsilon^2]\end{aligned}\tag{5.7}$$

where $N = \#$ of securities

$K = \#$ of factors

$\mathbf{E} = (N \times K)$ exposure to factors

$\mathbf{F} = (K \times K)$ covariance between factors

$\epsilon^2 = (N \times 1)$ idiosyncratic variance

Parameter values, often treated as known³, are estimates. Examples where they aren't point-estimated but occupy a range exist in the robust optimization literature, e.g., Goldfarb and Iyengar [2003] and Lu [2011]. In this paper, rather than falling in a range, parameters are regarded as following a distribution whose mean and covariance are given by the forecaster. With assumptions on independence and distributional shape, the information implies the mean and covariance of Σ .

The second part of paper enlists the result to compute the variance of portfolio utility and maximize utility a number of standard deviations below the mean, akin to robust optimization with the uncertainty set a function of portfolio weights and defined on portfolio utility.

5.3.1 Background

Within the large literature on portfolio optimization written since Markowitz [1952] introduced the notion of an efficient frontier, this paper connects to work on estimation error and modeling covariance uncertainty.

Topics around estimation error fall into 3 categories: 1–frequentist problems, 2–improving point estimates by Bayesian and other techniques, 3–optimizing beyond point estimates. For the first, early especially insightful papers are Jobson and Korkie [1980] and Broadie [1993]. For the second, on converting frequentist to Bayes estimates are Jorion [1986], Frost and Savarino [1986], and Black and Litterman

³e.g., commercial risk models from MSCI Barra, Axioma, and Northfield

[1992]. Robust statistics, as in Huber [1981], is used to infer point estimated inputs in Cavadini et al. [2001]. In its essence Bayesian, Ledoit and Wolf [2003] collapses a full covariance matrix toward a single factor model projection. Forms of shrinkage also appear in Fan et al. [2013], Abadir et al. [2014], Lam [2016], Ledoit and Wolf [2017], Goldberg et al. [2018], De Nard et al. [2019], and Engle et al. [2019]. Michaud [1998], which is critiqued by Scherer [2002], introduces ‘resampling’—repeatedly generating a fixed length sample of a normal distribution under the estimated parameters and optimizing using parameters inferred from the sample, then averaging the output. Particulars aside, it is akin to optimizing with a softened covariance matrix. For the third category, Goldfarb and Iyengar [2003] explores portfolio construction by robust optimization (no relation to robust statistics), associating a range with each point and finding the point with the best worst case. Its sophisticated max of min optimization restricts application to particular structures of utility functions and constraints. A more recent survey is Bertsimas et al. [2011]. Lu [2011] looks at joint uncertainty in return and variance. DeMiguel and Nogales [2009] and Yang et al. [2015] maximize portfolio utility as measured by robust statistics.

On covariance uncertainty, standard Bayesian time-series techniques, though unwieldy and requiring simulation, can treat covariance beyond point estimates. Models also appear in the robust optimization literature. Halldórsson and Tütüncü [2003] defines the range as all positive semi-definite matrices that fall within bounds on each entry. Goldfarb and Iyengar [2003] operates under the factor structure of equation (5.7) and bounds the separate pieces.

5.4 Probability model

As a foundation for the randomness implied by forecasts, imagine embedding equation (5.7) in a probability space $(\Omega, \mathcal{F}, P)$ and that covariance conditionally follows a factor

model.

$$\boldsymbol{\Sigma} = \text{cov} [\mathbf{r} | \mathbf{E}, \mathbf{F}, \boldsymbol{\varepsilon}^2] \quad (5.8)$$

$$= \mathbf{E} \mathbf{F} \mathbf{E}^T + \text{diag} [\boldsymbol{\varepsilon}^2] \quad (5.9)$$

where $\mathbf{r}$ = security returns

$\mathbf{r}_F$ = factor returns

$\mathbf{r}_\varepsilon$ = idiosyncratic returns

$$\mathbf{r} = \mathbf{E} \mathbf{r}_F + \mathbf{r}_\varepsilon$$

Note that the unconditional covariance is expected factor-modeled covariance plus (relatively small) variance in mean.

$$\text{cov} [\mathbf{r}] = \text{E} [\text{cov} [\mathbf{r} | \mathbf{E}, \mathbf{F}, \boldsymbol{\varepsilon}^2]] + \text{cov} [\text{E} [\mathbf{r} | \mathbf{E}, \mathbf{F}, \boldsymbol{\varepsilon}^2]] \quad (5.10)$$

$$= \text{E} [\boldsymbol{\Sigma}] + \text{cov} [\text{E} [\mathbf{E} \mathbf{r}_F + \mathbf{r}_\varepsilon | \mathbf{E}, \mathbf{F}, \boldsymbol{\varepsilon}^2]]$$

Returning to equation (5.9), writing the dependence on ω highlights that these are random variables:

$$\text{for } \omega \in \Omega, \quad \boldsymbol{\Sigma}(\omega) = \mathbf{E}(\omega) \mathbf{F}(\omega) \mathbf{E}^T(\omega) + \text{diag} [\boldsymbol{\varepsilon}^2(\omega)]. \quad (5.11)$$

The probability measure P represents the forecaster's beliefs. For example, a user of point-estimates effectively puts all the mass on

$$\{w : \mathbf{E}(\omega) = \hat{\mathbf{E}}, \mathbf{F}(\omega) = \hat{\mathbf{F}}, \boldsymbol{\varepsilon}^2(\omega) = \hat{\boldsymbol{\varepsilon}}^2\},$$

in which case,

$$\text{E}_P [\boldsymbol{\Sigma}] = \hat{\mathbf{E}} \hat{\mathbf{F}} \hat{\mathbf{E}}^T + \text{diag} [\hat{\boldsymbol{\varepsilon}}^2], \text{ and } \text{cov}_P [\boldsymbol{\Sigma}] = \mathbf{0}.$$

When parameter values aren't fixed, P isn't a point mass. The first half of this paper is interested in the implications on $\boldsymbol{\Sigma}$, in particular, $\text{E}_P [\boldsymbol{\Sigma}]$ and $\text{cov}_P [\boldsymbol{\Sigma}]$.

Though notation will be dropped for brevity, this is the conceptual structure for $\boldsymbol{\Sigma}$, and expectation and variance are henceforth with respect to P .

5.5 Uncertain covariance model

Let covariance follow the factor model of equation (5.7) with a parameter γ added to connect the level of idiosyncratic variance to $\mathbf{F}$.

$$\Sigma = \text{cov} [\mathbf{r} | \mathbf{E}, \mathbf{F}, \boldsymbol{\varepsilon}^2, \gamma] \quad (5.12)$$

$$= \mathbf{E} \mathbf{F} \mathbf{F}^T + \gamma \text{diag} [\boldsymbol{\varepsilon}^2] \quad (5.13)$$

The goal is to compute $E[\Sigma]$ and $\text{cov}[\Sigma]$ from forecasts of means and covariances of $\mathbf{E}$, $\boldsymbol{\varepsilon}^2$ and $(\mathbf{F}, \gamma)$. This is possible through three assumptions:

1. $\mathbf{E}$, $\boldsymbol{\varepsilon}^2$, and $(\mathbf{F}, \gamma)$ are jointly independent.
2. $\mathbf{E}$ is Gaussian.
3. The elements of $\boldsymbol{\varepsilon}^2$ are independent.

The first, perhaps contentious, is consistent with the separation of uncertainty sets in the robust optimization literature⁴. The second need hold only at the portfolio level for portfolio variance computations. The third could be relaxed if correlations between idiosyncratic variances were provided.

Formulas for $E[\Sigma]$ and $\text{cov}[\Sigma]$ are derived in section 5.10. From equation (5.56),

$$E[\Sigma] = \hat{\mathbf{E}} \hat{\mathbf{F}} \hat{\mathbf{E}}^T + \hat{\gamma} \text{diag} [\hat{\boldsymbol{\varepsilon}}^2] + \sum_{k,l} \hat{F}_{k,l} \underbrace{\text{cov} [E_{\cdot,k}, E_{\cdot,l}]}_{N \times N}, \quad (5.14)$$

which equals the typical point-estimated covariance plus a term. The term raises the eigenvalues by adding $K(K+1)/2$ symmetric matrices whose sum is positive semi-definite. For instance, with $\mathbf{E}$ uncorrelated across securities, each of the matrices is diagonal.

5.6 Portfolio variance with orthogonal factors

Portfolio variance is the covariance quantity of interest in portfolio optimization, and less computation is incurred under the special case of orthogonal factors with random variances. Formulas are derived below for use in optimization.

⁴See Goldfarb and Iyengar [2003].

The structured form divides the factor covariance matrix $\mathbf{F}$ of equation (5.13) into $\mathbf{M} \text{diag} [\theta_1, \dots, \theta_K] \mathbf{M}^T$ where $\mathbf{M}$ is fixed and $\mathbb{E}[\mathbf{F}] = \mathbf{M} \mathbf{M}^T$. $\mathbf{M}$ expresses the original factors as linear combinations of orthonormal factors and limits the deformations of $\mathbf{F}$. Possible choices for $\mathbf{M}$ are infinite, e.g., $\mathbb{E}[\mathbf{F}]^{\frac{1}{2}} \times \text{any rotation}$ ⁵.

$$\Sigma = \mathbf{E} \left(\overbrace{\mathbf{M} \text{diag} [\theta_1, \dots, \theta_K] \mathbf{M}^T}^{\mathbf{F}} \right) \mathbf{E}^T + \theta_{K+1} \text{diag} [\epsilon^2] \quad (5.15)$$

where $\mathbf{M} = (K \times K)$

$$\boldsymbol{\theta} = (K + 1 \times 1)$$

Estimates of the mean and variance of the parameters are supplied by the forecaster.

$$\begin{array}{ll} \hat{\mathbf{E}} = (N \times K) & \hat{\mathbf{E}} = \mathbb{E}[\mathbf{E}] \\ \hat{\mathbf{X}} = (N \times K \times N \times K) & \hat{X}_{m,i,n,j} = \text{cov}[E_{m,i}, E_{n,j}] \\ \hat{\epsilon}^2 = (N \times 1) & \hat{\epsilon}^2 = \mathbb{E}[\epsilon^2] \\ \hat{\omega} = (N \times 1) & \hat{\omega}_n = \text{var}[\epsilon_n^2] \\ \hat{\boldsymbol{\theta}} = (K + 1 \times 1) & \hat{\boldsymbol{\theta}} = \mathbb{E}[\boldsymbol{\theta}] \\ \hat{\boldsymbol{\Omega}} = (K + 1 \times K + 1) & \hat{\Omega}_{i,j} = \text{cov}[\theta_i, \theta_j] \end{array}$$

As a pre-processing step, $\mathbf{M}$ is subsumed into $\mathbf{E}$,

$$\mathbf{E} \leftarrow \mathbf{E} \mathbf{M} \quad (5.16)$$

$$\hat{\mathbf{E}} \leftarrow \hat{\mathbf{E}} \mathbf{M} \quad (5.17)$$

$$\hat{\mathbf{X}}_{n,\cdot,m,\cdot} \leftarrow \mathbf{M}^T \hat{\mathbf{X}}_{n,\cdot,m,\cdot} \mathbf{M} \quad (5.18)$$

and equation (5.15) becomes

$$\Sigma = \mathbf{E} \text{diag} [\theta_1 \dots \theta_K] \mathbf{E}^T + \theta_{K+1} \text{diag} [\epsilon^2] \quad (5.19)$$

⁵If $\mathbf{Q} \mathbf{Q}^T = \mathbf{I}$, $(\mathbf{M} \mathbf{Q})(\mathbf{M} \mathbf{Q})^T = \mathbf{M} \mathbf{Q} \mathbf{Q}^T \mathbf{M}^T = \mathbf{M} \mathbf{M}^T$.

⁶ γ of equation (5.13) has been renamed θ_{K+1} to simplify notation.

Proposition 5.1. *Mean and variance of portfolio variance under orthogonal factors*
If

1. $\Sigma = \mathbf{E} \text{diag} [\theta_1 \dots \theta_K] \mathbf{E}^T + \theta_{K+1} \text{diag} [\epsilon^2]$
2. $\mathbf{E}, \epsilon^2$, and $\boldsymbol{\theta}$ are jointly independent
3. The elements of ϵ^2 are independent
4. $\mathbf{E}$ is Gaussian
5. Mean and covariance of $\mathbf{E}, \epsilon^2$, and $\boldsymbol{\theta}$ are $\hat{\mathbf{E}}, \hat{\mathbf{X}}, \hat{\epsilon}^2, \hat{\omega}, \hat{\boldsymbol{\theta}}, \hat{\Omega}$ as defined above

$\mathbf{w} = (N \times 1)$ portfolio weights or holdings

$\sigma_{\mathbf{w}}^2 = \text{variance of portfolio } \mathbf{w}$

$$\mathbb{E} [\sigma_{\mathbf{w}}^2] = \mathbf{a}^T \hat{\boldsymbol{\theta}} \quad (5.20)$$

$$\text{var} [\sigma_{\mathbf{w}}^2] = \mathbf{a}^T \hat{\Omega} \mathbf{a} + \hat{\boldsymbol{\theta}}^T \mathbf{B} \hat{\boldsymbol{\theta}} + \sum_{i,j} \hat{\Omega}_{i,j} B_{i,j} \quad (5.21)$$

$$\text{where } a_i = \begin{cases} \hat{\mu}_i^2 + \hat{\Psi}_{i,i} & 1 \leq i \leq K \\ \sum_{n=1}^N w_n^2 \hat{\epsilon}_n^2 & i = K + 1 \end{cases} \quad (5.22)$$

$$B_{i,j} = \begin{cases} 2 \hat{\Psi}_{i,j}^2 + 4 \hat{\mu}_i \hat{\mu}_j \hat{\Psi}_{i,j} & 1 \leq i, j \leq K \\ \sum_{n=1}^N w_n^4 \hat{\omega}_n & i = j = K + 1 \\ 0 & \text{otherwise} \end{cases} \quad (5.23)$$

$\mathbf{e} = \text{exposures of portfolio } \mathbf{w}$

$$\begin{aligned} \hat{\boldsymbol{\mu}}(\mathbf{w}) &= \mathbb{E} [\mathbf{e}] \\ &= \hat{\mathbf{E}}^T \mathbf{w} \end{aligned} \quad (5.24)$$

$$\begin{aligned} \hat{\Psi}(\mathbf{w}) &= \text{cov} [\mathbf{e}] \\ \hat{\Psi}_{i,j}(\mathbf{w}) &= \text{cov} [e_i, e_j] \\ &= \mathbf{w}^T \hat{\mathbf{X}}_{\cdot, i, \cdot, j} \mathbf{w} \end{aligned} \quad (5.25)$$

Proof.

$$\sigma_{\mathbf{w}}^2 = \mathbf{w}^T \boldsymbol{\Sigma} \mathbf{w} \quad (5.26)$$

$$= \mathbf{w}^T \left(\mathbf{E} \operatorname{diag} [\theta_1 \dots \theta_K] \mathbf{E}^T + \theta_{K+1} \operatorname{diag} [\boldsymbol{\varepsilon}^2] \right) \mathbf{w} \quad (5.27)$$

$$= \sum_{i=1}^K \theta_i (\mathbf{w}^T \mathbf{E})_i^2 + \theta_{K+1} \sum_{n=1}^N w_n^2 \varepsilon_n^2 \quad (5.28)$$

This can be written as the inner product

$$\sigma_{\mathbf{w}}^2 = \boldsymbol{\theta}^T \boldsymbol{\nu} \quad (5.29)$$

where $\mathbf{e}(\mathbf{w}) =$ exposures of portfolio $\mathbf{w}$

$$= \mathbf{E}^T \mathbf{w} \quad (5.30)$$

$$\text{and } \nu_i(\mathbf{w}) = \begin{cases} e_i^2 & 1 \leq i \leq K \\ \sum_{n=1}^N w_n^2 \varepsilon_n^2 & i = K + 1 \end{cases} \quad (5.31)$$

$\boldsymbol{\theta}$ and $\boldsymbol{\nu}$ are independent because $\boldsymbol{\theta}$ is independent of $\mathbf{E}, \boldsymbol{\varepsilon}$.

Mean and variance of portfolio variance are the expected value and variance of equation (5.29).

$$\mathbb{E} [\sigma_{\mathbf{w}}^2] = \mathbb{E} [\boldsymbol{\theta}^T \boldsymbol{\nu}] \quad (5.32)$$

$$= \mathbb{E} [\boldsymbol{\theta}]^T \mathbb{E} [\boldsymbol{\nu}] \quad \text{since } \boldsymbol{\theta} \text{ ind. of } \boldsymbol{\nu} \quad (5.33)$$

$$\operatorname{var} [\sigma_{\mathbf{w}}^2] = \operatorname{var} [\boldsymbol{\theta}^T \boldsymbol{\nu}] \quad (5.34)$$

$$\begin{aligned} &= \mathbb{E} [\boldsymbol{\nu}]^T \operatorname{cov} [\boldsymbol{\theta}] \mathbb{E} [\boldsymbol{\nu}] + \mathbb{E} [\boldsymbol{\theta}]^T \operatorname{cov} [\boldsymbol{\nu}] \mathbb{E} [\boldsymbol{\theta}] \\ &\quad + \sum_{i,j} (\operatorname{cov} [\boldsymbol{\theta}])_{i,j} (\operatorname{cov} [\boldsymbol{\nu}])_{i,j} \quad \text{by Lemma 5.1} \end{aligned} \quad (5.35)$$

Values are computed via equations (5.33) and (5.35) from the mean and covariance of $\boldsymbol{\theta}$ and of $\boldsymbol{\nu}$. $\boldsymbol{\theta}$ is covered by the forecasts $\hat{\boldsymbol{\theta}}$ and $\hat{\boldsymbol{\Omega}}$. What remains is to work out $\mathbb{E} [\boldsymbol{\nu}]$ and $\operatorname{cov} [\boldsymbol{\nu}]$.

Step one is computing the mean and covariance of portfolio exposures.

$$\text{Let } \hat{\boldsymbol{\mu}}(\mathbf{w}) = \mathbb{E}[\mathbf{e}] \quad (5.36)$$

$$= \hat{\mathbf{E}}^T \mathbf{w} \quad (5.37)$$

$$\text{Let } \hat{\boldsymbol{\Psi}}(\mathbf{w}) = \text{cov}[\mathbf{e}] \quad (5.38)$$

$$\begin{aligned} \hat{\Psi}_{i,j}(\mathbf{w}) &= \text{cov}[e_i, e_j] \\ &= \mathbf{w}^T \hat{\mathbf{X}}_{:,i,:j} \mathbf{w} \end{aligned} \quad (5.39)$$

The expectation of $\boldsymbol{\nu}$ defined in equation (5.31) is a function of $\hat{\boldsymbol{\mu}}$, $\hat{\boldsymbol{\Psi}}$, and the forecast $\hat{\boldsymbol{\varepsilon}}^2$:

$$\mathbb{E}[\nu_i] = \begin{cases} \mathbb{E}[e_i^2] & 1 \leq i \leq K \\ \mathbb{E}\left[\sum_{n=1}^N w_n^2 \varepsilon_n^2\right] & i = K+1 \end{cases} \quad (5.40)$$

$$= \begin{cases} \mathbb{E}^2[e_i] + \text{var}[e_i] & 1 \leq i \leq K \\ \sum_{n=1}^N w_n^2 \mathbb{E}[\varepsilon_n^2] & i = K+1 \end{cases} \quad (5.41)$$

$$= \begin{cases} \hat{\mu}_i^2 + \Psi_{i,i} & 1 \leq i \leq K \\ \sum_{n=1}^N w_n^2 \hat{\varepsilon}_n^2 & i = K+1 \end{cases} \quad (5.42)$$

Note that equation (5.41) contains an ‘uncertainty correction’ ($+\text{var}[e_i]$) arising from nonlinearity, $\mathbb{E}[y^2] = \mathbb{E}^2[y] + \text{var}[y]$. Exposure variance adds to a portfolio’s expected variance, i.e., the common practice of squaring a point estimate of exposure underestimates risk when uncertainty in exposures doesn’t diversify away, for example, in concentrated portfolios or with exposure error correlated across securities.

$\text{cov}[\nu_i, \nu_j]$ separates into three cases:

1. $i \leq K, j \leq K$:

$$\text{cov}[\nu_i, \nu_j] = \text{cov}[e_i^2, e_j^2] \quad (5.43)$$

$$= 2(\text{cov}[e_i, e_j])^2 + 4\mathbb{E}[e_i]\mathbb{E}[e_j]\text{cov}[e_i, e_j] \quad (5.44)$$

by Corollary 7.1 since $\mathbf{e}$ Gaussian

$$= 2\hat{\Psi}_{i,j}^2 + 4\hat{\mu}_i\hat{\mu}_j\hat{\Psi}_{i,j} \quad (5.45)$$

2. $i = j = K + 1$:

$$\text{cov} [\nu_i, \nu_j] = \text{var} \left[\sum_{n=1}^N w_n^2 \varepsilon_n^2 \right] \quad (5.46)$$

$$= \sum_{n=1}^N w_n^4 \text{var} [\varepsilon_n^2] \quad \text{since elements of } \boldsymbol{\varepsilon} \text{ ind.} \quad (5.47)$$

$$= \sum_{n=1}^N w_n^4 \hat{\omega}_n \quad (5.48)$$

3. $i \leq K, j = K + 1$:

$$\text{cov} [\nu_i, \nu_j] = \text{cov} \left[e_i^2, \sum_{n=1}^N w_n^2 \varepsilon_n^2 \right] = 0 \quad \text{since } \mathbf{E} \text{ ind. of } \boldsymbol{\varepsilon} \quad (5.49)$$

□

5.7 Robust optimization on portfolio utility

Vanilla Markowitz utility is

$$U_\lambda(w) = w^T \alpha - \frac{\lambda}{2} w^T \Sigma w \quad (5.50)$$

$$\text{So, } \text{var} [U(w)] = \text{var} [w^T \alpha] + \frac{\lambda^2}{4} \text{var} [w^T \Sigma w] - \lambda \text{cov} [w^T \alpha, w^T \Sigma w] \quad (5.51)$$

The first term comes from the forecaster⁷, the second from the uncertain covariance model; assuming the third negligible yields

$$\text{var} [U(w)] \approx \text{var} [w^T \alpha] + \frac{\lambda^2}{4} \text{var} [w^T \Sigma w]. \quad (5.52)$$

Maximizing expected utility with a penalty on uncertainty,

$$w^* \equiv \arg \max_w \mathbb{E} [U(w)] - \kappa \times \text{stdev} [U(w)], \quad (5.53)$$

is effectively robust optimization with the uncertainty set defined—not on separate parameters as in conventional robust optimization—but on utility. It extends κ standard deviations below the mean and is a function of portfolio weights w . Cast in

⁷ $w^T \mathbf{\Lambda} w$ where $\mathbf{\Lambda}$ represents commonality in bets.

robust optimization language:

$$w^* \equiv \arg \max_w \min_{u_w \in S_w} u_w$$

$$\text{where } S_w = \{E[U(w)] - c \times \text{stdev}[U(w)] \mid c \leq \kappa\} \quad (5.54)$$

The maximin optimization becomes a simple max since the inner minimum is known when the uncertainty set is defined on the quantity being maximized (utility).

This formulation's success mitigating estimation error problems⁸ is tested in section 5.8.

5.7.1 Indifference region to reduce transaction costs

The literature on managing transaction costs, e.g., [Michaud \[1998\]](#), mentions an indifference region to prevent nuisance trading. Rebalancing can be skipped if the improvement is judged negligible.

Such a rule can be defined from mean and standard deviation of portfolio utility. The portfolios are indifferent if their utility ranges overlap, i.e., if the upper range of the existing exceeds the lower range of the changed,

$$E[U_0] + c \times \text{stdev}[U_0] > E[U] - c \times \text{stdev}[U]. \quad (5.55)$$

5.8 An example

Data are 106 periods of weekly Wed–Wed returns from Wed Dec 20, 2017 through Tue Dec 31, 2019 for the 927 securities with complete series in (IWB) iShares Russell 1000 ETF as of Oct 1, 2020. The first 80 periods, 12/20/2017–7/3/2019, are the training set. The final 26 periods, 7/3/2019–12/31/2019, are the testing set.

An orthogonal model of section 5.6 is constructed from the training set as follows:

1. 5 principal components are inferred from return covariance with securities equally weighted. This yields $\mathbf{M}$ (an identity matrix), $\hat{\boldsymbol{\theta}}$ ⁹, and from OLS regression, a first estimate of exposures and idiosyncratic variances, $\hat{\mathbf{E}}_{ols}$ and $\hat{\boldsymbol{\varepsilon}}_{ols}^2$.

⁸[Taleb \[2012\]](#) writes that the problem with ‘Portfolio theory, mean-variance, etc.’ is ‘Assuming knowledge of the parameters, not integrating models across parameters, relying on (very unstable) correlations’.

⁹ $\hat{\theta}_{K+1}$, the mean scalar on idiosyncratic variance, is 1.

2. Using the universe median and variance of the OLS estimates as a prior—exposures Gaussian, idiosyncratic variance inverse gamma— $\mathbf{E}$ and $\boldsymbol{\epsilon}^2$ are inferred through Bayesian regression, yielding $\hat{\mathbf{E}}$, $\hat{\boldsymbol{\epsilon}}^2$, $\hat{\mathbf{X}}$, and $\hat{\boldsymbol{\omega}}$. ($\hat{\mathbf{X}}$ assumes no exposure covariance between securities.)
3. $\hat{\boldsymbol{\Omega}}$, the covariance of factor variances and idiosyncratic variance scalar, is built in two steps. The K rows and columns corresponding to factors are the standard error.¹⁰ For idiosyncratic variance, a measure of its level is the cross-sectional median of each period's squared OLS regression residuals, divided by the mean across time for normalization. Its uncertainty is the variance of this series. The level is assumed uncorrelated with the factor variances.

The experiment selects 50 stocks at random and constructs fully-invested, long-only minimum variance portfolios eight ways:

Equally-weighted

Two percent in each of the 50 securities.

OLS (Conventional)

Minimum variance portfolio under covariance point-estimated from OLS exposures and idiosyncratic variances.

Bayes (Conventional)

Minimum variance portfolio under covariance point-estimated from Bayes posterior mean exposures and idiosyncratic variances.

$\kappa = 0$ Expected variance

Minimum expected¹¹ variance portfolio but no aversion to uncertainty.

¹⁰For the usual unbiased variance estimate of Gaussian x, y from n points, $\text{cov}[\hat{\sigma}_x^2, \hat{\sigma}_y^2] = 2\sigma_{x,y}^2/(n-1)$. Standard error replaces $\sigma_{x,y}$ with $\hat{\sigma}_{x,y}$.

¹¹Adds the variance term in equation (5.41) to the Bayes point estimate.

$\kappa > 0$ Expected variance with uncertainty aversion

For each κ in $\{0.1, 1, 10, 100\}$, a local minimizer of $E[\text{variance}] + \kappa \times \text{stdev}[\text{variance}]$ reached by searching from the $\kappa = 0$ portfolio.

Realized volatility is measured on the testing set. The experiment is repeated 10,000 times.¹²

Table 5.1 shows realized volatility quantiles for each method. Equally-weighted was dominated by OLS, which was dominated by Bayes, which was dominated by the uncertainty methods, whose results improved uniformly with increasing κ .

Table 5.1: Realized portfolio std dev quantiles by method (annualized)

	Quantile						
	0.02	0.05	0.25	0.50	0.75	0.95	0.98
Equal wt	10.86	11.41	12.78	13.79	14.84	16.47	17.23
OLS	6.41	6.78	7.72	8.50	9.46	11.43	12.36
Bayes	6.29	6.64	7.51	8.21	9.07	10.81	11.61
$\kappa = 0$	6.29	6.63	7.49	8.19	9.03	10.73	11.52
$\kappa = 0.1$	6.28	6.62	7.48	8.18	9.01	10.70	11.48
$\kappa = 1$	6.26	6.59	7.43	8.09	8.88	10.47	11.22
$\kappa = 10$	6.21	6.51	7.24	7.82	8.51	9.78	10.51
$\kappa = 100$	6.17	6.47	7.15	7.70	8.32	9.49	10.16

Table 5.2 compares methods head-to-head. Values are the fraction of experiments where the column method realized lower volatility than the row method. Uncertain methods beat each certain method in more than 80% of experiments. Increasing uncertainty consideration almost uniformly improved winning percentage through $\kappa = 10$. $\kappa = 100$ won 80% of experiments against $\kappa = 10$, but the fraction against other methods fell slightly.

Table 5.3 shows median expected variance¹³, uncertainty, and realized variance. Uncertainty falls and expected variance rises with increasing κ , a mathematical truth from the optimization objectives. Repeating the information from table 5.1: *realized*

¹²Total run time for $10,000 \times 8$ portfolios is under 20 minutes using the SLSQP solver in `scipy.optimize.minimize` Virtanen et al. [2020] on a circa 2015 laptop with 4th Generation Intel i7 CPU (i7-4710HQ @ 2.5GHz).

¹³Expected variance and uncertainty are via the uncertain covariance model.

Table 5.2: Head-to-head comparison of realized volatility

	Fraction of tests method had lower realized volatility						
	OLS	Bayes	$\kappa = 0$	$\kappa = 0.1$	$\kappa = 1$	$\kappa = 10$	$\kappa = 100$
Equal wt	0.997	0.999	0.999	0.999	1.000	1.000	1.000
OLS		0.890	0.903	0.906	0.913	0.903	0.894
Bayes			0.808	0.824	0.842	0.850	0.847
$\kappa = 0$				0.842	0.848	0.852	0.846
$\kappa = 0.1$					0.848	0.850	0.845
$\kappa = 1$						0.842	0.837
$\kappa = 10$							0.798

variance fell. As with all the results, the cause could be simple regularization or, as tables 5.4 and 5.5 suggest, uncertainty modeling's skill.

Table 5.3: Median expected variance, uncertainty, and realized variance (weekly, not annualized)

Median	Equal wt	OLS	Bayes	$\kappa = 0$	$\kappa = 0.1$	$\kappa = 1$	$\kappa = 10$	$\kappa = 100$
Expected variance	4.99	1.72	1.70	1.70	1.70	1.71	1.74	1.79
Uncertainty, sd(var)	0.77	0.25	0.24	0.24	0.24	0.24	0.22	0.22
Realized variance	3.66	1.39	1.30	1.29	1.29	1.26	1.18	1.14

Table 5.4 shows quantiles of the uncertain covariance model's prediction bias. Higher κ overpredicted more but had a narrower error range and smaller errors of underprediction.

Table 5.5 shows quantiled z-scored predictions: (expected - realized variance) / (stdev of variance). The range between low and high quantile shrank slightly from ~ 8 with $\kappa = 0$ to ~ 6 with $\kappa = 100$. Note that the opposite would occur without an increase in accuracy since higher κ corresponds to lower uncertainty in the denominator as well as higher expected variance. That the range didn't collapse or blow up suggests uncertainty predictions are on the correct scale and not naive regularization.

Table 5.4: Expected minus realized annualized standard deviation

	Quantile						
	0.02	0.05	0.25	0.50	0.75	0.95	0.98
Equal wt	-0.68	-0.03	1.41	2.30	3.20	4.41	4.85
OLS	-3.21	-2.36	-0.10	0.99	1.90	3.05	3.56
Bayes	-2.52	-1.71	0.24	1.23	2.07	3.17	3.66
$\kappa = 0$	-2.42	-1.61	0.29	1.25	2.09	3.17	3.65
$\kappa = 0.1$	-2.37	-1.57	0.31	1.26	2.09	3.18	3.65
$\kappa = 1$	-2.06	-1.28	0.45	1.35	2.15	3.22	3.68
$\kappa = 10$	-0.97	-0.41	0.94	1.70	2.43	3.41	3.86
$\kappa = 100$	-0.46	0.09	1.25	1.97	2.66	3.59	4.01

Table 5.5: Z-scores: $\frac{\text{predicted} - \text{realized variance}}{\text{uncertainty (std dev) of variance}}$

	Quantile						
	0.02	0.05	0.25	0.50	0.75	0.95	0.98
Equal wt	-0.56	-0.03	1.09	1.72	2.31	3.04	3.29
OLS	-5.12	-3.61	-0.15	1.36	2.42	3.53	3.96
Bayes	-4.04	-2.69	0.36	1.70	2.70	3.78	4.16
$\kappa = 0$	-3.91	-2.54	0.43	1.74	2.74	3.82	4.20
$\kappa = 0.1$	-3.86	-2.48	0.45	1.77	2.75	3.83	4.22
$\kappa = 1$	-3.38	-2.08	0.67	1.90	2.88	3.94	4.31
$\kappa = 10$	-1.70	-0.71	1.47	2.53	3.37	4.38	4.70
$\kappa = 100$	-0.80	0.14	2.00	2.95	3.74	4.67	4.99

5.9 Summary

The framework in this paper is a practical and mathematically tight way to model covariance with uncertainty. Uncertainty about individual pieces—for example, that the market beta of stock X is more reliably forecasted than stock Y’s—propagates to the portfolio level, and a portfolio variance forecast, no longer a fixed number, has a mean and variance.

Models can be constructed for a broad set of securities and easily shared. The method requires no simulation and is computationally fast, and optimization can occur beyond agreeably structured cases. The method is also easily extended to incorporate uncertainty about future regimes (section 5.11).

Among the many applications, one can, for example, choose the most reliably predicted from possible hedging securities. Or articulating the issue of estimation error, a computer can penalize uncertainty during optimization to find portfolios believed more likely to perform as expected.

5.10 Uncertain factor-modeled $E[\Sigma_{i,j}]$ and $\text{cov}[\Sigma_{i,j}, \Sigma_{m,n}]$

Proposition 5.2. *Mean and covariance of entries of Σ*

If

1. $\Sigma = \mathbf{E}\mathbf{F}\mathbf{E}^T + \gamma \text{diag}[\boldsymbol{\varepsilon}^2]$
2. $\mathbf{E}, \boldsymbol{\varepsilon}^2$, and $(\mathbf{F}, \gamma)$ are jointly independent
3. The elements of $\boldsymbol{\varepsilon}^2$ are independent
4. $\mathbf{E}$ is Gaussian
5. Mean and covariance of $\mathbf{E}, \boldsymbol{\varepsilon}^2$, and $(\mathbf{F}, \gamma)$ are:

Exposures

$$\begin{aligned} \hat{\mathbf{E}} &= (N \times K) & \hat{\mathbf{E}} &= E[\mathbf{E}] \\ \hat{\mathbf{X}} &= (N \times K \times N \times K) & \hat{X}_{m,i,n,j} &= \text{cov}[E_{m,i}, E_{n,j}] \end{aligned}$$

Idiosyncratic effects

$$\begin{aligned} \hat{\boldsymbol{\varepsilon}}^2 &= (N \times 1) & \hat{\boldsymbol{\varepsilon}}^2 &= E[\boldsymbol{\varepsilon}^2] \\ \hat{\boldsymbol{\omega}} &= (N \times 1) & \hat{\omega}_n &= \text{var}[\varepsilon_n^2] \end{aligned}$$

Factor structure

$$\begin{aligned} \hat{\mathbf{F}} &= (K \times K) & \hat{\mathbf{F}} &= E[\mathbf{F}] \\ \hat{\mathbf{H}} &= (K \times K \times K \times K) & \hat{H}_{i,j,k,l} &= \text{cov}[F_{i,j}, F_{k,l}] \\ \hat{\gamma} &= (\text{scalar}) & \hat{\gamma} &= E[\gamma] \\ \hat{\eta} &= (\text{scalar}) & \hat{\eta} &= \text{var}[\gamma] \\ \hat{\mathbf{h}} &= (K \times K) & \hat{h}_{i,j} &= \text{cov}[F_{i,j}, \gamma] \end{aligned}$$

Then mean and covariance of the covariance matrix entries are

$$\mathbb{E}[\Sigma] = \hat{\mathbf{E}}\hat{\mathbf{F}}\hat{\mathbf{E}}^T + \hat{\gamma} \text{diag}[\hat{\varepsilon}^2] + \sum_{k,l} \hat{F}_{k,l} \hat{\mathbf{X}}_{\cdot,k,\cdot,l} \quad (5.56)$$

$$\begin{aligned} \text{cov}[\Sigma_{i,j}, \Sigma_{m,n}] = & \sum_{k,l,s,t} \left[\hat{E}_{i,k} \hat{E}_{m,s} \hat{X}_{j,l,n,t} + \hat{E}_{j,l} \hat{E}_{n,t} \hat{X}_{i,k,m,s} \right. \\ & + \hat{X}_{i,k,m,s} \hat{X}_{j,l,n,t} \\ & + \hat{E}_{i,k} \hat{E}_{n,t} \hat{X}_{j,l,m,s} + \hat{E}_{j,l} \hat{E}_{m,s} \hat{X}_{i,k,n,t} \\ & \left. + \hat{X}_{i,k,n,t} \hat{X}_{j,l,m,s} \right] \left(\hat{F}_{k,l} \hat{F}_{s,t} + \hat{H}_{k,l,s,t} \right) \\ & + \sum_{k,l,s,t} \left[\hat{E}_{i,k} \hat{E}_{j,l} + \hat{X}_{i,k,j,l} \right] \left[\hat{E}_{m,s} \hat{E}_{n,t} + \hat{X}_{m,s,n,t} \right] \hat{H}_{k,l,s,t} \\ & + 1_{\{i=j\}} \hat{\varepsilon}_i^2 \sum_{s,t} \left[\hat{E}_{m,s} \hat{E}_{n,t} + \hat{X}_{m,s,n,t} \right] \hat{h}_{s,t} \\ & + 1_{\{m=n\}} \hat{\varepsilon}_m^2 \sum_{k,l} \left[\hat{E}_{i,k} \hat{E}_{j,l} + \hat{X}_{i,k,j,l} \right] \hat{h}_{k,l} \\ & + 1_{\{i=j\}} 1_{\{m=n\}} \hat{\varepsilon}_i^2 \hat{\varepsilon}_m^2 \hat{\eta} \\ & + 1_{\{i=j=m=n\}} \hat{\omega}_i (\hat{\gamma}^2 + \hat{\eta}) \end{aligned} \quad (5.57)$$

Proof.

$$\Sigma_{i,j} = \sum_{k,l} E_{i,k} E_{j,l} F_{k,l} + 1_{\{i=j\}} \varepsilon_i^2 \gamma \quad (5.58)$$

$$\mathbb{E}[\Sigma_{i,j}] = \sum_{k,l} \mathbb{E}[E_{i,k} E_{j,l}] \mathbb{E}[F_{k,l}] + 1_{\{i=j\}} \mathbb{E}[\varepsilon_i^2] \mathbb{E}[\gamma] \quad (5.59)$$

$$\begin{aligned} &= \sum_{k,l} (\mathbb{E}[E_{i,k}] \mathbb{E}[E_{j,l}] + \text{cov}[E_{i,k}, E_{j,l}]) \mathbb{E}[F_{k,l}] \\ &\quad + 1_{\{i=j\}} \mathbb{E}[\varepsilon_i^2] \mathbb{E}[\gamma] \end{aligned} \quad (5.60)$$

$$= \sum_{k,l} \left[\hat{E}_{i,k} \hat{E}_{j,l} + \hat{X}_{i,k,j,l} \right] \hat{F}_{k,l} + 1_{\{i=j\}} \hat{\varepsilon}_i^2 \hat{\gamma} \quad (5.61)$$

$$\text{so } \mathbb{E}[\Sigma] = \hat{\mathbf{E}}\hat{\mathbf{F}}\hat{\mathbf{E}}^T + \hat{\gamma} \text{diag}[\hat{\varepsilon}^2] + \sum_{k,l} \hat{F}_{k,l} \underbrace{\hat{\mathbf{X}}_{\cdot,k}(\cdot, l)}_{N \times N} \quad (5.62)$$

$$\text{cov} [\Sigma_{i,j}, \Sigma_{m,n}] = \underbrace{\text{cov} [\text{E} [\Sigma_{i,j} | \mathbf{F}, \gamma], \text{E} [\Sigma_{m,n} | \mathbf{F}, \gamma]]}_{\textcircled{1}} + \underbrace{\text{E} [\text{cov} [\Sigma_{i,j}, \Sigma_{m,n} | \mathbf{F}, \gamma]]}_{\textcircled{2}} \quad (5.63)$$

$$\begin{aligned} \textcircled{1} = \text{cov} & \left[\sum_{k,l} \text{E} [E_{i,k} E_{j,l}] F_{k,l} + 1_{\{i=j\}} \text{E} [\varepsilon_i^2] \gamma, \right. \\ & \left. \sum_{s,t} \text{E} [E_{m,s} E_{n,t}] F_{s,t} + 1_{\{m=n\}} \text{E} [\varepsilon_m^2] \gamma \right] \end{aligned} \quad (5.64)$$

$$\begin{aligned} &= \sum_{k,l,s,t} \text{E} [E_{i,k} E_{j,l}] \text{E} [E_{m,s} E_{n,t}] \text{cov} [F_{k,l}, F_{s,t}] \\ &+ 1_{\{i=j\}} \text{E} [\varepsilon_i^2] \sum_{s,t} \text{E} [E_{m,s} E_{n,t}] \text{cov} [F_{s,t}, \gamma] \\ &+ 1_{\{m=n\}} \text{E} [\varepsilon_m^2] \sum_{k,l} \text{E} [E_{i,k} E_{j,l}] \text{cov} [F_{k,l}, \gamma] \\ &+ 1_{\{i=j\}} 1_{\{m=n\}} \text{E} [\varepsilon_i^2] \text{E} [\varepsilon_m^2] \text{var} [\gamma] \end{aligned} \quad (5.65)$$

$$\begin{aligned} &= \sum_{k,l,s,t} \left[\hat{E}_{i,k} \hat{E}_{j,l} + \hat{X}_{i,k,j,l} \right] \left[\hat{E}_{m,s} \hat{E}_{n,t} + \hat{X}_{m,s,n,t} \right] \hat{H}_{k,l,s,t} \\ &+ 1_{\{i=j\}} \hat{\varepsilon}_i^2 \sum_{s,t} \left[\hat{E}_{m,s} \hat{E}_{n,t} + \hat{X}_{m,s,n,t} \right] \hat{h}_{s,t} \\ &+ 1_{\{m=n\}} \hat{\varepsilon}_m^2 \sum_{k,l} \left[\hat{E}_{i,k} \hat{E}_{j,l} + \hat{X}_{i,k,j,l} \right] \hat{h}_{k,l} \\ &+ 1_{\{i=j\}} 1_{\{m=n\}} \hat{\varepsilon}_i^2 \hat{\varepsilon}_m^2 \hat{\eta} \end{aligned} \quad (5.66)$$

$$\begin{aligned} \textcircled{2} = \mathbb{E} & \left[\text{cov} \left(\sum_{k,l} E_{i,k} E_{j,l} F_{k,l} + 1_{\{i=j\}} \varepsilon_i^2 \gamma, \right. \right. \\ & \left. \left. \sum_{s,t} E_{m,s} E_{n,t} F_{s,t} + 1_{\{m=n\}} \varepsilon_m^2 \gamma \mid \mathbf{F}, \gamma \right) \right] \end{aligned} \quad (5.67)$$

$$= \mathbb{E} \left[\sum_{k,l,s,t} \text{cov} [E_{i,k} E_{j,l}, E_{m,s} E_{n,t}] F_{k,l} F_{s,t} + 1_{\{i=j=m=n\}} \text{var} [\varepsilon_i^2] \gamma^2 \right] \quad (5.68)$$

$$\begin{aligned} &= \sum_{k,l,s,t} \text{cov} [E_{i,k} E_{j,l}, E_{m,s} E_{n,t}] \mathbb{E} [F_{k,l} F_{s,t}] \\ &\quad + 1_{\{i=j=m=n\}} \text{var} [\varepsilon_i^2] \mathbb{E} [\gamma^2] \end{aligned} \quad (5.69)$$

$$\begin{aligned} &= \sum_{k,l,s,t} \underbrace{\text{cov} [E_{i,k} E_{j,l}, E_{m,s} E_{n,t}]}_{\textcircled{3}} \left(\hat{F}_{k,l} \hat{F}_{s,t} + \hat{H}_{k,l,s,t} \right) \\ &\quad + 1_{\{i=j=m=n\}} \hat{\omega}_i (\hat{\gamma}^2 + \hat{\eta}) \end{aligned} \quad (5.70)$$

$$\begin{aligned} \textcircled{3} &= \text{cov} [E_{i,k} E_{j,l}, E_{m,s} E_{n,t}] \\ &= \hat{E}_{i,k} \hat{E}_{m,s} \hat{X}_{j,l,n,t} + \hat{E}_{j,l} \hat{E}_{n,t} \hat{X}_{i,k,m,s} + \hat{X}_{i,k,m,s} \hat{X}_{j,l,n,t} \\ &\quad + \hat{E}_{i,k} \hat{E}_{n,t} \hat{X}_{j,l,m,s} + \hat{E}_{j,l} \hat{E}_{m,s} \hat{X}_{i,k,n,t} + \hat{X}_{i,k,n,t} \hat{X}_{j,l,m,s} \end{aligned} \quad (5.71)$$

via Lemma 7.1 since $\mathbf{E}$ Gaussian

□

5.11 Incorporating structural uncertainty

Uncertainty so far has taken the flavor of estimation error, ambiguity in the value of parameters on an assumed-correct skeleton. A different, macro scale uncertainty exists in how events might unfold, the possibility of alternative regimes whose centers (configuration of factor variances) or rules (factor structures) differ by more than small perturbations so need to be modeled separately. Considering the possibilities is a view into a portfolio's robustness.

Suppose the future is a random draw from S scenarios, each described by a separate uncertain covariance model. Unconditional mean and variance of portfolio variance are assembled from forecasts for each scenario.

Assume mean returns are the same in each scenario.

Let π_s = probability of scenario s , $s = 1 \dots S$

$$m_s(\mathbf{w}) = \text{E} [\sigma_{\mathbf{w}}^2 | \text{scenario } s]$$

$$v_s(\mathbf{w}) = \text{var} [\sigma_{\mathbf{w}}^2 | \text{scenario } s]$$

$$\begin{aligned} \text{Then } \text{E} [\sigma_{\mathbf{w}}^2] &= \text{E} [\text{E} [\sigma_{\mathbf{w}}^2 | s]] \\ &= \sum_s \pi_s m_s \end{aligned} \quad (5.72)$$

$$\text{and } \text{var} [\sigma_{\mathbf{w}}^2] = \text{E} [\text{var} [\sigma_{\mathbf{w}}^2 | s]] + \text{var} [\text{E} [\sigma_{\mathbf{w}}^2 | s]] \quad (5.73)$$

$$= \text{E} [\text{var} [\sigma_{\mathbf{w}}^2 | s]] + \text{E} [\text{E}^2 [\sigma_{\mathbf{w}}^2 | s]] - \text{E}^2 [\text{E} [\sigma_{\mathbf{w}}^2 | s]] \quad (5.74)$$

$$= \sum_s \pi_s v_s + \sum_s \pi_s m_s^2 - \left[\sum_s \pi_s m_s \right]^2 \quad (5.75)$$

5.12 Mathematical results

Lemma 5.1. *Variance of inner product of independent variables*

For $\mathbf{a}, \mathbf{b}$ independent,

$$\begin{aligned} \text{var} [\mathbf{a}^T \mathbf{b}] &= \text{E} [\mathbf{b}]^T \text{cov} [\mathbf{a}] \text{E} [\mathbf{b}] \\ &\quad + \text{E} [\mathbf{a}]^T \text{cov} [\mathbf{b}] \text{E} [\mathbf{a}] + \sum_{i,j} (\text{cov} [\mathbf{a}])_{i,j} (\text{cov} [\mathbf{b}])_{i,j} \end{aligned} \quad (5.76)$$

Proof.

$$\text{var} [\mathbf{a}^T \mathbf{b}] = \text{var} [\mathbb{E} [\mathbf{a}^T \mathbf{b} | \mathbf{a}]] + \mathbb{E} [\text{var} [\mathbf{a}^T \mathbf{b} | \mathbf{a}]] \quad (5.77)$$

$$= \text{var} [\mathbf{a}^T \mathbb{E} [\mathbf{b}]] + \mathbb{E} [\mathbf{a}^T \text{cov} [\mathbf{b}] \mathbf{a}] \text{ since } \mathbf{a} \text{ ind. of } \mathbf{b} \quad (5.78)$$

$$= \mathbb{E} [\mathbf{b}]^T \text{cov} [\mathbf{a}] \mathbb{E} [\mathbf{b}] + \mathbb{E} [\mathbf{a}^T \text{cov} [\mathbf{b}] \mathbf{a}] \quad (5.79)$$

$$= \mathbb{E} [\mathbf{b}]^T \text{cov} [\mathbf{a}] \mathbb{E} [\mathbf{b}] + \mathbb{E} \left[\sum_{i,j} a_i a_j (\text{cov} [\mathbf{b}])_{i,j} \right] \quad (5.80)$$

$$= \mathbb{E} [\mathbf{b}]^T \text{cov} [\mathbf{a}] \mathbb{E} [\mathbf{b}] + \sum_{i,j} \mathbb{E} [a_i a_j] (\text{cov} [\mathbf{b}])_{i,j} \quad (5.81)$$

$$\begin{aligned} &= \mathbb{E} [\mathbf{b}]^T \text{cov} [\mathbf{a}] \mathbb{E} [\mathbf{b}] \\ &\quad + \sum_{i,j} (\mathbb{E} [a_i] \mathbb{E} [a_j] + \text{cov} [a_i, a_j]) (\text{cov} [\mathbf{b}])_{i,j} \end{aligned} \quad (5.82)$$

$$\begin{aligned} &= \mathbb{E} [\mathbf{b}]^T \text{cov} [\mathbf{a}] \mathbb{E} [\mathbf{b}] \\ &\quad + \mathbb{E} [\mathbf{a}]^T \text{cov} [\mathbf{b}] \mathbb{E} [\mathbf{a}] + \sum_{i,j} (\text{cov} [\mathbf{a}])_{i,j} (\text{cov} [\mathbf{b}])_{i,j} \end{aligned} \quad (5.83)$$

□

Lemma 5.2.

$$For \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma}),$$

$$\begin{aligned} E[x_1 x_2 x_3 x_4] &= \mu_1 \mu_2 \mu_3 \mu_4 \\ &\quad + \mu_1 \mu_2 \sigma_{3,4} + \mu_3 \mu_4 \sigma_{1,2} + \sigma_{1,2} \sigma_{3,4} \\ &\quad + \mu_1 \mu_3 \sigma_{2,4} + \mu_2 \mu_4 \sigma_{1,3} + \sigma_{1,3} \sigma_{2,4} \\ &\quad + \mu_1 \mu_4 \sigma_{2,3} + \mu_2 \mu_3 \sigma_{1,4} + \sigma_{1,4} \sigma_{2,3} \end{aligned} \quad (5.84)$$

$$\begin{aligned} \text{COV}[x_1 x_2, x_3 x_4] &= \mu_1 \mu_3 \sigma_{2,4} + \mu_2 \mu_4 \sigma_{1,3} + \sigma_{1,3} \sigma_{2,4} \\ &\quad + \mu_1 \mu_4 \sigma_{2,3} + \mu_2 \mu_3 \sigma_{1,4} + \sigma_{1,4} \sigma_{2,3} \end{aligned} \quad (5.85)$$

$$\begin{aligned} &= (\mu_1 \mu_3 + \sigma_{1,3}) (\mu_2 \mu_4 + \sigma_{2,4}) \\ &\quad + (\mu_1 \mu_4 + \sigma_{1,4}) (\mu_2 \mu_3 + \sigma_{2,3}) - 2\mu_1 \mu_3 \mu_2 \mu_4 \end{aligned} \quad (5.86)$$

Proof.

$$\begin{aligned}
 \mathbb{E}[x_1 x_2 x_3 x_4] &= \mathbb{E} \left[\prod_{i=1}^4 [\mu_i + (x_i - \mu_i)] \right] \\
 &= \mu_1 \mu_2 \mu_3 \mu_4 \\
 &\quad + \mu_1 \mu_2 \sigma_{3,4} + \mu_3 \mu_4 \sigma_{1,2} \\
 &\quad + \mu_1 \mu_3 \sigma_{2,4} + \mu_2 \mu_4 \sigma_{1,3} \\
 &\quad + \mu_1 \mu_4 \sigma_{2,3} + \mu_2 \mu_3 \sigma_{1,4} \\
 &\quad + \mathbb{E} \left[\prod_{i=1}^4 (x_i - \mu_i) \right] \\
 &\text{since odd moments are 0}
 \end{aligned}$$

$$\mathbb{E} \left[\prod_{i=1}^4 (x_i - \mu_i) \right] = \sigma_{1,2} \sigma_{3,4} + \sigma_{1,3} \sigma_{2,4} + \sigma_{1,4} \sigma_{2,3}$$

by Isserlis' theorem since 0 mean Gaussians

$$\text{cov}[x_1 x_2, x_3 x_4] = \mathbb{E}[x_1 x_2 x_3 x_4] - \underbrace{\mathbb{E}[x_1 x_2]}_{(\mu_1 \mu_2 + \sigma_{12})} \underbrace{\mathbb{E}[x_3 x_4]}_{(\mu_3 \mu_4 + \sigma_{34})}$$

□

Corollary 5.1. *Covariance of squared Gaussians*

$$\text{For } \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma}), \text{cov}[x_1^2, x_2^2] = 4\mu_1 \mu_2 \sigma_{1,2} + 2\sigma_{1,2}^2$$

CHAPTER 6

UNCERTAIN RISK PARITY

Introduced in the alternative portfolio construction methods of section 2.2.3, risk parity determines portfolio weights using covariance alone. It is typically done at the index or asset class level rather than on individual stocks, e.g., one security is the S&P 500 index, another is a corporate bond index. Risk parity finds the unique solution where each security matches its assigned risk budget (repeating equation 2.10),

$$w^*(\Sigma, \sigma^2, f) \equiv w > 0 \text{ s.t. } w \odot \Sigma w = \sigma^2 f \quad (6.1)$$

where $\sigma^2 \equiv$ target variance of portfolio

$f \equiv$ fraction of risk budgeted to each security, $f \in \mathbb{R}_+^n$

$\odot \equiv$ element-by-element multiplication

which happens to be equivalent to minimizing variance with a regularization penalty against small weights (repeating equation 2.11),

$$w^*(\Sigma, \sigma^2, f) = \arg \min_{w > 0} \frac{1}{2} w^T \Sigma w - \sigma^2 \sum_i f_i \log w_i \quad (6.2)$$

Working from a point-estimate of covariance leaves the method vulnerable to estimation error.

6.1 Literature review

Risk parity without uncertainty is covered in section 2.2.3 under ‘Risk parity’. Starting from the regularized minimum variance formulation of equation (6.2) with the penalty moved to constraints, Costa and Kwon [2019] minimizes the worst case variance where the covariance matrix comes from a Markov chain, each state having factor-modeled covariance with elliptical confidence sets on exposures and bounds on stock-specific variance. Costa and Kwon [2020a] minimizes the worst case error between variance risk contributions and budgeted risk when covariance falls within box

constraints. [Costa and Kwon \[2020b\]](#) adds a return forecast term having an elliptical uncertainty set. The authors' collected work on risk parity is in a dissertation [[Costa, 2021](#)]. [Kapsos et al. \[2018\]](#) considers robust optimization of risk parity under several uncertainty scenarios. [Kim and Kim \[2021\]](#) applies the resampling idea of [Michaud \[1998\]](#) with covariance assumed Wishart centered at the unbiased time series estimate and the final risk parity portfolio as the average of risk parity portfolios constructed from sampled matrices inversely weighted by sampled matrix condition number. The basis of this chapter, [Shah \[2021a\]](#), looks at the difference between variance contributions and budgeted risk as in [Costa and Kwon \[2020a\]](#), but models covariance as a distribution and minimizes functions of mean squared error.

6.2 Mathematical model

In the form of equation (6.1), the goal is a solution to n equations that are quadratic in the decision variable w and linear in the data Σ . Normally Σ is replaced by a point-estimate, and there is a unique solution since Σ is regarded as fixed. When Σ is random, there isn't a solution.

The paper develops an approach given $E[\Sigma]$ and $\text{cov}[\Sigma]$. These could be supplied directly or, since risk parity is generally at the index level, derived from portfolios of securities modeled by an uncertain covariance model as in chapter 5.

In the latter case, as in equation (5.4), the covariance of individual stock returns is expressed in terms of state variables,

$$\text{cov}[Y_t|\beta, \Theta, \gamma, \Lambda] \equiv \beta \text{diag}[\Theta] \beta^T + \gamma \Lambda$$

If each row of matrix p represents the portfolio weights of a different index,

$$\Sigma_t \equiv \text{cov}[p Y_t|\beta, \Theta, \gamma, \Lambda] \tag{6.3}$$

$$\equiv p (\beta \text{diag}[\Theta] \beta^T + \gamma \Lambda) p^T \tag{6.4}$$

$E[\Sigma]$ and $\text{cov}[\Sigma]$, after assumptions on independence and distributional shape, are functions of mean and covariance of the parameters.

from [Shah \[2021a\]](#) Uncertain Risk Parity

6.3 Abstract

Risk parity is a portfolio construction technique that scales sections of a portfolio—e.g., stocks, bonds, currencies, commodities—so that forecasted contributions to net portfolio risk match the budget. Because risks are measured from a point-estimate of covariance, the method is subject to problems of estimation error. This paper treats covariance as uncertain in order to find a risk parity weighting that doesn't count on perfectly optimized hedges and is robust to changes in regime.

Separately, of general interest are the uncertain risk contributions calculated en route. Reporting a portfolio's uncertain risk decomposition puts a band around numbers and reveals fragility. For example, market could seem hedged in a long-short portfolio but surface as the biggest risk when parameters are considered across their error range.

6.4 Introduction

Risk parity, named by [Qian \[2005\]](#), is a method of portfolio construction that, from risk alone, scales sections of a portfolio—e.g., stocks, levered bonds, currencies, commodities—so each section's forecasted contribution to net portfolio risk matches the (typically equally-weighted) risk budget. The idea without considering correlation dates to 1996 in Bridgewater All Weather Fund's weighting investments by inverse of volatility.

One motivation, as in [Qian \[2011\]](#), is capturing the risk premium across a diversified basket of asset classes. Another is correcting problems of the canonical example of risk-only portfolio construction, minimum variance, whose composition is sensitive to inputs and can be concentrated.

Absent bounds, minimum variance is the point where marginal variances for investments are equal. When risk parity equalizes risk contributions—proportional to

marginal variance $\times$ *weight*¹—small positions grow, large ones shrink, and the portfolio is less concentrated and less sensitive to inputs. Yet, operating from a single estimated covariance matrix, it remains vulnerable to estimation error.

While this paper investigates risk parity under uncertainty, risk parity itself can be explicitly cast as an approach to minimum variance under uncertainty. [Maillard et al. \[2010\]](#) shows that risk parity is the solution to minimum variance with a regularization penalty. The penalty can be regarded as Bayesian tempering data with a prior. From this perspective, [de Jong \[2018\]](#) sees risk parity as an intermediate optimal minimum risk portfolio on the spectrum from minimum variance (no uncertainty in estimates, solution dominated by data) to equal-weighted (complete uncertainty, solution dominated by non-informative prior).

This paper contributes four ideas to the literature: 1—justification for estimation error concerns in measuring risk contributions, 2—a framework and algorithm to solve risk parity with uncertainty, 3—a formula to measure risk contributions (decompose risk) with uncertainty, and 4—a way to model covariance uncertainty between portfolios from the securities within.

6.4.1 Literature review

[Maillard et al. \[2010\]](#) and [Bai et al. \[2016\]](#) explore mathematical properties and alternative formulations of risk parity. Instead of equalizing risk contributions across investments, [Lohre et al. \[2014\]](#) equalizes them across principal component directions of the covariance matrix, and [Roncalli and Weisang \[2016\]](#) across directions specified by any factor model. [Bernardi et al. \[2019\]](#) looks at risk parity’s ‘second order risk’, the ratio of realized to predicted risk, essentially the bias from estimate-optimized selection. Closest to this paper, [Costa and Kwon \[2019\]](#) uses factor-modeled covariance and solves risk parity under uncertainty and regime switching through robust (minimax) optimization with elliptical confidence sets on exposures and bounds on stock-specific variance. On risk parity but not related is the large literature on empirical performance in both return and volatility. This paper concerns construction, volatility but not return, and individual volatility contributions as much as the combined portfolio’s.

¹See equation (6.5)

6.4.2 Problem statement

Budgeted fractions of total portfolio risk are assigned to N investments, each of which can itself be a portfolio. Using an estimate of covariance to measure risk, risk parity finds a weighting that meets the budget. The inputs and decision variables are:

$\mathbf{f} = (N \times 1)$ budgeted fractions of portfolio risk

$$\mathbf{f} > 0, \sum_i f_i = 1$$

$\mathbf{C} = (N \times N)$ covariance between investments

$\mathbf{w} = (N \times 1)$ weights, usually constrained ≥ 0

One other possible input is target portfolio variance, σ^2 , to fix the scale: $\mathbf{w}^T \mathbf{C} \mathbf{w} = \sigma^2$. Otherwise, scale is set by forcing weights to sum to 1, $\mathbf{w}^T \mathbf{1} = 1$, which without uncertainty is the same as picking any target variance and dividing the solution by the sum of its weights.

Budgets can be enforced on contributions to portfolio standard deviation or variance, arriving at the same result.² The derivation is repeated here because both forms appear in what follows.

Contributions to standard deviation are calculated by Euler's theorem applied to portfolio standard deviation, which is homogenous of degree 1 in the weights³.

$s_i = \text{std dev contribution of } i\text{'th investment}$

$$\begin{aligned} &= w_i \frac{\partial}{\partial w_i} (\text{portfolio std dev}) \\ &= w_i \frac{\partial}{\partial w_i} \sqrt{\mathbf{w}^T \mathbf{C} \mathbf{w}} \\ &= \frac{w_i}{2\sqrt{\mathbf{w}^T \mathbf{C} \mathbf{w}}} \frac{\partial}{\partial w_i} \mathbf{w}^T \mathbf{C} \mathbf{w} \end{aligned} \tag{6.5}$$

$$= w_i (\mathbf{C} \mathbf{w})_i / \sqrt{\mathbf{w}^T \mathbf{C} \mathbf{w}} \tag{6.6}$$

So risk parity in standard deviation is

$$\sigma f_i = w_i (\mathbf{C} \mathbf{w})_i / \sqrt{\mathbf{w}^T \mathbf{C} \mathbf{w}} \quad i = 1..N, \mathbf{w} > 0 \tag{6.7}$$

where $\sigma \mathbf{f} = (N \times 1)$ standard deviation budgets

²See section 3.2.3 of [Qian et al. \[2007\]](#)

³See [Maillard et al. \[2010\]](#)

Contributions to variance are covariance with the portfolio.

$$\begin{aligned}
 v_i &= \text{variance contribution of } i\text{'th investment} \\
 &= \begin{bmatrix} 0 \cdots 0 & w_i & 0 \cdots 0 \end{bmatrix} \mathbf{C}\mathbf{w} \\
 &= w_i (\mathbf{C}\mathbf{w})_i
 \end{aligned} \tag{6.8}$$

Risk parity in variance is

$$\begin{aligned}
 \sigma^2 f_i &= w_i (\mathbf{C}\mathbf{w})_i \quad i = 1..N, \mathbf{w} > 0 \\
 \text{where } \sigma^2 \mathbf{f} &= (N \times 1) \text{ variance budgets}
 \end{aligned} \tag{6.9}$$

Since $\mathbf{w}^T \mathbf{C}\mathbf{w} = \sigma^2$, equations (6.7) and (6.9) are identical.

6.5 Sensitivity to covariance forecast

Emerging entirely from a point-estimate of covariance taken as true, risk parity could be sensitive to estimation error. Here are two arguments the concern is valid.

6.5.1 As solution to minimum variance with regularization

[Maillard et al. \[2010\]](#) shows that risk parity is the solution to minimum variance with a regularization penalty against small positions.

$$\arg \min_{\mathbf{w} > 0} \frac{1}{2} \mathbf{w}^T \mathbf{C}\mathbf{w} - \sigma^2 \sum_i f_i \log(w_i) \tag{6.10}$$

Setting the gradient to 0 yields $(\mathbf{C}\mathbf{w})_i - \sigma^2 f_i / w_i = 0$ for each i , so $w_i (\mathbf{C}\mathbf{w})_i = \sigma^2 f_i$, the risk parity equations (6.9).

Since it can be cast as optimization, risk parity might encounter issues described in the large literature on optimization's bias and estimation error problems.

The optimization characterization is also the easiest way to compute risk parity, minimizing a smooth, strictly convex (when restricted to any orthant, $\log |w_i|$ in place of $\log w_i$) function that has a unique solution.

6.5.2 Non-robustness in risk contribution assignment

[Menchero and Davis \[2011\]](#) reach the risk contribution assignment of (6.6) by measuring behavior in the direction of the net portfolio.

σ_i = standard dev of i'th investment

$\rho_{i,w}$ = correlation of i'th investment and portfolio

$\sigma_i \rho_{i,w}$ = std dev of i'th investment $\times$ corr with portfolio

$$= \sigma_i \text{cov}(i, w) / (\sigma_i \sigma) \quad (6.11)$$

$$= w_i (\mathbf{C}\mathbf{w})_i / \sqrt{\mathbf{w}^T \mathbf{C} \mathbf{w}} \quad (6.12)$$

$$= s_i$$

What about behavior orthogonal to the portfolio?

$$\begin{aligned} s_i^\perp &= \text{length of orthogonal behavior of i'th investment} \\ &= \sigma_i \sqrt{1 - \rho_{i,w}^2} \end{aligned} \quad (6.13)$$

By definition netting to zero across holdings, *risks orthogonal to the net portfolio are not counted* however consequential they might be.

As an extreme case, imagine a market-neutral equity portfolio broken into long and short. Market is the greatest risk within each side and, perceived as net hedged, not counted in risk contributions. However, since estimates are imperfect and the true beta is not exactly zero, the imperfect alignment can make market the portfolio's biggest real risk. The phenomenon occurs, to different degrees, in every hedge in any portfolio.

Thus, in conventional risk parity—particularly since it's the product of optimization and optimization finds perceived hedges—the most important real risk contributions can go unmeasured.

6.6 Adding uncertainty

To counter sensitivity to a particular estimate, especially linear combinations that seem to annihilate a direction of risk, covariance is no longer regarded as fixed and

known. Instead, entries have a mean and covariance: $E[C_{ij}]$ and $\text{cov}[C_{i,j}, C_{k,l}]$ for every i, j, k, l .

Equation (6.8) gives an investment's variance contribution, $v_i = w_i (\mathbf{C}\mathbf{w})_i = w_i \sum_j C_{i,j} w_j$. A portfolio's variance is the sum of the contributions, $v_P = \mathbf{w}^T \mathbf{C} \mathbf{w} = \sum_i v_i$. Written in terms of the entries of $\mathbf{C}$, means and covariances are

$$E[v_i] = w_i \sum_j E[C_{i,j}] w_j \quad (6.14)$$

$$E[v_P] = \sum_i E[v_i] \quad (6.15)$$

$$\text{cov}[v_i, v_j] = w_i w_j \sum_{k,l} w_k w_l \text{cov}[C_{i,k}, C_{j,l}] \quad (6.16)$$

$$\text{cov}[v_i, v_P] = \sum_j \text{cov}[v_i, v_j] \quad (6.17)$$

$$\text{var}[v_P] = \sum_i \text{cov}[v_i, v_P] \quad (6.18)$$

6.6.1 Risk parity restated with uncertainty

Variability in uncertain quantities makes meeting targets, e.g., equation (6.9), impossible. A loss function assigns costs to deviations. Risk parity weights are found by minimizing expected loss.

Uncertainty also means that portfolio variance is random, so there is a difference between aiming toward fixed *target contributions* (target fraction $\times$ target portfolio variance) and *target fractions* of random portfolio variance (target fraction $\times$ random portfolio variance) with scale set by budget constraint or portfolio variance itself aimed toward a target. The deviations for each case:

1. Target contributions

$$\begin{aligned} \delta_i &= \text{investment } i\text{'s deviation from budgeted variance, } i = 1 \dots N \\ &= v_i - f_i \sigma^2 \end{aligned} \quad (6.19)$$

2. Target fractions

$$\begin{aligned} \delta_i &= \text{investment } i\text{'s deviation from budgeted fraction of variance, } i = 1 \dots N \\ &= v_i - f_i v_P \end{aligned} \quad (6.20)$$

normalized by a portfolio variance target

$$\delta_{N+1} = v_P - \sigma^2 \quad (6.21)$$

or a budget constraint

$$\sum_i w_i = 1 \quad (6.22)$$

Measuring loss by summed mean-squared error yields

$$\begin{aligned} l(\mathbf{w}) &= \text{risk parity loss under } \mathbf{w} \\ &= \sum_i \mathbb{E} [\delta_i^2] \end{aligned} \quad (6.23)$$

$$= \sum_i \mathbb{E}^2 [\delta_i] + \text{var} [\delta_i] \quad (6.24)$$

With no uncertainty, the variance terms are zero, and the solution is the same as standard risk parity.

Loss is computed from the quantities in equations (6.14) – (6.18):

Target contributions

$$\mathbb{E} [\delta_i] = \mathbb{E} [v_i] - f_i \sigma^2 \quad (6.25)$$

$$\text{var} [\delta_i] = \text{var} [v_i] \quad (6.26)$$

Target fractions

$$\mathbb{E} [\delta_i] = \mathbb{E} [v_i] - f_i \mathbb{E} [v_P] \quad (6.27)$$

$$\text{var} [\delta_i] = \text{var} [v_i] + f_i^2 \text{var} [v_P] - 2f_i \text{cov} [v_i, v_P] \quad (6.28)$$

$$\mathbb{E} [\delta_{N+1}] = \mathbb{E} [v_P] - \sigma^2 \quad (6.29)$$

$$\text{var} [\delta_{N+1}] = \text{var} [v_P] \quad (6.30)$$

Both characterizations are 4th degree polynomials of portfolio weights, whose (local) minimizer can be found using nonlinear optimization techniques.

6.6.2 Multiple regimes

In determining allocations, regimes—e.g., risk-on vs. risk-off—and scenarios naturally enter. Rather than representing variation in the entries, each has its own covariance

estimate to measure loss. With a probability assigned to each scenario, the distribution of $\mathbf{C}$ is a mixture. Combined loss is the probability weighted average of the losses.

$$\begin{aligned} l(\mathbf{w}) &= \text{risk parity loss under } \mathbf{w} \\ &= \sum_i \mathbb{E} [\delta_i^2] \end{aligned} \tag{6.31}$$

$$= \sum_i \mathbb{E} [\mathbb{E} [\delta_i^2 | \text{regime's estimates}]] \tag{6.32}$$

6.7 Uncertain risk decomposition

A risk decomposition report attributes a portfolio's risk to sources, which are none other than risk parity's contributions. Applying the same machinery measures the effect of quantities' varying within their uncertainty. It surfaces the risks buried in fragile hedges.

Equation (6.8) is the convention for determining an investment's risk contribution to variance, v_i , and divided by portfolio standard deviation as in equation (6.6), standard deviation.

With uncertainty added, the single number generated from a point-estimate is replaced by the more informative $\mathbb{E}[v_i] \pm \text{stdev}[v_i]$ from equations (6.14) and (6.16).

For groups of securities—e.g., to summarize by industry, capitalization, or style—the risk contribution of a weighted basket, $\sum_i a_i v_i$, has mean $\sum_i a_i \mathbb{E}[v_i]$ and variance $\sum_{i,j} a_i a_j \text{cov}[v_i, v_j]$.

Note that with covariance modeled by factors as in section 6.8, uncertainty affects even the expected contribution, as in equation (6.41) in section 6.11.2 (repeated here):

$$\mathbb{E}[C_{i,j}] = \sum_{k=1 \dots K} \left[\hat{E}_{i,k} \hat{E}_{j,k} + \underbrace{\hat{X}_{i,k}(j, k)}_{\text{from uncertainty}} \right] \hat{\theta}_k + 1_{\{i=j\}} \hat{\varepsilon}_i^2 \hat{\theta}_{K+1} \tag{6.33}$$

This happens because the map from exposures to risk is nonlinear. Recall the earlier long-short market neutral example. The additional term is the average consequence of beta taking values in the range of uncertainty around its mean of zero, all of which involve risk (while the mean has none).

6.8 Factor-modeled covariance uncertainty

The estimates needed for uncertain risk parity—mean and covariance of the covariance matrix entries—can be generated assorted ways. For example, parametrically as in section 6.12, one could use the covariance of estimated covariance of samples assumed Gaussian⁴. Or nonparametrically, one could repeatedly sample data and estimate a matrix, then compute moments across the estimates. Instead of sampling, the same could be done across a modeled set of possible covariance matrices, perhaps assigning probabilities to differently structured versions, e.g., full covariance, diagonalized, single index, constant correlation.

A useful framework for expressing covariance—especially among investments whose composition changes over time, e.g., S&P500’s shift from industrials to financials to tech, thus not properly seen by history—is a linear factor model. To a factor model, a large universe does not present a problem. Covariance between investments that are themselves portfolios can be built up from their factor-modeled individual holdings to create forecasts based on present composition.

A conventional factor model expresses covariance as

$$\mathbf{C} = \mathbf{E}\mathbf{F}\mathbf{E}^T + \text{diag}[\boldsymbol{\varepsilon}^2] \quad (6.34)$$

where $\mathbf{E} = (N \times K)$ exposures

$\mathbf{F} = (K \times K)$ covariance between factors

$\boldsymbol{\varepsilon}^2 = (N \times 1)$ stock-specific risk

The uncertain form of Shah [2015], with parameters no longer fixed, is

$$\mathbf{C} = \mathbf{E} \left(\overbrace{\mathbf{M} \text{diag}[\theta_1, \dots, \theta_k] \mathbf{M}^T}^{\mathbf{F}} \right) \mathbf{E}^T + \theta_{K+1} \text{diag}[\boldsymbol{\varepsilon}^2] \quad (6.35)$$

where $\boldsymbol{\theta} = (K + 1 \times 1)$

$$\mathbf{M} = (K \times K) \quad (6.36)$$

As in that paper, $\mathbf{E}$, $\mathbf{M}$, $\boldsymbol{\theta}$, and $\boldsymbol{\varepsilon}^2$ are assumed jointly independent, and $\mathbf{E}$ redefined to subsume $\mathbf{M}$: $\mathbf{E} \leftarrow \mathbf{E}\mathbf{M}$.

⁴For the usual unbiased estimates of covariance made from T samples of an i.i.d. multivariate Gaussian $\mathbf{x}$, $\text{cov}[\hat{\sigma}_{x_1x_2}, \hat{\sigma}_{x_3x_4}] = (\sigma_{x_1x_3} \sigma_{x_2x_4} + \sigma_{x_1x_4} \sigma_{x_2x_3}) / (T - 1)$

$$\text{Then } \mathbf{C} = \mathbf{E} \text{diag} [\theta_1 \dots \theta_k] \mathbf{E}^T + \theta_{K+1} \text{diag} [\boldsymbol{\varepsilon}^2] \quad (6.37)$$

$$\text{and } C_{i,j} = \sum_{k=1 \dots K} E_{i,k} E_{j,k} \theta_k + 1_{\{i=j\}} \varepsilon_i^2 \theta_{K+1} \quad (6.38)$$

Given estimates of the mean and covariance of the pieces $(\mathbf{E}, \theta, \boldsymbol{\varepsilon}^2)$, equations (6.46) and (6.47) in section 6.11 yield $E[C_{p_1, p_2}]$ and $\text{cov}[C_{p_1, p_2}, C_{p_3, p_4}]$ for any portfolios p_1, p_2, p_3, p_4 .

6.9 An example

Python code for this example is on the author's github:

<https://github.com/AnishRSIGM/UncertainRiskParity>.

Conventional and uncertain risk-parity portfolios are constructed from 4 indices proxied by ETFs: AGG-iShares Core US Aggregate Bond, EEM-iShares MSCI Emerging Markets, EFA-iShares MSCI EAFE, SPY-S&P500. From 160 weeks of Wed–Wed weekly returns, 2017 Aug 23 – 2020 Sept 16, 75% (120 periods) are used for estimation (training set), and 25% (40 periods) for out-of-sample performance (testing set).

The data are examined split two ways. Experiment 1 cuts past and future—the first 120 weeks are for training, the last 40 testing. Experiment 2 separates observations within the interval—the first 3 of every 4 weeks are for training, the 4th for testing. The experiments are unusually different because markets became more volatile and tightly correlated after the pandemic hit in late February 2020. Standard deviation and correlation of weekly returns are in tables 6.1 and 6.2. The difference between training and test volatility for AGG shows bonds experienced periods of extreme returns, which fall in the test set of experiment 1 and training set of experiment 2.

Portfolios are constructed from parameters estimated on the training set and evaluated on the testing set. Expected covariance is taken as the sample covariance over the training set. Covariance of the covariance is estimated ad-hoc by cutting the 120 training periods every 10 periods, calculating the covariance over each group, and finally, the covariance across the matrices. The results are in tables 6.3 and 6.4.

Table 6.1: Experiment 1: Std Dev & Correlation
Training – 8/23/2017–12/11/2019

	Weekly SD	AGG	EEM	EFA	SPY
AGG	0.41	1.00	-0.19	-0.19	-0.19
EEM	2.52	-0.19	1.00	0.89	0.78
EFA	1.73	-0.19	0.89	1.00	0.87
SPY	1.83	-0.19	0.78	0.87	1.00

Testing – 12/11/2019–9/16/2020

	Weekly SD	AGG	EEM	EFA	SPY
AGG	1.75	1.00	0.69	0.74	0.53
EEM	4.82	0.69	1.00	0.93	0.84
EFA	4.52	0.74	0.93	1.00	0.89
SPY	4.18	0.53	0.84	0.89	1.00

Notice in uncertain risk parity, the training variance contribution has a mean and standard deviation instead of taking a single value.

Results reflect the historic events. In experiment 1, the pandemic is entirely in the testing set, and predicted contributions and their standard deviation are low. In experiment 2, much of the pandemic is in the training set, and predicted contributions and, more so, their standard deviations are high.

The example is a simplified illustration. In practice, forecasts go through complex procedures of conditioning, forecasting, and points of manual intervention. Using the linked Python code, readers can construct both conventional and uncertain risk parity portfolios from their conditioned forecasts and compare results.

6.10 Summary

Risk parity is a standard method of portfolio construction. Because it is an optimization and optimized alignments can conceal latent risks, there is a strong conceptual argument for working from more than point-estimates of covariance. Considering first and second moments, uncertain risk parity is solved using a loss function that weighs the variation around estimates. A secondary benefit of the machinery is a robust view of portfolio risk, uncertain risk decomposition, that applies to all portfolios regardless

Table 6.2: Experiment 2: Std Dev & Correlation
 Training – 8/23/2017–9/16/2019 excluding every 4th period

	Weekly SD	AGG	EEM	EFA	SPY
AGG	1.06	1.00	0.48	0.59	0.38
EEM	3.44	0.48	1.00	0.91	0.82
EFA	2.89	0.59	0.91	1.00	0.88
SPY	2.80	0.38	0.82	0.88	1.00

Testing – 8/23/2017–9/16/2019 every 4th period

	Weekly SD	AGG	EEM	EFA	SPY
AGG	0.39	1.00	-0.07	-0.14	-0.05
EEM	2.53	-0.07	1.00	0.86	0.79
EFA	1.96	-0.14	0.86	1.00	0.91
SPY	1.91	-0.05	0.79	0.91	1.00

Table 6.3: Experiment 1: Weight and Variance Contributions

Training – 8/23/2017–12/11/2019

Testing – 12/11/2019–9/16/2020

Conventional Risk Parity

	AGG	EEM	EFA	SPY
weight	0.75	0.07	0.09	0.09
v training	0.06	0.06	0.06	0.06
v testing	2.67	0.62	0.86	0.66

Uncertain Risk Parity

	AGG	EEM	EFA	SPY
weight	0.79	0.06	0.08	0.07
$E[v]$ training	0.08	0.05	0.05	0.04
$SD[v]$ training	0.08	0.67	0.54	0.61
v testing	2.77	0.52	0.72	0.45

of construction. To accommodate the changes in composition that occur over time, a factor model, in uncertain form, can be enlisted to model uncertain covariance between portfolios.

Table 6.4: Experiment 2: Weight and Variance Contributions

Training – 8/23/2017–9/16/2019 ex every 4th period

Testing – 8/23/2017–9/16/2019 every 4th period

Conventional Risk Parity				
	AGG	EEM	EFA	SPY
weight	0.54	0.13	0.15	0.17
v training	0.71	0.71	0.71	0.71
v testing	0.03	0.29	0.26	0.28

Uncertain Risk Parity				
	AGG	EEM	EFA	SPY
weight	0.44	0.16	0.17	0.22
$E[v]$ training	0.60	0.99	0.92	1.04
$SD[v]$ training	4.75	11.31	10.83	8.27
v testing	0.01	0.43	0.36	0.44

6.11 Uncertain covariance between portfolios from securities

Mean and covariance of portfolio covariance can be computed from holdings and the mean and covariance of security covariance.

Let $\mathbf{W} = (N \times M)$ holdings of N portfolios

$\mathbf{H} = (M \times M)$ covariance between M securities

$\mathbf{C} = (N \times N)$ covariance between N portfolios

6.11.1 Model-free

Without assuming an underlying model, computations are sums of 1st and 2nd moments.

$$\mathbb{E}[C_{i,j}] = \sum_{a,b} W_{i,a} W_{j,b} \mathbb{E}[H_{a,b}] \quad (6.39)$$

$$\begin{aligned} \text{cov}[C_{i,j}, C_{m,n}] &= \text{cov} \left[\sum_{a,b} W_{i,a} W_{j,b} H_{a,b}, \sum_{c,d} W_{m,c} W_{n,d} H_{c,d} \right] \\ &= \sum_{a,b,c,d} W_{i,a} W_{j,b} W_{m,c} W_{n,d} \text{cov}[H_{a,b}, H_{c,d}] \end{aligned} \quad (6.40)$$

There is one problem—the dimension of $\text{cov}[\mathbf{H}]$ is M^4 , quickly occupying a tremendous amount of memory. A universe of 1000 stocks, accounting for symmetries, would involve over $\frac{(M^2/2)^2}{2} = 125$ billion numbers.

6.11.2 Uncertain covariance modeled

Space complexity issues can be avoided by imposing structure. The uncertain covariance model of [Shah \[2015\]](#) is as follows:

1. $\mathbf{H} = \mathbf{E} \text{diag}[\theta_1 \dots \theta_K] \mathbf{E}^T + \theta_{K+1} \text{diag}[\boldsymbol{\varepsilon}^2]$
2. $\mathbf{E}$, θ , and $\boldsymbol{\varepsilon}^2$ are jointly independent
3. The elements of $\boldsymbol{\varepsilon}^2$ are independent
4. $\mathbf{E}$ is Gaussian
5. Mean and covariance of $\mathbf{E}$, θ , and $\boldsymbol{\varepsilon}^2$ are

$\hat{\mathbf{E}} = (M \times K)$	$\hat{\mathbf{E}} = \mathbb{E}[\mathbf{E}]$
$\hat{\mathbf{X}} = (M \times K \times M \times K)$	$\hat{\mathbf{X}}_{a,i,b,j} = \text{cov}[\mathbf{E}_{a,i}, \mathbf{E}_{b,j}]$
$\hat{\boldsymbol{\varepsilon}}^2 = (M \times 1)$	$\hat{\boldsymbol{\varepsilon}}^2 = \mathbb{E}[\boldsymbol{\varepsilon}^2]$
$\hat{\boldsymbol{\omega}} = (M \times 1)$	$\hat{\omega}_a = \text{var}[\varepsilon_a^2]$
$\hat{\boldsymbol{\theta}} = (K + 1 \times 1)$	$\hat{\boldsymbol{\theta}} = \mathbb{E}[\boldsymbol{\theta}]$
$\hat{\boldsymbol{\Omega}} = (K + 1 \times K + 1)$	$\hat{\Omega}_{i,j} = \text{cov}[\theta_i, \theta_j]$

Then

$$\mathbb{E}[H_{a,b}] = \sum_{k=1 \dots K} \left[\hat{E}_{a,k} \hat{E}_{b,k} + \hat{X}_{a,k,b,k} \right] \hat{\theta}_k + 1_{\{a=b\}} \hat{\varepsilon}_a^2 \hat{\theta}_{K+1} \quad (6.41)$$

$$\mathbb{E}[\mathbf{H}] = \hat{\mathbf{E}} \text{diag} \left[\hat{\theta}_1 \dots \hat{\theta}_K \right] \hat{\mathbf{E}}^T + \sum_{k=1 \dots K} \hat{\theta}_k \hat{\mathbf{X}}_{\cdot,k,\cdot,k} + \hat{\theta}_{K+1} \text{diag} \left[\hat{\varepsilon}^2 \right] \quad (6.42)$$

$$\begin{aligned} \text{cov}[H_{a,b}, H_{c,d}] = & \sum_{1 \leq k, l \leq K} \left[\hat{E}_{a,k} \hat{E}_{c,l} \hat{X}_{b,k,d,l} + \hat{E}_{b,k} \hat{E}_{d,l} \hat{X}_{a,k,c,l} \right. \\ & + \hat{X}_{a,k,c,l} \hat{X}_{b,k,d,l} + \hat{E}_{a,k} \hat{E}_{d,l} \hat{X}_{b,k,c,l} \\ & \left. + \hat{E}_{b,k} \hat{E}_{c,l} \hat{X}_{a,k,d,l} + \hat{X}_{a,k,d,l} \hat{X}_{b,k,c,l} \right] \left(\hat{\theta}_k \hat{\theta}_l + \hat{\Omega}_{k,l} \right) \\ & + \sum_{1 \leq k, l \leq K} \left[\hat{E}_{a,k} \hat{E}_{b,k} \hat{E}_{c,l} \hat{E}_{d,l} + \hat{E}_{a,k} \hat{E}_{b,k} \hat{X}_{c,l,d,l} \right. \\ & \left. + \hat{E}_{c,l} \hat{E}_{d,l} \hat{X}_{a,k,b,k} + \hat{X}_{a,k,b,k} \hat{X}_{c,l,d,l} \right] \hat{\Omega}_{k,l} \\ & + 1_{\{a=b\}} \hat{\varepsilon}_a^2 \sum_{k=1 \dots K} \left[\hat{E}_{c,k} \hat{E}_{d,k} + \hat{X}_{c,k,d,k} \right] \hat{\Omega}_{k,K+1} \\ & + 1_{\{c=d\}} \hat{\varepsilon}_c^2 \sum_{k=1 \dots K} \left[\hat{E}_{a,k} \hat{E}_{b,k} + \hat{X}_{a,k,b,k} \right] \hat{\Omega}_{k,K+1} \\ & + 1_{\{a=b\}} 1_{\{c=d\}} \hat{\varepsilon}_a^2 \hat{\varepsilon}_c^2 \hat{\Omega}_{K+1,K+1} \\ & + 1_{\{a=b=c=d\}} \hat{\omega}_a \left(\hat{\theta}_{K+1}^2 + \hat{\Omega}_{K+1,K+1} \right) \end{aligned} \quad (6.43)$$

Using $\mathbf{C} = \mathbf{W}\mathbf{H}\mathbf{W}^T$ with this structured form and applying formulas (6.39) and (6.40) to equations (6.42) and (6.43) yields the expected value and covariance of

portfolio covariance.

Let $\mathbf{G} = (N \times K)$ factor exposures of portfolios

$$\begin{aligned}\hat{\mathbf{G}} &= \mathbf{E}[\mathbf{G}] \\ &= \mathbf{W}\hat{\mathbf{E}}\end{aligned}\tag{6.44}$$

$$\hat{\mathbf{Y}} = (N \times K \times N \times K) \text{ cov}[\mathbf{G}]$$

$$\begin{aligned}\hat{\mathbf{Y}}_{m,i,n,j} &= \text{cov}[\mathbf{G}_{m,i}, \mathbf{G}_{n,j}] \\ \hat{\mathbf{Y}}_{\cdot,i,\cdot,j} &= \mathbf{W}\hat{\mathbf{X}}_{\cdot,i,\cdot,j}\mathbf{W}^T\end{aligned}\tag{6.45}$$

$$\mathbf{E}[\mathbf{C}] = \hat{\mathbf{G}} \text{diag}[\hat{\theta}_1 \dots \hat{\theta}_K] \hat{\mathbf{G}}^T + \sum_{k=1 \dots K} \hat{\theta}_k \hat{\mathbf{Y}}_{\cdot,k,\cdot,k} + \hat{\theta}_{K+1} \mathbf{W} \text{diag}[\hat{\varepsilon}^2] \mathbf{W}^T \tag{6.46}$$

$$\begin{aligned}\text{cov}[C_{i,j}, C_{m,n}] &= \sum_{1 \leq k, l \leq K} \left[\hat{G}_{i,k} \hat{G}_{m,l} \hat{Y}_{j,k,n,l} + \hat{G}_{j,k} \hat{G}_{n,l} \hat{Y}_{i,k,m,l} \right. \\ &\quad \left. + \hat{Y}_{i,k,m,l} \hat{Y}_{j,k,n,l} + \hat{G}_{i,k} \hat{G}_{n,l} \hat{Y}_{j,k,m,l} \right. \\ &\quad \left. + \hat{G}_{j,k} \hat{G}_{m,l} \hat{Y}_{i,k,n,l} + \hat{Y}_{i,k,n,l} \hat{Y}_{j,k,m,l} \right] \left(\hat{\theta}_k \hat{\theta}_l + \hat{\Omega}_{k,l} \right) \\ &+ \sum_{1 \leq k, l \leq K} \left[\hat{G}_{i,k} \hat{G}_{j,k} \hat{G}_{m,l} \hat{G}_{n,l} + \hat{G}_{i,k} \hat{G}_{j,k} \hat{Y}_{m,l,n,l} \right. \\ &\quad \left. + \hat{G}_{m,l} \hat{G}_{n,l} \hat{Y}_{i,k,j,k} + \hat{Y}_{i,k,j,k} \hat{Y}_{m,l,n,l} \right] \hat{\Omega}_{k,l} \\ &+ \left(\sum_a W_{i,a} W_{j,a} \hat{\varepsilon}_a^2 \right) \left(\sum_{k=1 \dots K} \left[\hat{G}_{m,k} \hat{G}_{n,k} + \hat{Y}_{m,k,n,k} \right] \hat{\Omega}_{k,K+1} \right) \\ &+ \left(\sum_a W_{m,a} W_{n,a} \hat{\varepsilon}_a^2 \right) \left(\sum_{k=1 \dots K} \left[\hat{G}_{i,k} \hat{G}_{j,k} + \hat{Y}_{i,k,j,k} \right] \hat{\Omega}_{k,K+1} \right) \\ &+ \left(\sum_a W_{i,a} W_{j,a} \hat{\varepsilon}_a^2 \right) \left(\sum_a W_{m,a} W_{n,a} \hat{\varepsilon}_a^2 \right) \hat{\Omega}_{K+1,K+1} \\ &+ \left(\sum_a W_{i,a} W_{j,a} W_{m,a} W_{n,a} \hat{\omega}_a \right) \left(\hat{\theta}_{K+1}^2 + \hat{\Omega}_{K+1,K+1} \right)\end{aligned}\tag{6.47}$$

6.12 Parametric uncertainty covariance of unbiased covariance estimates

Lemma 6.1. *Parametric estimate of covariance uncertainty*

If $\hat{\sigma}_{ab}$ and $\hat{\sigma}_{cd}$ are the usual unbiased estimates of covariance made from N samples

of an i.i.d. multivariate Gaussian, $\left\{ \begin{pmatrix} a_i \\ b_i \\ c_i \\ d_i \end{pmatrix} \right\}_{i=1..N}$,

then $\text{cov} [\hat{\sigma}_{ab}, \hat{\sigma}_{cd}] = (\sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}) / (N - 1)$

$$\begin{aligned} \text{where } \hat{\sigma}_{ab} &= \frac{1}{N-1} \sum_i \left[a_i - \frac{1}{N} \sum_k a_k \right] \left[b_i - \frac{1}{N} \sum_l b_l \right] \\ &= \frac{1}{N-1} \left[\sum_i a_i b_i - \frac{1}{N} \sum_k a_k \sum_l b_l \right] \end{aligned}$$

Proof. Since $E[a_i] = E[\frac{1}{N} \sum_k a_k]$, the means are irrelevant so can be assumed to be 0.

$$\text{cov} [\hat{\sigma}_{ab}, \hat{\sigma}_{cd}] = E [\hat{\sigma}_{ab} \hat{\sigma}_{cd}] - E [\hat{\sigma}_{ab}] E [\hat{\sigma}_{cd}] \quad (6.48)$$

$$= E [\hat{\sigma}_{ab} \hat{\sigma}_{cd}] - \sigma_{ab} \sigma_{cd} \quad (6.49)$$

$$(N-1)^2 E [\hat{\sigma}_{ab} \hat{\sigma}_{cd}] = E \left[\left(\sum_i a_i b_i - \frac{1}{N} \sum_k a_k \sum_l b_l \right) \left(\sum_j c_j d_j - \frac{1}{N} \sum_m c_m \sum_n d_n \right) \right] \quad (6.50)$$

$$\begin{aligned} &= E \left[\underbrace{\sum_i a_i b_i \sum_j c_j d_j}_{\textcircled{1}} - \frac{1}{N} \underbrace{\left[\sum_i a_i b_i \sum_m c_m \sum_n d_n \right]}_{\textcircled{2}} \right] \\ &\quad - \frac{1}{N} E \left[\underbrace{\sum_j c_j d_j \sum_k a_k \sum_l b_l}_{\textcircled{3}} \right] \\ &\quad + \frac{1}{N^2} E \left[\underbrace{\sum_k a_k \sum_l b_l \sum_m c_m \sum_n d_n}_{\textcircled{4}} \right] \quad (6.51) \end{aligned}$$

$$\textcircled{2} = \text{E} \left[\sum_{i,m,n} a_i b_i c_m d_n \right] \quad (6.52)$$

$$= \text{E} \left[\underbrace{\sum_{i,m,n \neq m} a_i b_i c_m d_n}_{=0} \right] + \text{E} \left[\sum_{i,m} a_i b_i c_m d_m \right] \quad (6.53)$$

since first and third moments are zero

$$= \text{E} \left[\sum_i a_i b_i \sum_m c_m d_m \right] \quad (6.54)$$

$$= \textcircled{1} \quad (6.55)$$

$$\text{Likewise, } \textcircled{3} = \textcircled{1} \quad (6.56)$$

$$\textcircled{1} = \text{E} \left[\sum_{i,j} a_i b_i c_j d_j \right] \quad (6.57)$$

$$= \text{E} \left[\sum_{i,j \neq i} a_i b_i c_j d_j + \sum_i a_i b_i c_i d_i \right] \quad (6.58)$$

$$= \sum_{i,j \neq i} \text{E} [a_i b_i] \text{E} [c_j d_j] + \sum_i \text{E} [a_i b_i c_i d_i] \quad (6.59)$$

$$= \sum_{i,j \neq i} \sigma_{a,b} \sigma_{c,d} + \sum_i \text{E} [a_i b_i c_i d_i] \quad (6.60)$$

$$= N(N-1) \sigma_{a,b} \sigma_{c,d} + \sum_i \text{E} [a_i b_i c_i d_i] \quad (6.61)$$

$$= N(N-1) \sigma_{a,b} \sigma_{c,d} + \sum_i \sigma_{a,b} \sigma_{c,d} + \sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c} \quad (6.62)$$

by Isserlis' theorem [Isserlis \[1918\]](#)

$$= N^2 \sigma_{a,b} \sigma_{c,d} + N \sigma_{a,c} \sigma_{b,d} + N \sigma_{a,d} \sigma_{b,c} \quad (6.63)$$

$$\textcircled{4} = \text{E} \left[\sum_{k,l,m,n} a_k b_l c_m d_n \right] \quad (6.64)$$

$$= \text{E} \left[\sum_{k,m \neq k} a_k b_k c_m d_m \right] + \text{E} \left[\sum_{k,m \neq k} a_k b_m c_k d_m \right] + \text{E} \left[\sum_{k,m \neq k} a_k b_m c_m d_k \right] \\ + \text{E} \left[\sum_k a_k b_k c_k d_k \right] \quad \text{since first and third moments are zero} \quad (6.65)$$

$$= N(N-1) (\sigma_{a,b} \sigma_{c,d} + \sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}) \\ + \sum_k \sigma_{a,b} \sigma_{c,d} + \sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c} \quad (6.66)$$

$$= N^2 (\sigma_{a,b} \sigma_{c,d} + \sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}) \quad (6.67)$$

Substituting back into equations (6.49) and (6.51)

$$(N-1)^2 \text{E} [\hat{\sigma}_{ab} \hat{\sigma}_{cd}] = \frac{N-2}{N} (N^2 \sigma_{a,b} \sigma_{c,d} + N \sigma_{a,c} \sigma_{b,d} + N \sigma_{a,d} \sigma_{b,c}) \\ + (\sigma_{a,b} \sigma_{c,d} + \sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}) \quad (6.68)$$

$$= (N-1)^2 \sigma_{a,b} \sigma_{c,d} + (N-1) (\sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}) \quad (6.69)$$

$$\text{cov} [\hat{\sigma}_{ab}, \hat{\sigma}_{cd}] = \text{E} [\hat{\sigma}_{ab} \hat{\sigma}_{cd}] - \sigma_{ab} \sigma_{cd} \\ = \frac{\sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}}{N-1} \quad (6.70)$$

□

CHAPTER 7

RISK UNDER UNCERTAINTY AND PRICE MOVEMENT

Risk decomposition separates a portfolio’s variance or standard deviation into sources. It is a standard tool for analyzing the risks of an investment portfolio. The portfolio is divided into notional parts—e.g., individual securities, holdings by sector or region, factor exposures—whose contributions to net risk are estimated and reported. The procedure is described in [Menchero and Davis \[2011\]](#) and [O’Cinneide \[2018\]](#).

For risk measured in variance, a part’s contribution is its covariance with whole; in standard deviation, this number is divided by portfolio standard deviation. For example, for the simple case of measuring security contributions to portfolio risk against a risk-free benchmark,

$$\begin{aligned} v &\equiv \text{security contributions to portfolio variance} \in \mathbb{R}^n \\ &\equiv w \odot \Sigma w \end{aligned} \tag{7.1}$$

$$\begin{aligned} s &\equiv \text{security contributions to portfolio std dev} \in \mathbb{R}^n \\ &\equiv v / \sqrt{w^T \Sigma w} \end{aligned} \tag{7.2}$$

where $\odot \equiv$ element-by-element multiplication

Convention regards the inputs—portfolio weights and covariance—as fixed and known, but composition changes with price and estimates have errors. In equation (7.1), neither w nor Σ is actually fixed.

Section 6 of [O’Cinneide \[2018\]](#) considers perturbations of Σ . The machinery developed in chapter 5 of this dissertation makes it possible to deform Σ according to beliefs and compute $E[\Sigma]$ and $\text{cov}[\Sigma]$. The paper underlying this chapter, [Shah \[2020\]](#), additionally lets security weights vary and, assuming weight movement independent of parameter uncertainty, computes risk contributions’ mean and variance.

Evaluating across the range of possibility measures risks more accurately and

surfaces latent fragility in hedging. Presentations of the paper have been well-received by a broad audience. Soon after giving an online talk, I was invited to present to professional/academic associations in New York¹, Amsterdam², and London³.

7.1 Mathematical model

As in the two preceding chapters, covariance is modeled via factors, this time not orthogonal,

$$\begin{aligned}\Sigma &\equiv \text{cov} [Y|E, F, \gamma, \varepsilon^2] \\ &\equiv EFE^T + \gamma \text{diag}(\varepsilon^2)\end{aligned}\tag{7.3}$$

In the mathematical framework of chapter 3, we are interested in the center and spread of functions $f(\Sigma_t, w_t) | X_t$ where $Z_t = \{E_t, F_t, \gamma_t, \varepsilon_t^2, w_t\}$ and the side information $X_t = \{\text{mean and variance of } Z_t\}$ encapsulates all the information of $Y_{1:t-1}, X_{1:t-1}, Z_{1:t-1}$.

¹Shah [2021c]

²Shah [2021b]

³Shah [2022]

from [Shah \[2020\]](#) Risk under Uncertainty and Price Movement

7.2 Abstract

Risk decomposition is a standard tool for analyzing investment portfolio risk. The portfolio is divided into notional parts—e.g., individual securities, holdings by sector or region, factor exposures—whose contributions to net risk are estimated and reported. Convention regards the inputs—portfolio weights and covariance—as fixed and known, but portfolio composition changes with price movement and estimates have errors. Since behavior only in the direction of net risk is counted, hedged effects are invisible regardless of size. What if numbers aren’t exact? Hedge instability manifests in proportion to underlying gross (not net) exposure, akin to leverage. For example, in a market-neutral portfolio, market is the largest risk in each side, conventionally uncounted since hedged, and extremely consequential under small deviations. To solve the problem, this paper models weights and parameters as uncertain. Evaluating a portfolio across the range of possibility measures risks better and surfaces latent fragility. No longer point-estimated, contributions are reported with a center and spread.

7.3 Introduction

Risk decomposition is a tool used by portfolio managers, risk managers, allocators, and others to examine a portfolio’s risk profile. It breaks net risk into contributions from the portfolio’s notional parts, which are defined according to the perspective sought. Parts can be concrete—e.g., individual securities; holdings grouped by sector, geographical region, or separately managed subportfolio—or abstract, e.g., exposures to risk factors and idiosyncratic effects.

Knowledge of how parts contribute to or hedge portfolio volatility guides decisions throughout the investment process. A portfolio manager can see over/undersized and hidden bets, for example, that an investment tilt adds less risk than thought so can

be increased. Risk managers and allocators use the information to keep portfolios on mandate, analyze changes, and compute metrics, e.g., parametric VaR contributions. The information enters performance measurement in attribution computations. Client management and sales teams communicate and certify the nature of portfolios through risk contribution reports produced by commercial investment platforms. Beyond reporting, risk contributions are the quantities equalized in risk parity portfolio construction.

7.3.1 Problem with conventional method

Numbers are computed from portfolio weights and security covariance, both regarded as fixed and known. However, the composition of a live portfolio shifts continuously with price movement, and inferred quantities are imperfect.

While many tasks rely on point estimates, the problem in this setting is articulated in [Shah \[2021a\]](#), which investigates risk parity under parameter uncertainty. The argument, expanded to include changes in composition, is as follows.

[Menchero and Davis \[2011\]](#), in giving a geometric interpretation of risk contributions, shows that behavior *in the direction of net risk* is what’s counted. Behavior in all other directions cancels out by definition. At the same time, the gross exposure of these uncounted hedged behaviors can be large—picture the market effect contained in each side of a market-neutral⁴. If the system experiences any small deviation—through price movement or gap between estimated parameter and reality—the hedging balance is thrown off, and the risk of the gross exposure manifests.

The severe aspect of the problem is not that hedged behaviors inaccurately sum out, e.g., the short side’s market effect has a negative risk contribution which seems to offset the long’s positive. It’s that they go unperceived, i.e., market risk scarcely appears anywhere in the portfolio because market is nearly orthogonal to the net thus invisible. These alignments occur because portfolios are constructed to hedge risks.

⁴Or in benchmark-relative risks, the benchmark acts as the short side.

7.3.2 Contribution and central idea of this paper

This paper surfaces latent risks by considering inputs across the range of uncertainty. Parameters and future portfolio weights are mathematically perturbed according to one's beliefs about their distribution. The impact of gross exposure occurs and is measured. Structural assumptions allow computation analytically, without sampling. Risk contributions, no longer derived point estimates, are accompanied with an estimate of their uncertainty.

[O'Kinneide \[2018\]](#) looks at perturbations from parameter uncertainty to risk contributions. There is also a paragraph on variance contributions under parameter uncertainty in the aforementioned uncertain risk parity paper. The solution presented here leverages the uncertain covariance technique used by that paper, from [Shah \[2015\]](#). The author is aware of no other literature that considers both uncertainty and price movement.

7.4 Conventional risk contributions

Risk contributions are the portion of volatility, in either standard deviation or variance, attributed to the separate parts of the portfolio. They sum to portfolio volatility. Absent uncertainty and price movement, calculations are standard.⁵

In variance:

$$\begin{aligned}
 \sigma^2 &= \text{variance of portfolio} \\
 v_i &= \text{part } i\text{'s contribution to portfolio variance} \\
 &= \text{cov}[\text{part } i, \text{portfolio}] \\
 \sum_i v_i &= \sigma^2
 \end{aligned} \tag{7.4}$$

⁵For a derivation, see section 3.2.3 of [Qian et al. \[2007\]](#)

In standard deviation:

$$\begin{aligned}
 s_i &= \text{part } i\text{'s contribution to portfolio standard deviation} \\
 &= v_i/\sigma \\
 \sum_i s_i &= \sigma
 \end{aligned} \tag{7.5}$$

The contributed fraction of volatility is the same either way, $v_i/\sigma^2 = s_i/\sigma$.

7.4.1 Variance contributions of common partitions

Variance contributions of common portfolio partitions are stated here because they will be later modeled as uncertain. The corresponding standard deviation contribution comes from applying equation (7.5).

Individual securities, groups, or subportfolios

Consider the portfolio as a collection of subportfolios, $\mathbf{w} = \mathbf{p}^{(1)} + \mathbf{p}^{(2)} + \dots$, e.g., by individual security, group (sector, geographic region, investment type), or separately managed holdings. Subportfolio p 's variance contribution is

$$v_p = \mathbf{p}^T \mathbf{C} \mathbf{w} \tag{7.6}$$

where $\mathbf{w} = (N \times 1)$ portfolio security weights

$\mathbf{C} = (N \times N)$ covariance

If subportfolios are individual securities, the n 'th security contributes

$$v_n = w_n \sum_m C_{n,m} w_m \tag{7.7}$$

Factors

If covariance is factor-modeled as

$$\mathbf{C} = \mathbf{E}\mathbf{F}\mathbf{E}^T + \gamma \text{diag} [\boldsymbol{\varepsilon}^2] \quad (7.8)$$

where $\mathbf{E} = (N \times K)$ factor exposures

$\mathbf{F} = (K \times K)$ covariance between factors

$\boldsymbol{\varepsilon}^2 = (N \times 1)$ idiosyncratic variance

$\gamma = 1$ for now. Will be used later

let $\mathbf{e} = (K \times 1)$ exposures of the portfolio

$$= \mathbf{E}^T \mathbf{w} \quad (7.9)$$

$\mathbf{e}^p = (K \times 1)$ exposures from subportfolio p

$$= \mathbf{E}^T \mathbf{p} \quad (7.10)$$

Variance contributions are

by factor

$$\text{factor } i : \quad v_i = e_i (\mathbf{F}\mathbf{e})_i \quad (7.11)$$

$$\varepsilon^2 : \quad v_{\varepsilon^2} = \gamma \sum_m w_m^2 \varepsilon_m^2 \quad (7.12)$$

$$\text{total: } \sigma^2 = \sum_i v_i + v_{\varepsilon^2} \quad (7.13)$$

by factor and subportfolio

$$\text{subportfolio } p, \text{ factor } i : \quad v_{p,i} = e_i^p (\mathbf{F}\mathbf{e})_i \quad (7.14)$$

$$\text{subportfolio } p, \varepsilon^2 : \quad v_{p,\varepsilon^2} = \gamma \sum_m p_m w_m \varepsilon_m^2 \quad (7.15)$$

$$\text{total: } \sigma^2 = \sum_{k,i} v_{p^{(k)},i} + \sum_k v_{p^{(k)},\varepsilon^2} \quad (7.16)$$

7.5 Uncertain variance contributions

To add uncertainty—in $\mathbf{C}$ from estimation, in $\mathbf{w}$ from price movement—the point estimate and fixed weights are replaced by estimates of mean and variance⁶.

$$\mathbb{E}[\mathbf{C}] = (N \times N) \text{ expected value of } \mathbf{C}$$

$$\text{var}[\mathbf{C}] = (N \times N \times N \times N) \text{ covariance between elements of } \mathbf{C}$$

$$\mathbb{E}[\mathbf{w}] = (N \times 1) \text{ expected value of } \mathbf{w}$$

$$\text{var}[\mathbf{w}] = (N \times N) \text{ covariance between elements of } \mathbf{w}$$

Consider the individual security contributions of equation (7.7). With parameters and weights independent, mean and covariance are

$$\begin{aligned} \mathbb{E}[v_{\text{security } n}] &= \mathbb{E}\left[\sum_m w_n C_{n,m} w_m\right] \\ &= \sum_m \mathbb{E}[C_{n,m}] \mathbb{E}[w_n w_m] \\ &= \sum_m \mathbb{E}[C_{n,m}] (\mathbb{E}[w_n] \mathbb{E}[w_m] + \text{cov}[w_n, w_m]) \end{aligned} \quad (7.17)$$

$$\begin{aligned} \text{cov}[v_n, v_m] &= \text{cov}\left[\sum_i w_n C_{n,i} w_i, \sum_j w_m C_{m,j} w_j\right] \\ &= \sum_{i,j} \text{cov}[w_n C_{n,i} w_i, w_m C_{m,j} w_j] \\ &= \sum_{i,j} \mathbb{E}[w_n w_i] \mathbb{E}[w_m w_j] \text{cov}[C_{n,i}, C_{m,j}] \\ &\quad + \text{cov}[w_n w_i, w_m w_j] \mathbb{E}[C_{n,i}] \mathbb{E}[C_{m,j}] \\ &\quad + \text{cov}[w_n w_i, w_m w_j] \text{cov}[C_{n,i}, C_{m,j}] \end{aligned} \quad (7.18)$$

since $\mathbf{w}$, $\mathbf{C}$ independent

Though equation (7.18) is straightforward, dimensionality in $\text{var}[\mathbf{C}]$ and computational complexity can be a problem. As without uncertainty, both are reduced in the structure of a factor model.

⁶In this paper's notation, $\text{var}[\mathbf{x}]$ means covariance of (scalar, vector, or matrix) variable $\mathbf{x}$ with itself and $\text{cov}[\mathbf{x}, \mathbf{y}]$ is covariance between $\mathbf{x}$ and $\mathbf{y}$.

7.5.1 Uncertain factor-modeled covariance

Changing the parameters in the factor model of equation (7.8) from fixed to having mean and covariance yields an uncertain covariance model similar to Shah [2015] but more general⁷ because no structure is imposed on $\mathbf{F}$.

$$\mathbf{C} = \mathbf{E}\mathbf{F}\mathbf{E}^T + \gamma \text{diag} [\boldsymbol{\varepsilon}^2] \quad (7.19)$$

Three assumptions make computations possible.

1. $\mathbf{E}, \boldsymbol{\varepsilon}^2, \{\mathbf{F}, \gamma\}$, and $\mathbf{w}$ are jointly independent.
2. The elements of $\boldsymbol{\varepsilon}^2$ are independent.
3. $\mathbf{w}$ and factor exposures (with the effect of shifting composition) of the portfolio and subportfolios are Gaussian.

The model is defined by mean and covariance of the parameters, and the portfolio by mean and covariance of its weights.

Exposures

$$\begin{aligned} \bar{\mathbf{E}} &= (N \times K) & \bar{\mathbf{E}} &= \mathbb{E} [\mathbf{E}] \\ \mathbf{X} &= (N \times K \times N \times K) & X_{m,i,n,j} &= \text{cov} [E_{m,i}, E_{n,j}] \end{aligned}$$

Idiosyncratic effects

$$\begin{aligned} \bar{\boldsymbol{\varepsilon}}^2 &= (N \times 1) & \bar{\boldsymbol{\varepsilon}}^2 &= \mathbb{E} [\boldsymbol{\varepsilon}^2] \\ \boldsymbol{\omega} &= (N \times 1) & \omega_n &= \text{var} [\varepsilon_n^2] \end{aligned}$$

Factor structure

$$\begin{aligned} \bar{\mathbf{F}} &= (K \times K) & \bar{\mathbf{F}} &= \mathbb{E} [\mathbf{F}] \\ \mathbf{H} &= (K \times K \times K \times K) & H_{i,j,k,l} &= \text{cov} [F_{i,j}, F_{k,l}] \\ \bar{\gamma} &= (\text{scalar}) & \bar{\gamma} &= \mathbb{E} [\gamma] \\ \eta &= (\text{scalar}) & \eta &= \text{var} [\gamma] \\ \mathbf{h} &= (K \times K) & h_{i,j} &= \text{cov} [F_{i,j}, \gamma] \end{aligned}$$

⁷Appendix 7.11 converts that uncertain covariance structure to this paper's.

Portfolio weights

$$\bar{\mathbf{w}} = (N \times 1)$$

$$\mathbf{\Omega} = (N \times N)$$

$$\bar{\mathbf{w}} = \mathbb{E}[\mathbf{w}]$$

$$\Omega_{m,n} = \text{cov}[w_m, w_n]$$

7.5.2 Mean and covariance of contributions

Contributions from subportfolio and factor in equations (7.14) and (7.15) are the building blocks of most quantities of interest, e.g., contributions by factor, region, or sector. Their means and covariances are derived in the following propositions: 7.1 (expectations $\mathbb{E}[v_{p,i}]$, $\mathbb{E}[v_{p,\varepsilon^2}]$) and 7.2, 7.3, 7.4 (covariance pairs $\text{cov}[v_{p,i}, v_{q,k}]$, $\text{cov}[v_{p,\varepsilon^2}, v_{q,\varepsilon^2}]$, $\text{cov}[v_{p,i}, v_{q,\varepsilon^2}]$).

Let P stand for the set of securities notionally in subportfolio p , even those with zero weight⁸. To avoid more notation, assume if a subportfolio holds a security, the weight is the whole portfolio's weight, i.e., $n \in P \rightarrow p_n = w_n$. Nevertheless, the formulas below allow for securities held in common so that they can be used to compute covariance with the entire portfolio.

⁸If weights are benchmark-relative, price movement in other securities shifts these away from zero.

Proposition 7.1. $\mathbb{E}[v_{p,i}]$ and $\mathbb{E}[v_{p,\varepsilon^2}]$

$$\mathbb{E}[v_{p,i}] = \overbrace{\bar{e}_i^p (\bar{\mathbf{F}} \bar{\mathbf{e}})_i}^{\text{mean}} + \overbrace{\sum_j \bar{F}_{i,j} \text{cov}[e_i^p, e_j]}^{\text{uncertainty}} \quad (7.20)$$

$$\begin{aligned} &= \overbrace{\bar{e}_i^p (\bar{\mathbf{F}} \bar{\mathbf{e}})_i}^{\text{mean}} + \\ &\quad \sum_j \bar{F}_{i,j} \sum_{n \in P, m} \overbrace{X_{n,i,m,j} \bar{w}_n \bar{w}_m}^{E \text{ uncertainty}} + \overbrace{\bar{E}_{n,i} \bar{E}_{m,j} \Omega_{n,m}}^{w \text{ uncertainty}} + \overbrace{X_{n,i,m,j} \Omega_{n,m}}^{E,w \text{ unc.}} \end{aligned} \quad (7.21)$$

$$\mathbb{E}[v_{p,\varepsilon^2}] = \bar{\gamma} \sum_{n \in P} \left[\underbrace{\bar{w}_n^2}_{\text{mean}} + \underbrace{\Omega_{n,n}}_{w \text{ unc.}} \right] \bar{\varepsilon}_n^2 \quad (7.22)$$

$$\text{where } \bar{\mathbf{e}}^p = \mathbb{E}[\mathbf{e}^p] = \bar{\mathbf{E}}^T \bar{\mathbf{p}} \quad (7.23)$$

$$\bar{\mathbf{e}} = \bar{\mathbf{E}}^T \bar{\mathbf{w}} \quad (7.24)$$

$\mathbb{E}[v_{p,i}]$ computes in $O(K)$ given $\bar{\mathbf{e}}^p, \bar{\mathbf{e}}, \text{cov}[\mathbf{e}^p, \mathbf{e}]$. $\mathbb{E}[v_{p,\varepsilon^2}]$ computes in $O(N)$, $\bar{\mathbf{e}}^p$ and $\bar{\mathbf{e}}$ in $O(KN)$, and $\text{cov}[\mathbf{e}^p, \mathbf{e}]$ in $O(K^2 N^2)$.

Proof.

$$\mathbb{E}[v_{p,i}] = \mathbb{E}[e_i^p (\mathbf{F}\mathbf{e})_i] \quad (7.25)$$

$$\begin{aligned} &= \mathbb{E} \left[e_i^p \sum_j F_{i,j} e_j \right] \\ &= \sum_j \mathbb{E}[F_{i,j}] \mathbb{E}[e_i^p e_j] \quad \text{since } \mathbf{e}, \mathbf{F} \text{ ind} \\ &= \sum_j \bar{F}_{i,j} (\mathbb{E}[e_i^p] \mathbb{E}[e_j] + \text{cov}[e_i^p, e_j]) \\ &= \underbrace{\bar{e}_i^p (\bar{\mathbf{F}}\bar{\mathbf{e}})_i}_{\text{mean}} + \underbrace{\sum_j \bar{F}_{i,j} \text{cov}[e_i^p, e_j]}_{\text{uncertainty}} \end{aligned} \quad (7.26)$$

$$\begin{aligned} \text{cov}[e_i^p, e_j^q] &= \text{cov} \left[\sum_{n \in P} E_{n,i} w_n, \sum_{m \in Q} E_{m,j} w_m \right] \\ &= \sum_{n \in P, m \in Q} \text{cov}[E_{n,i} w_n, E_{m,j} w_m] \\ &= \sum_{n \in P, m \in Q} \underbrace{X_{n,i,m,j} \bar{w}_n \bar{w}_m}_{E \text{ uncertainty}} + \underbrace{\bar{E}_{n,i} \bar{E}_{m,j} \Omega_{n,m}}_{w \text{ uncertainty}} + \underbrace{X_{n,i,m,j} \Omega_{n,m}}_{E,w \text{ unc.}} \\ &\quad \text{since } \mathbf{E}, \mathbf{w} \text{ ind} \end{aligned} \quad (7.27)$$

$$\mathbb{E}[v_{p,\varepsilon^2}] = \mathbb{E} \left[\gamma \sum_{n \in P} w_n^2 \varepsilon_n^2 \right] \quad (7.28)$$

$$= \mathbb{E}[\gamma] \sum_{n \in P} \mathbb{E}[w_n^2] \mathbb{E}[\varepsilon_n^2] \quad \text{since } \gamma, \mathbf{w}, \boldsymbol{\varepsilon}^2 \text{ ind} \quad (7.29)$$

$$= \bar{\gamma} \sum_{n \in P} \left[\underbrace{\bar{w}_n^2}_{\text{mean}} + \underbrace{\Omega_{n,n}}_{w \text{ unc.}} \right] \bar{\varepsilon}_n^2 \quad (7.30)$$

□

Notice that equations (7.21) and (7.22) separate expectation into sources—mean and different uncertainties.

Proposition 7.2. $\text{cov}[v_{p,i}, v_{q,k}]$

If portfolio and subportfolio factor exposures $\mathbf{e}$, $\mathbf{e}^p$, and $\mathbf{e}^q$ are Gaussian,

$$\begin{aligned} \text{cov}[v_{p,i}, v_{q,k}] &= \sum_{j,l} H_{i,j,k,l} (\bar{e}_i^p \bar{e}_j + \text{cov}[e_i^p, e_j]) (\bar{e}_k^q \bar{e}_l + \text{cov}[e_k^q, e_l]) \\ &\quad + (\bar{F}_{i,j} \bar{F}_{k,l} + H_{i,j,k,l}) \left[\right. \\ &\quad (\bar{e}_i^p \bar{e}_k^q + \text{cov}[e_i^p, e_k^q]) (\bar{e}_j \bar{e}_l + \text{cov}[e_j, e_l]) \\ &\quad + (\bar{e}_i^p \bar{e}_l + \text{cov}[e_i^p, e_l]) (\bar{e}_k^q \bar{e}_j + \text{cov}[e_k^q, e_j]) \\ &\quad \left. - 2\bar{e}_i^p \bar{e}_k^q \bar{e}_j \bar{e}_l \right] \end{aligned} \quad (7.31)$$

and is computed in $O(K^2)$ given $\bar{\mathbf{e}}$, $\bar{\mathbf{e}}^p$, $\bar{\mathbf{e}}^q$, $\text{var}[\mathbf{e}]$, $\text{cov}[\mathbf{e}, \mathbf{e}^p]$, $\text{cov}[\mathbf{e}, \mathbf{e}^q]$, and $\text{cov}[\mathbf{e}^p, \mathbf{e}^q]$.

Proof.

$$\begin{aligned} \text{cov}[v_{p,i}, v_{q,k}] &= \text{cov} \left[e_i^p \sum_j F_{i,j} e_j, e_k^q \sum_l F_{k,l} e_l \right] \\ &= \sum_{j,l} \text{cov}[F_{i,j} e_i^p e_j, F_{k,l} e_k^q e_l] \\ &= \sum_{j,l} \text{cov}[F_{i,j}, F_{k,l}] \mathbb{E}[e_i^p e_j] \mathbb{E}[e_k^q e_l] \\ &\quad + (\bar{F}_{i,j} \bar{F}_{k,l} + \text{cov}[F_{i,j}, F_{k,l}]) \text{cov}[e_i^p e_j, e_k^q e_l] \\ &\quad \text{since } \mathbf{F}, \mathbf{e} \text{ ind} \end{aligned} \quad (7.32)$$

$$\begin{aligned} &= \sum_{j,l} H_{i,j,k,l} (\bar{e}_i^p \bar{e}_j + \text{cov}[e_i^p, e_j]) (\bar{e}_k^q \bar{e}_l + \text{cov}[e_k^q, e_l]) \\ &\quad + (\bar{F}_{i,j} \bar{F}_{k,l} + H_{i,j,k,l}) \underbrace{\text{cov}[e_i^p e_j, e_k^q e_l]}_{\textcircled{1}} \end{aligned} \quad (7.33)$$

$$\begin{aligned} \textcircled{1} &= (\bar{e}_i^p \bar{e}_k^q + \text{cov}[e_i^p, e_k^q]) (\bar{e}_j \bar{e}_l + \text{cov}[e_j, e_l]) \\ &\quad + (\bar{e}_i^p \bar{e}_l + \text{cov}[e_i^p, e_l]) (\bar{e}_k^q \bar{e}_j + \text{cov}[e_k^q, e_j]) - 2\bar{e}_i^p \bar{e}_k^q \bar{e}_j \bar{e}_l \\ &\quad \text{since } \mathbf{e}, \mathbf{e}^p, \text{ and } \mathbf{e}^q \text{ Gaussian} \end{aligned} \quad (7.34)$$

□

Proposition 7.3. $\text{cov}[v_{p,\varepsilon^2}, v_{q,\varepsilon^2}]$

If $\mathbf{w}$ is Gaussian,

$$\begin{aligned}
\text{cov} [v_{p,\varepsilon^2}, v_{q,\varepsilon^2}] &= \eta \mathbb{E} [v_{p,\varepsilon^2}] \mathbb{E} [v_{q,\varepsilon^2}] / \bar{\gamma}^2 \\
&\quad + (\bar{\gamma}^2 + \eta) \sum_{n \in P, m \in Q} (4\bar{w}_n \bar{w}_m \Omega_{n,m} + 2\Omega_{n,m}^2) \bar{\varepsilon}_n^2 \bar{\varepsilon}_m^2 \\
&\quad + (\bar{\gamma}^2 + \eta) \sum_{n \in P \cap Q} (\bar{w}_n^4 + 6\bar{w}_n^2 \Omega_{n,n} + 3\Omega_{n,n}^2) \omega_n
\end{aligned} \tag{7.35}$$

and is computed in $O(N^2)$ given $\mathbb{E} [v_{p,\varepsilon^2}]$ and $\mathbb{E} [v_{q,\varepsilon^2}]$.

Proof.

$$\text{cov} [v_{p,\varepsilon^2}, v_{q,\varepsilon^2}] = \text{cov} \left[\gamma \sum_{n \in P} w_n^2 \varepsilon_n^2, \gamma \sum_{m \in Q} w_m^2 \varepsilon_m^2 \right] \tag{7.36}$$

$$\begin{aligned}
&= \text{cov} [\gamma, \gamma] \mathbb{E} \left[\sum_{n \in P} w_n^2 \varepsilon_n^2 \right] \mathbb{E} \left[\sum_{m \in Q} w_m^2 \varepsilon_m^2 \right] \\
&\quad + (\mathbb{E}^2 [\gamma] + \text{cov} [\gamma, \gamma]) \text{cov} \left[\sum_{n \in P} w_n^2 \varepsilon_n^2, \sum_{m \in Q} w_m^2 \varepsilon_m^2 \right]
\end{aligned}$$

$$\text{since } \gamma, \{\mathbf{w}, \boldsymbol{\varepsilon}^2\} \text{ ind} \tag{7.37}$$

$$\begin{aligned}
&= \eta \frac{\mathbb{E} [v_{p,\varepsilon^2}]}{\bar{\gamma}} \frac{\mathbb{E} [v_{q,\varepsilon^2}]}{\bar{\gamma}} \\
&\quad + (\bar{\gamma}^2 + \eta) \underbrace{\text{cov} \left[\sum_{n \in P} w_n^2 \varepsilon_n^2, \sum_{m \in Q} w_m^2 \varepsilon_m^2 \right]}_{\textcircled{1}}
\end{aligned} \tag{7.38}$$

$$\textcircled{1} = \text{cov} \left[\sum_{n \in P} w_n^2 \varepsilon_n^2, \sum_{m \in Q} w_m^2 \varepsilon_m^2 \right] \quad (7.39)$$

$$= \sum_{n \in P, m \in Q} \text{cov} [w_n^2 \varepsilon_n^2, w_m^2 \varepsilon_m^2] \quad (7.40)$$

$$= \sum_{n \in P, m \in Q} \text{cov} [w_n^2, w_m^2] \bar{\varepsilon}_n^2 \bar{\varepsilon}_m^2 + (\mathbb{E} [w_n^2] \mathbb{E} [w_m^2] + \text{cov} [w_n^2, w_m^2]) \text{cov} [\varepsilon_n^2, \varepsilon_m^2] \quad (7.41)$$

since $\mathbf{w}, \boldsymbol{\varepsilon}^2$ ind

$$= \sum_{n \in P, m \in Q} \text{cov} [w_n^2, w_m^2] \bar{\varepsilon}_n^2 \bar{\varepsilon}_m^2 + \sum_{n \in P \cap Q} (\mathbb{E}^2 [w_n^2] + \text{var} [w_n^2]) \text{var} [\varepsilon_n^2] \quad \text{since elements of } \boldsymbol{\varepsilon}^2 \text{ ind} \quad (7.42)$$

$$= \sum_{n \in P, m \in Q} (4\bar{w}_n \bar{w}_m \Omega_{n,m} + 2\Omega_{n,m}^2) \bar{\varepsilon}_n^2 \bar{\varepsilon}_m^2 + \sum_{n \in P \cap Q} (\mathbb{E}^2 [w_n^2] + 4\bar{w}_n^2 \Omega_{n,n} + 2\Omega_{n,n}^2) \omega_n \quad \text{since } \mathbf{w} \text{ Gaussian} \quad (7.43)$$

$$= \sum_{n \in P, m \in Q} (4\bar{w}_n \bar{w}_m \Omega_{n,m} + 2\Omega_{n,m}^2) \bar{\varepsilon}_n^2 \bar{\varepsilon}_m^2 + \sum_{n \in P \cap Q} \left([\bar{w}_n^2 + \Omega_{n,n}]^2 + 4\bar{w}_n^2 \Omega_{n,n} + 2\Omega_{n,n}^2 \right) \omega_n \quad (7.44)$$

$$= \sum_{n \in P, m \in Q} (4\bar{w}_n \bar{w}_m \Omega_{n,m} + 2\Omega_{n,m}^2) \bar{\varepsilon}_n^2 \bar{\varepsilon}_m^2 + \sum_{n \in P \cap Q} (\bar{w}_n^4 + 6\bar{w}_n^2 \Omega_{n,n} + 3\Omega_{n,n}^2) \omega_n \quad (7.45)$$

□

Proposition 7.4. $\text{cov} [v_{p,i}, v_{q,\varepsilon^2}]$

If $\mathbf{w}$ is Gaussian,

$$\text{cov} [v_{p,i}, v_{q,\varepsilon^2}] = c + \sum_j (\bar{F}_{i,j} \bar{\gamma} + h_{i,j}) \text{cov} \left[e_i^p e_j, \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] \quad (7.46)$$

$$\approx c$$

$$\text{where } c = \frac{\text{E} [v_{q,\varepsilon^2}]}{\bar{\gamma}} \sum_j h_{i,j} (\bar{e}_i^p \bar{e}_j + \text{cov} [e_i^p, e_j]) \quad (7.47)$$

$$\begin{aligned} \text{and } \text{cov} \left[e_i^p e_j, \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] &= 2 \sum_{m \in P, n \in Q, s} (\bar{E}_{m,i} \bar{E}_{s,j} + X_{m,i,s,j}) \bar{\varepsilon}_n^2 \\ &\quad \left[\bar{w}_m \bar{w}_n \Omega_{s,n} + \bar{w}_s \bar{w}_n \Omega_{m,n} + \Omega_{m,n} \Omega_{s,n} \right] \end{aligned} \quad (7.48)$$

Computation is $O(K)$ for approximate covariance c and $O(KN^3)$ for exact $\text{cov} [v_{p,i}, v_{q,\varepsilon^2}]$ given $\bar{\mathbf{e}}^p, \bar{\mathbf{e}}, \text{cov} [\mathbf{e}^p, \mathbf{e}]$, and $\text{E} [v_{q,\varepsilon^2}]$.

The approximation ignores price movement's linking factor exposures and idiosyncratic risk. The effect is secondary and, in a model and one built from estimates, seems of needless precision. The result is what would be if idiosyncratic risk arose from an independent set of portfolio weights following the same distribution, i.e., if $\sigma^2 = \mathbf{w}^T \mathbf{E} \mathbf{F} \mathbf{F}^T \mathbf{w} + \gamma \mathbf{z}^T \text{diag} [\varepsilon^2] \mathbf{z}$ where $\mathbf{z}$ ind of $\mathbf{w}$, $\text{E} [\mathbf{z}] = \text{E} [\mathbf{w}]$, $\text{var} [\mathbf{z}] = \text{var} [\mathbf{w}]$. Thus, the matrix of contributions' approximated covariance is positive semidefinite.

Proof.

$$\text{cov} [v_{p,i}, v_{q,\varepsilon^2}] = \text{cov} \left[e_i^p \sum_j F_{i,j} e_j, \gamma \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] \quad (7.49)$$

$$= \sum_{n \in Q, j} \text{cov} [e_i^p F_{i,j} e_j, \gamma w_n^2 \varepsilon_n^2] \quad (7.50)$$

$$= \sum_j \text{cov} [F_{i,j}, \gamma] \text{E} [e_i^p e_j] \text{E} \left[\sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] \\ + (\bar{F}_{i,j} \bar{\gamma} + \text{cov} [F_{i,j}, \gamma]) \text{cov} \left[e_i^p e_j, \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] \quad (7.51)$$

since $\{\mathbf{F}, \gamma\}$ ind of $\{\mathbf{e}, \mathbf{w}, \varepsilon^2\}$

$$= \sum_j h_{i,j} (\bar{e}_i^p \bar{e}_j + \text{cov} [e_i^p, e_j]) \frac{\text{E} [v_{q,\varepsilon^2}]}{\bar{\gamma}} \\ + (\bar{F}_{i,j} \bar{\gamma} + h_{i,j}) \underbrace{\text{cov} \left[e_i^p e_j, \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right]}_{\textcircled{1}} \quad (7.52)$$

$$\textcircled{1} = \text{cov} \left[e_i^p e_j, \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] \\ = \text{cov} \left[\sum_{m \in P, s} E_{m,i} w_m E_{s,j} w_s, \sum_{n \in Q} w_n^2 \varepsilon_n^2 \right] \\ = \sum_{m \in P, n \in Q, s} \text{cov} [w_m w_s E_{m,i} E_{s,j}, w_n^2 \varepsilon_n^2] \\ = \sum_{m \in P, n \in Q, s} \text{E} [E_{m,i} E_{s,j}] \text{E} [\varepsilon_n^2] \text{cov} [w_m w_s, w_n^2] \quad (7.53)$$

since $\mathbf{E}, \mathbf{w}, \varepsilon^2$ ind

$$= \sum_{m \in P, n \in Q, s} (\bar{E}_{m,i} \bar{E}_{s,j} + X_{m,i,s,j}) \bar{\varepsilon}_n^2 \underbrace{\text{cov} [w_m w_s, w_n^2]}_{\textcircled{2}}$$

$$\textcircled{2} = 2 (\bar{w}_m \bar{w}_n \Omega_{s,n} + \bar{w}_s \bar{w}_n \Omega_{m,n} + \Omega_{m,n} \Omega_{s,n}) \quad (7.54)$$

since $\mathbf{w}$ Gaussian

□

7.5.3 Computational complexity

If the portfolio is divided into M subportfolios, the mean of the $M \times (K + 1)$ parts computes in $O(MK^2N^2)$ broken out in table 7.1. For variance, generally only the

Table 7.1: Computational complexity of reporting contributions' mean

Quantity	Complexity
$\bar{e}, \bar{e}^p$	$M \times O(KN)$
$\text{cov}[\mathbf{e}^p, \mathbf{e}]$	$M \times O(K^2N^2)$
$E[\mathbf{v}_p]$	$M \times O(K^2)$
$E[v_{p,\varepsilon^2}]$	$M \times O(N)$
$O(MK^2N^2)$	

range around a reported number is of interest; covariance with others is irrelevant with the exception of total portfolio variance, needed to convert numbers to standard deviation. Together, these compute in $O(MK^2N^2 + MK^4)$ broken out in table 7.2.

Table 7.2: Computational complexity of reporting contributions' variance

Quantity	Complexity
$\text{var}[\mathbf{e}^p], \text{var}[\mathbf{e}]$	$M \times O(K^2N^2)$
$\text{var}[\mathbf{v}_p], \text{var}[\mathbf{v}_{portfolio}]$	$M \times O(K^4)$
$\text{cov}[\mathbf{v}_p, \mathbf{v}_{portfolio}]$	$M \times O(K^4)$
$\text{var}[v_{p,\varepsilon^2}], \text{var}[v_{portfolio,\varepsilon^2}]$	$M \times O(N^2)$
$\text{cov}[v_{p,\varepsilon^2}, v_{portfolio,\varepsilon^2}]$	$M \times O(N^2)$
$\text{approx cov}[v_p, v_{p,\varepsilon^2}], \text{cov}[v_{portfolio}, v_{portfolio,\varepsilon^2}]$	$M \times O(K^2)$
$\text{approx cov}[v_p, v_{portfolio,\varepsilon^2}], \text{cov}[v_{portfolio}, v_{p,\varepsilon^2}]$	$M \times O(K^2)$
$O(MK^2N^2 + MK^4)$	

7.6 Uncertain standard deviation contributions

Volatility is more tangible in standard deviation than variance, e.g., a portfolio standard deviation of 15% vs. variance of 225%². Can uncertain standard deviation

contributions be computed? Equation (7.5), $s_i = v_i/\sigma$, converts variance contributions to standard deviation. However, nonlinearity makes the mean and covariance of $\mathbf{s}$, unlike $\mathbf{v}$, unavailable.

The usual estimate of center would be $E[\mathbf{s}] = E[\mathbf{v}/\sqrt{\sigma^2}]$. Instead, a notional center of

$$\text{ctr}[\mathbf{s}] = \frac{E[\mathbf{v}]}{\sqrt{E[\sigma^2]}} \quad (7.55)$$

agrees in the absence of uncertainty and properly sums to a notional center of portfolio volatility, $\sum_i \text{ctr}[s_i] = \sum_i E[v_i] / \sqrt{E[\sigma^2]} = \sqrt{E[\sigma^2]}$.

Approximate covariance comes from linearizing sensitivity to variance contributions and taking the variance of the linearization.

Proposition 7.5. *Approximate var $[\mathbf{s}]$*

If contributions to variance $\mathbf{v}$ have mean $E[\mathbf{v}]$ and covariance $\text{var}[\mathbf{v}]$, contributions to standard deviation $\mathbf{s}$ have approximate covariance

$$\begin{aligned} \text{var}[\mathbf{s}] \approx & \frac{\text{var}[\mathbf{v}]}{E[\sigma^2]} - \frac{E[\mathbf{v}] \text{cov}[\sigma^2, \mathbf{v}] + \text{cov}[\mathbf{v}, \sigma^2] E[\mathbf{v}]^T}{2(E[\sigma^2])^2} \\ & + \frac{\text{var}[\sigma^2] E[\mathbf{v}] E[\mathbf{v}]^T}{4E[\sigma^2] (E[\sigma^2])^2} \end{aligned} \quad (7.56)$$

$$\text{So } \text{var}[s_k] \approx \frac{\text{var}[v_k]}{E[\sigma^2]} - \frac{\text{cov}[v_k, \sigma^2] E[v_k]}{E^2[\sigma^2]} + \frac{\text{var}[\sigma^2] E^2[v_k]}{4E^3[\sigma^2]} \quad (7.57)$$

$$\text{var}[\sigma] \approx \frac{\text{var}[\sigma^2]}{4E[\sigma^2]} \quad (7.58)$$

$$\text{where } E[\sigma^2] = \mathbf{1}^T E[\mathbf{v}] \quad (7.59)$$

$$\text{cov}[\mathbf{v}, \sigma^2] = \text{var}[\mathbf{v}] \mathbf{1} \quad (7.60)$$

$$\text{var}[\sigma^2] = \mathbf{1}^T \text{var}[\mathbf{v}] \mathbf{1} \quad (7.61)$$

Proof.

First note that portfolio variance is the sum of the variance contributions, $\sigma^2 = \mathbf{1}^T \mathbf{v}$. So, $E[\sigma^2] = \mathbf{1}^T E[\mathbf{v}]$, $\text{var}[\sigma^2] = \mathbf{1}^T \text{var}[\mathbf{v}] \mathbf{1}$, and $\text{cov}[\mathbf{v}, \sigma^2] = \text{var}[\mathbf{v}] \mathbf{1}$.

For the covariance of $\mathbf{s}$, linearize around the mean of $\mathbf{v}$ and σ^2 .

$$\mathbf{s} = \mathbf{v}/\sqrt{\sigma^2} \quad (7.62)$$

$$s_i \approx \text{constant} + \sum_j \frac{\partial s_i}{\partial v_j} (v_j - \mathbb{E}[v_j]) + \frac{\partial s_i}{\partial \sigma^2} (\sigma^2 - \mathbb{E}[\sigma^2]) \quad (7.63)$$

$$\frac{\partial s_i}{\partial v_j} = \begin{cases} \frac{1}{\sqrt{\mathbb{E}[\sigma^2]}} & j = i \\ 0 & j \neq i \end{cases} \quad (7.64)$$

$$\frac{\partial s_i}{\partial \sigma^2} = -\frac{\mathbb{E}[v_i]}{2(\mathbb{E}[\sigma^2])^{3/2}} \quad (7.65)$$

Approximate the covariance as the covariance of the linearization.

$$\begin{aligned} \mathbf{J} &= \text{Jacobian of } \begin{bmatrix} \mathbf{s} \\ \sigma \end{bmatrix} \text{ to } \mathbf{v}, \sigma^2 \\ &= \frac{1}{\sqrt{\mathbb{E}[\sigma^2]}} \begin{pmatrix} 1 & 0 & \cdots & 0 & -\frac{\mathbb{E}[v_1]}{2\mathbb{E}[\sigma^2]} \\ 0 & 1 & \ddots & \vdots & -\frac{\mathbb{E}[v_2]}{2\mathbb{E}[\sigma^2]} \\ \vdots & \ddots & \ddots & 0 & \vdots \\ 0 & \cdots & 0 & 1 & -\frac{\mathbb{E}[v_K]}{2\mathbb{E}[\sigma^2]} \\ 0 & \cdots & \cdots & 0 & 1/2 \end{pmatrix} \end{aligned} \quad (7.66)$$

$$= \frac{1}{\sqrt{\mathbb{E}[\sigma^2]}} \begin{pmatrix} \mathbf{I} & -\frac{1}{2\mathbb{E}[\sigma^2]} \mathbb{E}[\mathbf{v}] \\ \mathbf{0} & 1/2 \end{pmatrix} \quad (7.67)$$

$$\text{var} \begin{bmatrix} \mathbf{s} \\ \sigma \end{bmatrix} \approx \mathbf{J} \text{var} \begin{bmatrix} \mathbf{v} \\ \sigma^2 \end{bmatrix} \mathbf{J}^T \quad (7.68)$$

$$= \mathbf{J} \begin{pmatrix} \text{var}[\mathbf{v}] & \text{cov}[\mathbf{v}, \sigma^2] \\ \text{cov}[\sigma^2, \mathbf{v}] & \text{var}[\sigma^2] \end{pmatrix} \mathbf{J}^T \quad (7.69)$$

Note: this correctly gives $\text{var}[\sum_i s_i] = \text{var}[\sigma]$.

$$\begin{aligned} \text{So } \text{var} [\mathbf{s}] &\approx \frac{\text{var} [\mathbf{v}]}{\text{E} [\sigma^2]} - \frac{\text{E} [\mathbf{v}] \text{cov} [\sigma^2, \mathbf{v}] + \text{cov} [\mathbf{v}, \sigma^2] \text{E} [\mathbf{v}]^T}{2 (\text{E} [\sigma^2])^2} \\ &\quad + \frac{\text{var} [\sigma^2]}{4\text{E} [\sigma^2]} \frac{\text{E} [\mathbf{v}] \text{E} [\mathbf{v}]^T}{(\text{E} [\sigma^2])^2} \end{aligned} \quad (7.70)$$

$$\text{cov} [\mathbf{s}, \sigma] \approx \frac{\text{cov} [\mathbf{v}, \sigma^2]}{2 \text{E} [\sigma^2]} - \frac{\text{var} [\sigma^2]}{4\text{E} [\sigma^2]} \frac{\text{E} [\mathbf{v}]}{\text{E} [\sigma^2]} \quad (7.71)$$

$$\text{var} [\sigma] \approx \frac{\text{var} [\sigma^2]}{4 \text{E} [\sigma^2]} \quad (7.72)$$

□

7.7 Modeling weight dynamics

The perturbation on portfolio weights is specified by their covariance over the horizon of interest. Not normally at hand, it can be derived from a model of return covariance, initial holdings, and if a benchmark is involved, benchmark weights and whether it is fixed weight or floating.

The return model enlisted is likely of shorter horizon than the risk model. For example, one might be interested in how the annual risk of the portfolio can vary until the next rebalance period. In that case, the return model's horizon is ideally the time until rebalancing.

Proposition 7.6. *Future weight covariance Ω*

$$\Omega \approx \mathbf{J} \text{var} [\mathbf{r}] \mathbf{J}^T \quad (7.73)$$

where

$$\mathbf{J} = \begin{cases} -(\bar{\mathbf{w}} + \mathbf{b})(\bar{\mathbf{w}} + \mathbf{b})^T + \text{diag} [\bar{\mathbf{w}} + \mathbf{b}] & \text{fixed benchmark} \\ -(\bar{\mathbf{w}} + \mathbf{b})(\bar{\mathbf{w}} + \mathbf{b})^T + \text{diag} [\bar{\mathbf{w}}] + \mathbf{b}\mathbf{b}^T & \text{floating benchmark} \end{cases} \quad (7.74)$$

$\bar{\mathbf{w}} = (n \times 1)$ *current benchmark-relative weights*

$\mathbf{b} = (n \times 1)$ *benchmark weights*

$\mathbf{r} = (n \times 1)$ *future returns*

If there isn't a benchmark, $\mathbf{b} = \mathbf{0}$; otherwise, $\mathbf{b}$ sums to 1. In both cases, $\bar{\mathbf{w}} + \mathbf{b}$ sums to 1.

Proof.

Let $\mathbf{w} = (n \times 1)$ future benchmark-relative weights.

With a fixed weight benchmark:

$$w_i = \frac{(1 + r_i)(\bar{w}_i + b_i)}{\sum_k (1 + r_k)(\bar{w}_k + b_k)} - b_i \quad (7.75)$$

$$= \alpha_i - b_i \quad (7.76)$$

With a floating benchmark:

$$w_i = \frac{(1 + r_i)(\bar{w}_i + b_i)}{\sum_k (1 + r_k)(\bar{w}_k + b_k)} - \frac{(1 + r_i)b_i}{\sum_k (1 + r_k)b_k} \quad (7.77)$$

$$= \alpha_i - \beta_i \quad (7.78)$$

$$\text{where } \alpha_i = \frac{(1 + r_i)(\bar{w}_i + b_i)}{\sum_k (1 + r_k)(\bar{w}_k + b_k)} \quad (7.79)$$

$$\beta_i = \frac{(1 + r_i)b_i}{\sum_k (1 + r_k)b_k} \quad (7.80)$$

The approximate weight covariance is the covariance of linearization around the mean return.

$$\begin{aligned} \mathbf{\Omega} &= \text{var}[\mathbf{w}] \\ &\approx \mathbf{J} \text{var}[\mathbf{r}] \mathbf{J}^T \end{aligned} \quad (7.81)$$

where $\mathbf{J}$ = Jacobian of $\mathbf{w}$ to $\mathbf{r}$ evaluated at $\mathbf{r} = \bar{\mathbf{r}}$

$$J_{i,j} = \frac{\partial w_i}{\partial r_j} = \begin{cases} \frac{\partial \alpha_i}{\partial r_j} & \text{fixed benchmark} \\ \frac{\partial \alpha_i}{\partial r_j} - \frac{\partial \beta_i}{\partial r_j} & \text{floating benchmark} \end{cases} \quad (7.82)$$

$$\frac{\partial \alpha_i}{\partial r_j} = -(\bar{w}_j + b_j) \frac{(1 + r_i)(\bar{w}_i + b_i)}{[\sum_k (1 + r_k)(\bar{w}_k + b_k)]^2} + \mathbf{1}_{\{i=j\}} \frac{(\bar{w}_i + b_i)}{[\sum_k (1 + r_k)(\bar{w}_k + b_k)]^2} \quad (7.83)$$

$$= -(\bar{w}_j + b_j)(\bar{w}_i + b_i) + \mathbf{1}_{\{i=j\}}(\bar{w}_i + b_i) \quad (7.84)$$

evaluated at $\mathbf{r} = \bar{\mathbf{r}} = \mathbf{0}$ ⁹ since $\sum_k \bar{w}_k + b_k = 1$

Likewise,

$$\frac{\partial \beta_i}{\partial r_j} = -b_j b_i + \mathbf{1}_{\{i=j\}} b_i \quad (7.85)$$

$$\mathbf{J} = \begin{cases} -(\bar{\mathbf{w}} + \mathbf{b})(\bar{\mathbf{w}} + \mathbf{b})^T + \text{diag}[\bar{\mathbf{w}} + \mathbf{b}] & \text{fixed benchmark} \\ -(\bar{\mathbf{w}} + \mathbf{b})(\bar{\mathbf{w}} + \mathbf{b})^T + \text{diag}[\bar{\mathbf{w}} + \mathbf{b}] + \mathbf{b}\mathbf{b}^T - \text{diag}[\mathbf{b}] & \text{floating benchmark} \end{cases} \quad (7.86)$$

□

7.8 An example

Data are 106 periods of weekly Wed–Wed returns from Wed Dec 20, 2017 through Tue Dec 31, 2019. The first 80 periods, 12/20/2017–7/3/2019, are the training set. The final 26 periods, 7/3/2019–12/31/2019, are the testing set. The universe consists of the 474 securities in the S&P 500 as of Dec 20, 2017 with complete series. Each security belongs to one of eleven sectors. ETFs—IWV-iShares Russell 3000, IWB-iShares Russell 1000, IWM-iShares Russell 2000, IWF-iShares Russell 1000 Growth, IWD-iShares Russell 1000 Value—are used in the risk factors.

An uncertain covariance model is constructed from the training set as follows:

1. Three factors are from ETF returns: market = r_{3000} , size = $r_{1000} - r_{2000}$, and growth-value = $r_{1000G} - r_{1000V}$. Inferred on security returns stripped of these,

⁹Omitted here for clarity, considering drift ($\bar{\mathbf{r}} \neq \mathbf{0}$) adds scalars.

four principal components are added for a total of 7 factors. Their mean covariance $\bar{\mathbf{F}}$ is the sample covariance of the series.

2. Exposures $\mathbf{E}$ and idiosyncratic variances $\boldsymbol{\varepsilon}^2$ are first point-estimated as the coefficients and variance of the residual in OLS regression of security returns against factors. Using the universe median and variance of these as a prior—exposures Gaussian, idiosyncratic variance inverse gamma—they are re-estimated with mean and variance via Bayesian regression, yielding $\bar{\mathbf{E}}, \mathbf{X}, \bar{\boldsymbol{\varepsilon}}^2, \boldsymbol{\omega}$. ($\mathbf{X}$ assumes no exposure covariance between securities.)
3. For simplicity, $\mathbf{H}$, the variance of factor covariance $\mathbf{F}$, is taken as the standard sampling error.¹⁰ Note that this underestimates since the real uncertainty isn't from sample size but changes in state. For example, VIX^2 , a good proxy for $F_{1,1}$, observed weekly over the same training period has variance $275 \times H_{1,1,1,1}$.
4. A measure of the idiosyncratic variance scalar γ is the cross-sectional median of each period's squared OLS regression residuals, divided by the mean across time for normalization. γ 's variance η is the variance of this series; its mean $\bar{\gamma}$ is 1.
5. $\mathbf{h}$, the covariance between $\mathbf{F}$ and γ , is for simplicity assumed to be zero.
6. To report in annualized variance, $\mathbf{F}$ and $\bar{\gamma}$ are multiplied by 52, and $\mathbf{H}$, η , and $\mathbf{h}$ by 52^2 .

Portfolio weight covariance $\boldsymbol{\Omega}$ is computed by applying the method in section 7.7 to the horizon-scaled sample covariance of the training set.

7.8.1 Portfolio 1: Diversified long, cash benchmark

The 472 securities are equal-weighted (0.21%) and analyzed with a 26 week horizon and no benchmark. Subportfolios are holdings in each sector.

Table 7.3 shows the mean and standard deviation of factor exposures contributed by sector. Standard deviations are small because the portfolio is diversified and exposure uncertainty was modeled as uncorrelated across securities.

¹⁰For the usual unbiased variance estimate of multivariate Gaussian (a, b, c, d) from N points, $\text{cov}[\hat{\sigma}_{ab}, \hat{\sigma}_{cd}] = (\sigma_{a,c} \sigma_{b,d} + \sigma_{a,d} \sigma_{b,c}) / (N - 1)$.

Table 7.3: Diversified long—mean and SD of factor exposures contributed by sector

	N	wt %	Factor						
			mkt	size	gv	pc1	pc2	pc3	pc4
All	474	100	1.059	-0.098	-0.333	0.326	-0.301	0.190	0.080
±			0.023	0.015	0.023	0.050	0.040	0.029	0.033
Comm Svcs	5	1.1	0.009	-0.002	-0.005	0.003	-0.009	0.002	0.002
			0.001	0.001	0.002	0.002	0.002	0.002	0.002
Cons Discr	76	16.0	0.174	-0.036	-0.038	0.125	-0.087	0.042	0.065
			0.010	0.006	0.007	0.012	0.011	0.010	0.012
Cons Staples	32	6.8	0.045	0.005	-0.032	-0.032	-0.061	0.033	0.010
			0.004	0.003	0.005	0.006	0.007	0.006	0.005
Energy	28	5.9	0.080	-0.023	-0.048	0.007	0.070	0.006	0.106
			0.011	0.005	0.007	0.005	0.011	0.005	0.016
Financials	65	13.7	0.166	-0.007	-0.106	0.110	-0.028	-0.073	-0.058
			0.011	0.004	0.009	0.010	0.005	0.008	0.006
Healthcare	57	12.0	0.129	-0.019	-0.020	-0.004	-0.057	-0.013	0.029
			0.009	0.004	0.006	0.006	0.008	0.007	0.008
Industrials	66	13.9	0.175	-0.001	-0.061	0.077	0.018	0.026	-0.056
			0.009	0.004	0.006	0.009	0.007	0.007	0.008
Tech	64	13.5	0.164	-0.005	0.048	0.108	-0.016	0.035	-0.023
			0.013	0.004	0.007	0.013	0.008	0.009	0.008
Materials	22	4.6	0.059	-0.003	-0.023	0.037	-0.003	0.011	-0.006
			0.004	0.002	0.003	0.004	0.004	0.004	0.004
Real Estate	32	6.8	0.037	-0.013	-0.020	-0.034	-0.085	0.064	0.004
			0.004	0.003	0.003	0.005	0.010	0.008	0.004
Utilities	27	5.7	0.021	0.007	-0.030	-0.072	-0.044	0.056	0.008
			0.003	0.002	0.005	0.011	0.008	0.009	0.004

Table 7.4 shows risk decomposed by factor conventionally and with uncertainty. Conventional uses the point estimate model left after discarding uncertainty parameters:

$$C_{\text{conventional}} = \bar{\mathbf{E}}\bar{\mathbf{F}}\bar{\mathbf{E}}^T + \bar{\gamma} \text{diag}(\bar{\mathbf{e}}^2) \quad (7.87)$$

Notice that uncertain contributions come with an approximate standard deviation, e.g., portfolio variance of 14.47 ± 1.20 vs. the conventional's point estimate of 14.45. The table's final 4 columns split the uncertain's center into sources—means and different uncertainties as described by equations (7.21) and (7.22).

Table 7.5 shows ex-ante modeled and realized risk decomposed by sector. Realized

Table 7.4: Diversified long—conventional and uncertain risk contributions by factor

	Conventional	Uncertain		means	$\mathbf{w}$ unc	$\mathbf{E}$ unc	$\mathbf{w}, \mathbf{E}$ unc
		center	$\pm$				
All	14.45	14.47	1.20	14.44	0.03	0	0
mkt	14.68	14.67	2.39	14.66	0.01	0	0
size	0.24	0.24	0.10	0.24	0.00	0	0
gv	-1.11	-1.11	0.34	-1.11	0.00	0	0
pc1	0.30	0.31	0.26	0.30	0.01	0	0
pc2	0.20	0.21	0.20	0.20	0.00	0	0
pc3	0.07	0.07	0.11	0.07	0.00	0	0
pc4	0.01	0.01	0.05	0.01	0.00	0	0
ε^2	0.06	0.07	0.02	0.06	0.00		

contributions are measured with two sets of portfolio weights: $\mathbf{w}_0$ -as constructed, $\mathbf{w}_T$ -at the horizon (26 weeks in the future) adjusted for price movement. Realized covariance is the sample covariance of returns over the test set (the same 26 weeks).

Table 7.5: Diversified long—ex-ante modeled and realized risk contributions by sector

	Conventional	Uncertain		Realized	
		center	$\pm$	$\mathbf{w}_0$	$\mathbf{w}_T$
All	14.45	14.47	1.20	13.84	13.78
Comm Svcs	0.13	0.13	0.02	0.13	0.12
Cons Discr	2.59	2.60	0.17	2.76	2.77
Cons Staples	0.54	0.54	0.07	0.30	0.30
Energy	0.99	0.99	0.16	1.61	1.41
Financials	2.06	2.06	0.17	2.43	2.45
Healthcare	1.81	1.81	0.13	1.30	1.31
Industrials	2.29	2.29	0.14	2.44	2.44
Tech	2.58	2.58	0.23	1.95	2.08
Materials	0.79	0.79	0.06	0.70	0.69
Real Estate	0.53	0.53	0.09	0.13	0.13
Utilities	0.16	0.16	0.08	0.09	0.09

7.8.2 Portfolio 2: 44 stock long-short, cash benchmark

A 44 stock long-short portfolio is constructed from the the top 4 securities alphabetically by name in each sector. The first two are long; the remaining two are short.

Weights come from minimizing variance as measured by the conventional factor model of equation (7.87) with longs constrained positive and summing to 1 and shorts negative and summing to -1. The horizon is again 26 weeks for modeled future weight covariance Ω and realized weights w_T .

Table 7.6 shows factor exposures by side. Tables 7.7, 7.8, and 7.9 are as the preceding tables 7.3, 7.4, and 7.5.

Table 7.6: 44 stock long-short—mean and SD of factor exposures contributed by side

	N	wt %	Factor						
			mkt	size	gv	pc1	pc2	pc3	pc4
All	44	0	0.008	0.016	0.003	0.014	0.061	-0.005	0.022
±			0.062	0.068	0.079	0.093	0.099	0.106	0.110
Long	22	100	0.851	0.031	-0.286	-0.105	-0.570	0.351	-0.088
			0.127	0.048	0.066	0.085	0.096	0.087	0.080
Short	22	-100	-0.843	-0.015	0.289	0.119	0.631	-0.356	0.110
			0.146	0.051	0.070	0.112	0.107	0.088	0.087

In this less diversified portfolio, uncertainty has a measurable effect on contributions' center—reported in the last four columns of table 7.8—accounting for 14% ($\frac{0.37+0.34+0.02}{5.22}$) of portfolio volatility.

7.9 Summary

Risk contribution information enters throughout investment management. Because what's measured is behavior in the direction of net risk, the portfolio's actual risks, viewed statically, can be hidden by hedging alignments. Looking instead with parameters and portfolio weights modeled as random evaluates the system across the distribution of what's possible. Latent fragility surfaces, and contributions go from a point estimate to having a center and range, offering robust and realistic assessments of risk.

Table 7.7: 44 stock long-short—mean and SD of factor exposures contributed by sector

	N	wt %	Factor						
			mkt	size	gv	pc1	pc2	pc3	pc4
All	44	0.0	0.008	0.016	0.003	0.014	0.061	-0.005	0.022
±			0.062	0.068	0.079	0.093	0.099	0.106	0.110
Comm Svcs	4	3.0	0.021	0.001	-0.059	-0.015	-0.038	-0.029	0.006
			0.017	0.015	0.020	0.022	0.024	0.025	0.025
Cons Discr	4	-1.2	-0.020	-0.011	0.055	-0.041	-0.032	-0.026	0.018
			0.020	0.017	0.024	0.027	0.026	0.027	0.028
Cons Staples	4	3.9	0.031	-0.020	-0.031	0.001	-0.007	0.047	0.058
			0.015	0.017	0.021	0.022	0.025	0.030	0.028
Energy	4	-2.0	-0.017	0.013	0.011	-0.000	0.027	0.016	-0.039
			0.014	0.012	0.014	0.013	0.018	0.017	0.028
Financials	4	-6.4	-0.065	-0.011	0.005	-0.046	0.082	-0.069	0.007
			0.026	0.025	0.030	0.031	0.037	0.039	0.041
Healthcare	4	1.2	-0.015	0.025	0.014	-0.028	-0.026	-0.019	-0.048
			0.021	0.018	0.019	0.023	0.025	0.027	0.030
Industrials	4	1.5	0.027	0.017	-0.003	0.011	0.054	0.024	0.026
			0.017	0.013	0.015	0.020	0.022	0.021	0.023
Tech	4	6.0	0.070	-0.007	-0.016	0.011	-0.036	0.005	-0.060
			0.028	0.024	0.030	0.031	0.034	0.035	0.039
Materials	4	-3.8	-0.046	-0.009	0.015	0.030	0.009	0.027	0.028
			0.022	0.021	0.024	0.027	0.029	0.032	0.034
Real Estate	4	-1.1	-0.018	0.021	0.016	0.023	-0.016	0.026	-0.009
			0.018	0.027	0.030	0.033	0.041	0.042	0.041
Utilities	4	-1.2	0.041	-0.003	-0.004	0.067	0.045	-0.005	0.036
			0.021	0.030	0.033	0.041	0.042	0.044	0.045

Table 7.8: 44 stock long-short—conventional and uncertain risk contributions by factor

	Conventional	Uncertain		means	<i>w</i> unc	<i>E</i> unc	<i>w</i> , <i>E</i> unc
		center	±				
All	4.85	5.22	1.03	4.50	0.37	0.34	0.02
mkt	0.00	0.13	0.20	0.00	0.08	0.05	0
size	0.00	0.05	0.08	0.00	0.00	0.04	0
gv	0.00	0.06	0.11	0.00	0.01	0.05	0
pc1	0.00	0.07	0.10	0.00	0.01	0.05	0
pc2	0.02	0.08	0.12	0.02	0.01	0.05	0
pc3	0.00	0.06	0.08	0.00	0.00	0.05	0
pc4	0.00	0.06	0.09	0.00	0.01	0.05	0
ε ²	4.81	4.71	1.12	4.46	0.25		

Table 7.9: 44 stock long-short—ex-ante modeled and realized risk contributions by sector

	Conventional	Uncertain		Realized	
		center	±	<i>w</i> ₀	<i>w</i> _{<i>T</i>}
All	4.85	5.22	1.03	4.47	4.30
Comm Svcs	0.26	0.29	0.11	0.23	0.16
Cons Discr	0.32	0.35	0.14	0.38	0.54
Cons Staples	0.34	0.36	0.11	0.21	0.27
Energy	0.14	0.17	0.10	1.11	0.59
Financials	0.64	0.68	0.22	-0.16	-0.28
Healthcare	0.32	0.37	0.15	0.36	0.32
Industrials	0.23	0.25	0.11	0.02	0.03
Tech	0.55	0.61	0.29	0.19	0.48
Materials	0.45	0.49	0.17	1.25	1.18
Real Estate	0.71	0.74	0.18	0.34	0.44
Utilities	0.87	0.90	0.27	0.54	0.56

7.10 Mathematical results

Lemma 7.1.

For $\begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma}),$

$$\begin{aligned} \text{E}[x_1x_2x_3x_4] &= \mu_1\mu_2\mu_3\mu_4 \\ &\quad + \mu_1\mu_2\sigma_{3,4} + \mu_3\mu_4\sigma_{1,2} + \sigma_{1,2}\sigma_{3,4} \\ &\quad + \mu_1\mu_3\sigma_{2,4} + \mu_2\mu_4\sigma_{1,3} + \sigma_{1,3}\sigma_{2,4} \\ &\quad + \mu_1\mu_4\sigma_{2,3} + \mu_2\mu_3\sigma_{1,4} + \sigma_{1,4}\sigma_{2,3} \end{aligned}$$

Proof.

$$\begin{aligned}
\mathbb{E}[x_1 x_2 x_3 x_4] &= \mathbb{E} \left[\prod_{i=1}^4 [\mu_i + (x_i - \mu_i)] \right] \\
&= \mu_1 \mu_2 \mu_3 \mu_4 \\
&\quad + \mu_1 \mu_2 \sigma_{3,4} + \mu_3 \mu_4 \sigma_{1,2} \\
&\quad + \mu_1 \mu_3 \sigma_{2,4} + \mu_2 \mu_4 \sigma_{1,3} \\
&\quad + \mu_1 \mu_4 \sigma_{2,3} + \mu_2 \mu_3 \sigma_{1,4} \\
&\quad + \mathbb{E} \left[\prod_{i=1}^4 (x_i - \mu_i) \right] \\
&\text{since odd moments are 0}
\end{aligned}$$

$$\mathbb{E} \left[\prod_{i=1}^4 (x_i - \mu_i) \right] = \sigma_{1,2} \sigma_{3,4} + \sigma_{1,3} \sigma_{2,4} + \sigma_{1,4} \sigma_{2,3}$$

by Isserlis' theorem [Isserlis \[1918\]](#) since 0 mean Gaussians

$$\text{cov}[x_1 x_2, x_3 x_4] = \mathbb{E}[x_1 x_2 x_3 x_4] - \underbrace{\mathbb{E}[x_1 x_2]}_{(\mu_1 \mu_2 + \sigma_{12})} \underbrace{\mathbb{E}[x_3 x_4]}_{(\mu_3 \mu_4 + \sigma_{34})}$$

□

Corollary 7.1. *Covariance of squared Gaussians*

$$\text{For } \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma}), \text{cov}[x_1^2, x_2^2] = 4\mu_1 \mu_2 \sigma_{1,2} + 2\sigma_{1,2}^2$$

Lemma 7.2. *For $\{x_1, x_2\}, \{y_1, y_2\}$ jointly independent,*

$$\begin{aligned}
\text{cov}[x_1 y_1, x_2 y_2] &= \mathbb{E}[x_1] \mathbb{E}[x_2] \text{cov}[y_1, y_2] + \text{cov}[x_1, x_2] \mathbb{E}[y_1] \mathbb{E}[y_2] \\
&\quad + \text{cov}[x_1, x_2] \text{cov}[y_1, y_2]
\end{aligned} \tag{7.91}$$

$$\begin{aligned}
&= (\mathbb{E}[x_1] \mathbb{E}[x_2] + \text{cov}[x_1, x_2]) \text{cov}[y_1, y_2] \\
&\quad + \text{cov}[x_1, x_2] \mathbb{E}[y_1] \mathbb{E}[y_2]
\end{aligned} \tag{7.92}$$

Proof.

$$\begin{aligned}
\text{cov}[x_1 y_1, x_2 y_2] &= \text{E}[x_1 y_1 x_2 y_2] - \text{E}[x_1 y_1] \text{E}[x_2 y_2] \\
&= \text{E}[x_1 x_2] \text{E}[y_1 y_2] - \text{E}[x_1] \text{E}[y_1] \text{E}[x_2] \text{E}[y_2] \\
&= \left(\text{E}[x_1] \text{E}[x_2] + \text{cov}[x_1, x_2] \right) \left(\text{E}[y_1] \text{E}[y_2] + \text{cov}[y_1, y_2] \right) \\
&\quad - \text{E}[x_1] \text{E}[y_1] \text{E}[x_2] \text{E}[y_2] \\
&= \text{E}[x_1] \text{E}[x_2] \text{cov}[y_1, y_2] + \text{cov}[x_1, x_2] \text{E}[y_1] \text{E}[y_2] \\
&\quad + \text{cov}[x_1, x_2] \text{cov}[y_1, y_2]
\end{aligned}$$

□

7.11 Mean and variance of factor covariance from uncertain covariance model

The covariance model in this paper is a generalization of the form in [Shah \[2015\]](#), which imposes the following structure on $\mathbf{F}$:

$$\mathbf{C} = \mathbf{E} \left(\overbrace{\mathbf{M} \text{diag}[\boldsymbol{\theta}] \mathbf{M}^T}^{\mathbf{F}} \right) \mathbf{E}^T + \gamma \text{diag}[\boldsymbol{\varepsilon}^2] \quad (7.93)$$

where $\mathbf{M} = (K \times K)$ fixed

$\boldsymbol{\theta} = (K \times 1)$ random and dependent on γ

The values $\text{E}[\mathbf{F}]$, $\text{var}[\mathbf{F}]$, and $\text{cov}[\mathbf{F}, \gamma]$ used in this paper are easily computed from that form.

$$\text{E}[\mathbf{F}] = \mathbf{M} \text{diag}[\text{E}[\boldsymbol{\theta}]] \mathbf{M}^T \quad (7.94)$$

$$\text{cov}[\mathbf{F}, \gamma] = \mathbf{M} \text{diag}[\text{cov}[\boldsymbol{\theta}, \gamma]] \mathbf{M}^T \quad (7.95)$$

$$\text{cov}[F_{i,j}, F_{k,l}] = \sum_{s,t} M_{i,s} M_{j,s} M_{k,t} M_{l,t} \text{cov}[\theta_s, \theta_t] \quad (7.96)$$

Proof.

$$\begin{aligned}
\mathbf{F} &= \mathbf{M} \operatorname{diag} [\boldsymbol{\theta}] \mathbf{M}^T \\
\mathbb{E} [\mathbf{F}] &= \mathbb{E} [\mathbf{M} \operatorname{diag} [\boldsymbol{\theta}] \mathbf{M}^T] \\
&= \mathbf{M} \operatorname{diag} [\mathbb{E} [\boldsymbol{\theta}]] \mathbf{M}^T \\
\operatorname{cov} [\mathbf{F}, \boldsymbol{\gamma}] &= \operatorname{cov} [\mathbf{M} \operatorname{diag} [\boldsymbol{\theta}] \mathbf{M}^T, \boldsymbol{\gamma}] \\
&= \mathbf{M} \operatorname{diag} [\operatorname{cov} [\boldsymbol{\theta}, \boldsymbol{\gamma}]] \mathbf{M}^T \\
F_{i,j} &= \sum_k M_{i,k} M_{j,k} \theta_k \\
\operatorname{cov} [F_{i,j}, F_{k,l}] &= \operatorname{cov} \left[\sum_s M_{i,s} M_{j,s} \theta_s, \sum_t M_{k,t} M_{l,t} \theta_t \right] \\
&= \sum_{s,t} M_{i,s} M_{j,s} M_{k,t} M_{l,t} \operatorname{cov} [\theta_s, \theta_t]
\end{aligned}$$

□

CHAPTER 8

CONCLUSION

Motivated largely by the problems caused by estimation error in investment portfolio optimization, this dissertation developed techniques to improve covariance forecasting, first by incorporating side information, then by projecting uncertainty in the values of parameters of factor-modeled covariance to uncertainty in covariance between securities.

The latter result made it possible to approach several investment problems explicitly modeling uncertainty — mean-variance portfolio optimization, risk parity portfolio construction, reporting risks under uncertainty and price movement. Algorithmic trading would be the next, and surely there are others.

At the same time, applications of the uncertain covariance model in [chapter 5](#) aren't restricted to quantitative finance. Covariance appears everywhere. Mean and variance are the first two terms of a Taylor series, and forecasting higher moments (particularly for the purpose of making decisions) is hard. There are surely many non-finance problems that could benefit from a tractable, predictive distribution of covariance.

BIBLIOGRAPHY

- Karim M Abadir, Walter Distaso, and Filip Žikeš. Design-free estimation of variance matrices. *Journal of Econometrics*, 181(2):165–180, 2014. ISSN 0304-4076. doi: 10.1016/j.jeconom.2014.03.010.
- Evan W Anderson and Ai-Ru Cheng. Robust Bayesian portfolio choices. *The Review of Financial Studies*, 29(5):1330–1375, 2016.
- Werner Antweiler and Murray Z Frank. Is all that talk just noise? The information content of internet stock message boards. *The Journal of Finance*, 59(3):1259–1294, 2004. doi: 10.1111/j.1540-6261.2004.00662.x.
- Adam Atkins, Mahesan Niranjana, and Enrico Gerding. Financial news predicts stock market volatility better than close price. *The Journal of Finance and Data Science*, 4(2):120–137, 2018.
- Doron Avramov and Guofu Zhou. Bayesian portfolio analysis. *Annual Review of Financial Economics*, 2(1):25–47, 2010. doi: 10.1146/annurev-financial-120209-133947.
- Xi Bai, Katya Scheinberg, and Reha Tutuncu. Least-squares approach to risk parity in portfolio selection. *Quantitative Finance*, 16(3):357–376, 2016. doi: 10.1080/14697688.2015.1031815.
- Ole E Barndorff-Nielsen and Neil Shephard. Non-Gaussian Ornstein–Uhlenbeck-based models and some of their uses in financial economics. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 63(2):167–241, 2001.
- Christopher B Barry. Portfolio analysis under uncertain means, variances, and covariances. *The Journal of Finance*, 29(2):515–522, 1974.
- David Bauder, Taras Bodnar, Nestor Parolya, and Wolfgang Schmid. Bayesian mean–variance analysis: optimal portfolio selection under parameter uncertainty. *Quantitative Finance*, 21(2):221–242, 2021.

- Luc Bauwens, Sébastien Laurent, and Jeroen VK Rombouts. Multivariate GARCH models: a survey. *Journal of Applied Econometrics*, 21(1):79–109, 2006.
- Stan Beckers. Variances of security price returns based on high, low, and closing prices. *The Journal of Business*, pages 97–112, 1983.
- Simone Bernardi, Markus Leippold, and Harald Lohre. Second-order risk of alternative risk parity strategies. *Journal of Risk*, 21(3):1–25, 2019. doi: 10.21314/jor.2018.401.
- Dimitris Bertsimas, David B Brown, and Constantine Caramanis. Theory and applications of robust optimization. *SIAM review*, 53(3):464–501, 2011. doi: 10.1137/080734510.
- Fischer Black and Robert Litterman. Global portfolio optimization. *Financial Analysts Journal*, 48(5):28–43, 1992. doi: 10.2469/faj.v48.n5.28.
- Tim Bollerslev. Generalized autoregressive conditional heteroskedasticity. *Journal of Econometrics*, 31(3):307–327, 1986.
- Tim Bollerslev. Modelling the coherence in short-run nominal exchange rates: a multivariate generalized ARCH model. *The Review of Economics and Statistics*, pages 498–505, 1990.
- Tim Bollerslev, Robert F. Engle, and Daniel B. Nelson. ARCH models. In *Handbook of Econometrics*, volume 4 of *Handbook of Econometrics*, chapter 49, pages 2959–3038. Elsevier, 1994. doi: 10.1016/S1573-4412(05)80018-2.
- Jean-Philippe Bouchaud. Radical complexity. *arXiv preprint arXiv:2103.09692*, 2021.
- Michael W Brandt and Christopher S Jones. Volatility forecasting with range-based EGARCH models. *Journal of Business & Economic Statistics*, 24(4):470–486, 2006.
- L Breiman. Optimal gambling systems for favorable games. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Contributions to the Theory of Statistics*. The Regents of the University of California, 1961.

- Mark Broadie. Computing efficient frontiers using estimated parameters. *Annals of Operations Research*, 45(1):21–58, 1993. doi: 10.1007/bf02282040.
- Joël Bun, Jean-Philippe Bouchaud, and Marc Potters. Cleaning large correlation matrices: tools from random matrix theory. *Physics Reports*, 666:1–109, 2017.
- Fabiano Cavadini, Alessandro Sbuelz, and Fabio Trojani. A simplified way of incorporating model risk, estimation risk and robustness in mean variance portfolio management. *Working paper, Tilburg University*, 2001.
- Sebastián Ceria and Robert A Stubbs. Incorporating estimation errors into portfolio selection: Robust portfolio construction. *Journal of Asset Management*, 7(2):109–127, 2006.
- Gary Chamberlain. A characterization of the distributions that imply meanvariance utility functions. *Journal of Economic Theory*, 29(1):185–201, 1983. ISSN 0022-0531. doi: 10.1016/0022-0531(83)90129-1.
- Tun-Jen Chang, Sang-Chin Yang, and Kuang-Jung Chang. Portfolio optimization problems in different risk measures using genetic algorithm. *Expert Systems with Applications*, 36(7):10529–10537, 2009. ISSN 0957-4174. doi: 10.1016/j.eswa.2009.02.062.
- James Chilson, Raymond Ng, Alan Wagner, and Ruben Zamar. Parallel computation of high-dimensional robust correlation and covariance matrices. *Algorithmica*, 45(3):403–431, 2006.
- Ray Yeutien Chou. Forecasting financial volatilities with extreme values: the conditional autoregressive range (carr) model. *Journal of Money, Credit and Banking*, pages 561–582, 2005.
- Ray Yeutien Chou, Chun-Chou Wu, and Nathan Liu. Forecasting time-varying covariance with a range-based dynamic conditional correlation model. *Review of Quantitative Finance and Accounting*, 33(4):327–345, 2009.
- Ray Yeutien Chou, Hengchih Chou, and Nathan Liu. Range volatility models and their applications in finance. In *Handbook of Quantitative Finance and Risk Management*, pages 1273–1281. Springer, 2010.

- Giorgio Costa. *Advances in Risk Parity Portfolio Optimization*. PhD thesis, Mechanical and Industrial Engineering, University of Toronto, 2021. <https://hdl.handle.net/1807/106376>.
- Giorgio Costa and Roy Kwon. A robust framework for risk parity portfolios. *Journal of Asset Management*, 21(5):447–466, 2020a.
- Giorgio Costa and Roy H Kwon. Risk parity portfolio optimization under a Markov regime-switching framework. *Quantitative Finance*, 19(3):453–471, 2019.
- Giorgio Costa and Roy H Kwon. Generalized risk parity portfolio optimization: an admm approach. *Journal of Global Optimization*, 78(1):207–238, 2020b.
- Peter Cotton. *Microprediction: Building an Open AI Network*. MIT Press, 2022a.
- Peter Cotton. Schur complementary portfolios – a unification of machine learning and optimization-based allocation, Nov 2022b. URL <https://medium.com/geekculture/schur-complementary-portfolios-fix-hierarchical-risk-parity-28b0efa1f35f>.
- Aditi Dandapani, Paul Jusselin, and Mathieu Rosenbaum. From quadratic Hawkes processes to super-Heston rough volatility models with Zumbach effect. *Quantitative Finance*, pages 1–13, 2021.
- Michael J Daniels and Robert E Kass. Shrinkage estimators for covariance matrices. *Biometrics*, 57(4):1173–1184, 2001.
- Sanjiv R Das and Mike Y Chen. Yahoo! for Amazon: Sentiment extraction from small talk on the web. *Management Science*, 53(9):1375–1388, 2007.
- Mark Davis and Sébastien Lleo. Black–litterman in continuous time: the case for filtering. *Quantitative Finance Letters*, 1(1):30–35, 2013.
- Marielle de Jong. Portfolio optimisation in an uncertain world. *Journal of Asset Management*, 19(4):216–221, 2018. doi: 10.1057/s41260-017-0066-3.
- Beatriz Vaz de Melo Mendes and Ricardo Pereira Câmara Leal. Robust multivariate modeling in finance. *International Journal of Managerial Finance*, 2005.

- Gianluca De Nard, Olivier Ledoit, and Michael Wolf. Factor models for portfolio selection in large dimensions: The good, the better and the ugly. *Journal of Financial Econometrics*, 01 2019. ISSN 1479-8409. doi: 10.1093/jfinec/nby033.
- Gianluca De Nard, Robert F Engle, Olivier Ledoit, and Michael Wolf. Large dynamic covariance matrices: enhancements based on intraday data. *University of Zurich, Department of Economics, Working Paper*, (356), 2021.
- Marcos Lopez De Prado. Building diversified portfolios that outperform out of sample. *The Journal of Portfolio Management*, 42(4):59–69, 2016.
- Victor DeMiguel and Francisco J Nogales. Portfolio selection with robust estimation. *Operations Research*, 57(3):560–577, 2009. doi: 10.1287/opre.1080.0566.
- Dan diBartolomeo and Sandy Warrick. Making covariance-based portfolio risk models sensitive to the rate at which markets reflect new information. In John Knight and Stephen Satchell, editors, *Linear Factor Models in Finance*, Quantitative Finance, chapter 12, pages 249 – 261. Butterworth-Heinemann, Oxford, 2005. ISBN 978-0-7506-6006-8. doi: 10.1016/B978-075066006-8.50013-6.
- Thomas Dimpfl and Stephan Jank. Can internet search queries help to predict stock market volatility? *European Financial Management*, 22(2):171–192, 2016.
- Robert Engle. Dynamic conditional correlation: A simple class of multivariate generalized autoregressive conditional heteroskedasticity models. *Journal of Business & Economic Statistics*, 20(3):339–350, 2002.
- Robert Engle. Risk and volatility: Econometric models and financial practice. *American Economic Review*, 94(3):405–420, 2004.
- Robert F Engle. Autoregressive conditional heteroscedasticity with estimates of the variance of united kingdom inflation. *Econometrica: Journal of the Econometric Society*, pages 987–1007, 1982. doi: 10.2307/1912773.
- Robert F. Engle and Kenneth F. Kroner. Multivariate simultaneous generalized ARCH. *Econometric Theory*, 11(1):122–150, 1995. ISSN 02664666, 14694360. URL <http://www.jstor.org/stable/3532933>.

- Robert F Engle, Victor K Ng, and Michael Rothschild. Asset pricing with a factor-ARCH covariance structure: Empirical estimates for treasury bills. *Journal of Econometrics*, 45(1-2):213–237, 1990.
- Robert F Engle, Olivier Ledoit, and Michael Wolf. Large dynamic covariance matrices. *Journal of Business & Economic Statistics*, 37(2):363–375, 2019. doi: 10.1080/07350015.2017.1345683.
- Frank J Fabozzi, Petter N Kolm, Dessislava A Pachamanova, and Sergio M Focardi. *Robust Portfolio Optimization and Management*. John Wiley & Sons, 2007.
- Frank J Fabozzi, Dashan Huang, and Guofu Zhou. Robust portfolios: contributions from operations research and finance. *Annals of Operations Research*, 176(1):191–220, 2010.
- Jianqing Fan, Yingying Fan, and Jinchi Lv. High dimensional covariance matrix estimation using a factor model. *Journal of Econometrics*, 147(1):186–197, 2008.
- Jianqing Fan, Yuan Liao, and Martina Mincheva. Large covariance estimation by thresholding principal orthogonal complements. *Journal of the Royal Statistical Society. Series B, Statistical methodology*, 75(4), 2013. doi: 10.1111/rssb.12016.
- William Feller. The asymptotic distribution of the range of sums of independent random variables. *The Annals of Mathematical Statistics*, 22(3):427–432, 1951.
- Peter A Frost and James E Savarino. An empirical bayes approach to efficient portfolio selection. *Journal of Financial and Quantitative Analysis*, 21(03):293–305, 1986. doi: 10.2307/2331043.
- Lorenzo Garlappi, Raman Uppal, and Tan Wang. Portfolio selection with parameter and model uncertainty: A multi-prior approach. *The Review of Financial Studies*, 20(1):41–81, 2007.
- Nicolae Gârleanu and Lasse Heje Pedersen. Dynamic trading with predictable returns and transaction costs. *The Journal of Finance*, 68(6):2309–2340, 2013.
- Mark B Garman and Michael J Klass. On the estimation of security price volatilities from historical data. *The Journal of Business*, pages 67–78, 1980.

- Lisa R Goldberg, Alex Papanicolaou, and Alexander Shkolnik. The dispersion bias. *Available at SSRN 3071328*, 2018. doi: 10.2139/ssrn.3071328.
- Donald Goldfarb and Garud Iyengar. Robust portfolio selection problems. *Mathematics of Operations Research*, 28(1):1–38, 2003. doi: 10.1287/moor.28.1.1.14260.
- Bjarni V Halldórsson and Reha H Tütüncü. An interior-point method for a class of saddle-point problems. *Journal of Optimization Theory and Applications*, 116(3): 559–590, 2003. doi: 10.1023/a:1023065319772.
- Lars Peter Hansen and Thomas J Sargent. *Robustness*. Princeton University Press, 2011.
- Peter R Hansen and Asger Lunde. A forecast comparison of volatility models: does anything beat a GARCH(1, 1)? *Journal of Applied Econometrics*, 20(7):873–889, 2005.
- Peter Reinhard Hansen, Zhuo Huang, and Howard Howan Shek. Realized GARCH: a joint model for returns and realized measures of volatility. *Journal of Applied Econometrics*, 27(6):877–906, 2012.
- Peter Huber. *Robust statistics*. John Wiley & Sons, 1981. doi: 10.1002/0471725250.
- Leon Isserlis. On a formula for the product-moment coefficient of any order of a normal frequency distribution in any number of variables. *Biometrika*, 12(1/2): 134–139, 1918.
- J David Jobson and Bob Korkie. Estimation for markowitz efficient portfolios. *Journal of the American Statistical Association*, 75(371):544–554, 1980. doi: 10.2307/2287643.
- Iain M Johnstone. High dimensional statistical inference and random matrices. In *Proceedings of the International Congress of Mathematicians Madrid, August 22–30, 2006*, pages 307–333, 2007.
- Philippe Jorion. Bayes-Stein estimation for portfolio analysis. *Journal of Financial and Quantitative Analysis*, 21(03):279–292, 1986. doi: 10.2307/2331042.

- Philippe Jorion. Risk management lessons from Long-Term Capital Management. *European financial management*, 6(3):277–300, 2000. doi: 10.1111/1468-036X.00125.
- Marcin Kacperczyk and Paul Damien. Asset allocation under distribution uncertainty. *McCombs Research Paper Series No. IROM-01*, 11, 2011.
- Petko S Kalev, Wai-Man Liu, Peter K Pham, and Elvis Jarnecic. Public information arrival and volatility of intraday stock returns. *Journal of Banking & Finance*, 28(6):1441–1467, 2004.
- Michalis Kapsos, Nicos Christofides, and Berc Rustem. Robust risk budgeting. *Annals of Operations Research*, 266(1):199–221, 2018.
- Ioannis Karatzas and Steven Shreve. *Brownian Motion and Stochastic Calculus*, volume 113. Springer, 1991.
- J. Kelly. A new interpretation of information rate. *IRE Transactions on Information Theory*, 2(3):185–189, 1956. doi: 10.1109/TIT.1956.1056803.
- Jafar A Khan, Stefan Van Aelst, and Ruben H Zamar. Robust linear model selection based on least angle regression. *Journal of the American Statistical Association*, 102(480):1289–1299, 2007.
- Hyuksoo Kim and Saejoon Kim. Reduction of estimation error impact in the risk parity strategies. *Quantitative Finance*, pages 1–14, 2021.
- Seung-Jean Kim and Stephen Boyd. Robust efficient frontier analysis with a separable uncertainty model. Technical report, Technical report, Department of Electrical Engineering, Stanford University, 2007.
- Roger W Klein and Vijay S Bawa. The effect of estimation risk on optimal portfolio choice. *Journal of Financial Economics*, 3(3):215–231, 1976.
- Petter Kolm and Gordon Ritter. On the Bayesian interpretation of Black-Litterman. *European Journal of Operational Research*, 258(2):564–572, 2017.
- Petter Kolm and Gordon Ritter. Factor investing with Black-Litterman-Bayes: Incorporating factor views and priors in portfolio construction. *The Journal of Portfolio Management*, 47(2):113–126, 2020.

- Petter Kolm, Gordon Ritter, and Joseph Simonian. Black-Litterman and beyond: The Bayesian paradigm in investment management. *The Journal of Portfolio Management*, 47(5):91–113, 2021.
- Hiroshi Konno and Hiroaki Yamazaki. Mean-absolute deviation portfolio optimization model and its applications to tokyo stock market. *Management Science*, 37(5):519–531, 1991.
- Jaya Krishnakumar and Elvezio Ronchetti. Robust estimators for simultaneous equations models. *Journal of Econometrics*, 78(1):295–314, 1997.
- Naoto Kunitomo. Improving the Parkinson method of estimating security price volatilities. *The Journal of Business*, pages 295–302, 1992.
- Tze Leung Lai, Haipeng Xing, and Zehao Chen. Mean–variance portfolio optimization when means and covariances are unknown. *The Annals of Applied Statistics*, 5(2A):798–823, 2011.
- Clifford Lam. Nonparametric eigenvalue-regularized precision or covariance matrix estimator. *The Annals of Statistics*, 44(3):928–953, 2016. doi: 10.1214/15-AOS1393.
- Nathan Lassance, Victor DeMiguel, and Frédéric Vrans. Optimal portfolio diversification via independent component analysis. *Operations Research*, 70(1):55–72, 2022. doi: 10.1287/opre.2021.2140.
- GJ Lauprete, AM Samarov, and RE Welsch. Robust portfolio optimization. In *Developments in Robust Statistics*, pages 235–245. Springer, 2003.
- Olivier Ledoit and Michael Wolf. Improved estimation of the covariance matrix of stock returns with an application to portfolio selection. *Journal of Empirical Finance*, 10(5):603–621, 2003. doi: 10.1016/S0927-5398(03)00007-0.
- Olivier Ledoit and Michael Wolf. Honey, i shrunk the sample covariance matrix. *The Journal of Portfolio Management*, 30(4):110–119, 2004. doi: 10.3905/jpm.2004.110.
- Olivier Ledoit and Michael Wolf. Nonlinear shrinkage estimation of large-dimensional covariance matrices. *The Annals of Statistics*, 40(2):1024–1060, 2012.

- Olivier Ledoit and Michael Wolf. Nonlinear shrinkage of the covariance matrix for portfolio selection: Markowitz meets Goldilocks. *The Review of Financial Studies*, 30(12):4349–4388, 06 2017. ISSN 0893-9454. doi: 10.1093/rfs/hhx052.
- Pui Lam Leung and Foon Yip Ng. Improved estimation of a covariance matrix in an elliptically contoured matrix distribution. *Journal of Multivariate Analysis*, 88(1): 131–137, 2004.
- Harald Lohre, Heiko Opfer, and Gabor Orszag. Diversifying risk parity. *Journal of Risk*, 16(5):53–79, 2014. doi: 10.21314/jor.2014.284.
- Zhaosong Lu. Robust portfolio selection based on a joint ellipsoidal uncertainty set. *Optimization Methods & Software*, 26(1):89–104, 2011. doi: 10.1080/10556780903334682.
- Sébastien Maillard, Thierry Roncalli, and Jérôme Teïletche. The properties of equally weighted risk contribution portfolios. *The Journal of Portfolio Management*, 36(4): 60–70, 2010. doi: 10.3905/jpm.2010.36.4.060.
- Benoit Mandelbrot. The variation of certain speculative prices. *The Journal of Business*, 36(4):394–419, 1963.
- Harry Markowitz. Portfolio selection. *The Journal of Finance*, 7(1):77–91, 1952. doi: 10.1111/j.1540-6261.1952.tb01525.x.
- Harry Markowitz. Mean–variance approximations to expected utility. *European Journal of Operational Research*, 234(2):346–355, 2014.
- Harry M Markowitz. *Portfolio Selection*. Yale University Press, jul 1959. ISBN 9780300013726.
- Ricardo A Maronna and Ruben H Zamar. Robust estimates of location and dispersion for high-dimensional datasets. *Technometrics*, 44(4):307–317, 2002.
- Jose Menchero and Ben Davis. Risk contribution is exposure times volatility times correlation: Decomposing risk using the x-sigma-rho formula. *The Journal of Portfolio Management*, 37(2):97–106, 2011. ISSN 0095-4918. doi: 10.3905/jpm.2011.37.2.097.

- Robert C Merton. Optimum consumption and portfolio rules in a continuous-time model. *Journal of Economic Theory*, 3(4):373–413, 1971. ISSN 0022-0531. doi: 10.1016/0022-0531(71)90038-X.
- Attilio Meucci. Fully flexible views: Theory and practice. *Risk*, 21(10):97–102, 2008.
- Attilio Meucci. Review of statistical arbitrage, cointegration, and multivariate Ornstein-Uhlenbeck. *SSRN*, 2009. doi: 10.2139/ssrn.1404905.
- Attilio Meucci. 'P' versus 'Q': Differences and commonalities between the two areas of quantitative finance. *GARP Risk Professional*, pages 47–50, 2 2011. Available at SSRN: <https://ssrn.com/abstract=1717163>.
- Richard O. Michaud. The markowitz optimization enigma: Is 'optimized' optimal? *Financial Analysts Journal*, 45(1):31–42, 1989. doi: 10.2469/faj.v45.n1.31.
- Richard O. Michaud. *Efficient Asset Management: A Practical Guide to Stock Portfolio Optimization and Asset Allocation*. Oxford University Press, 1998.
- Luiza Miranyan. Incorporating forward-looking market data into linear multifactor fundamental models. *Journal of Risk*, 14(4):3–34, 2012. doi: 10.21314/JOR.2012.245.
- Leela Mitra, Gautam Mitra, and Dan Dibartolomeo. Equity portfolio risk estimation using market information and sentiment. *Quantitative Finance*, 9(8):887–895, 2009. doi: 10.1080/14697680903448361.
- Colm A. O’Cinneide. Risk contributions: duality and sensitivity. *Quantitative Finance*, 18(12):2023–2033, 2018. doi: 10.1080/14697688.2017.1423175.
- Michael Parkinson. The extreme value method for estimating the variance of the rate of return. *The Journal of Business*, 53(1):61–65, 1980. doi: 10.1086/296071.
- Markus Pelger. Deep learning statistical arbitrage. NYU Courant, M.S. in Mathematics in Finance Seminars, Feb 2023.
- Cdric Perret-Gentil and Maria-Pia Victoria-Feser. Robust mean-variance portfolio selection. fame research paper 140. *International Center for Financial Asset Management and Engineering, Geneva*, 2004.

- Vasiliki Plerou, Parameswaran Gopikrishnan, Bernd Rosenow, Luis A Nunes Amaral, Thomas Guhr, and H Eugene Stanley. Random matrix approach to cross correlations in financial data. *Physical Review E*, 65(6):066126, 2002.
- Marc Potters and Jean-Philippe Bouchaud. *A First Course in Random Matrix Theory: For Physicists, Engineers and Data Scientists*. Cambridge University Press, 2020.
- Marc Potters, Jean-Philippe Bouchaud, and Laurent Laloux. Financial applications of random matrix theory: Old laces and new pieces. *Acta Physica Polonica B*, 36(9):2767, 2005.
- Edward Qian. Risk parity portfolios: Efficient portfolios through true diversification. *Panagora Asset Management*, 2005.
- Edward Qian. Risk parity and diversification. *The Journal of Investing*, 20(1):119–127, 2011. ISSN 1068-0896. doi: 10.3905/joi.2011.20.1.119.
- Edward E Qian, Ronald H Hua, and Eric H Sorensen. *Quantitative equity portfolio management: modern techniques and applications*. CRC Press, 2007. doi: 10.1201/9781420010794.
- Calum Robertson, Shlomo Geva, and Rodney C Wolff. Can the content of public news be used to forecast abnormal stock market behaviour? In *Seventh IEEE International Conference on Data Mining (ICDM 2007)*, pages 637–642. IEEE, 2007.
- R Tyrrell Rockafellar and Stanislav Uryasev. Optimization of conditional value-at-risk. *Journal of Risk*, 2(3):21–41, 2000. ISSN 1465-1211. doi: 10.21314/JOR.2000.038.
- R Tyrrell Rockafellar and Stanislav Uryasev. Conditional value-at-risk for general loss distributions. *Journal of Banking & Finance*, 26(7):1443–1471, 2002.
- David M Rocke. Robustness properties of S-estimators of multivariate location and shape in high dimension. *The Annals of Statistics*, pages 1327–1345, 1996.

- L Christopher G Rogers and Stephen E Satchell. Estimating variance from high, low and closing prices. *The Annals of Applied Probability*, pages 504–512, 1991.
- Leonard CG Rogers and Fanyin Zhou. Estimating correlation from high, low, opening and closing prices. *The Annals of Applied Probability*, 18(2):813–823, 2008.
- Thierry Roncalli and Guillaume Weisang. Risk parity portfolios with risk factors. *Quantitative Finance*, 16(3):377–388, 2016. doi: 10.1080/14697688.2015.1046907.
- Peter J Rousseeuw. Multivariate estimation with high breakdown point. In W Grossman, G Pflug, I Vincze, and W Wertz, editors, *Mathematical Statistics and Applications*, volume 8, pages 283–297. Reidel Publishing Company, 1985.
- Peter J Rousseeuw and Katrien Van Driessen. A fast algorithm for the minimum covariance determinant estimator. *Technometrics*, 41(3):212–223, 1999.
- Berc Rustem, Robin G Becker, and Wolfgang Marty. Robust min–max portfolio strategies for rival forecast and risk scenarios. *Journal of Economic Dynamics and Control*, 24(11-12):1591–1621, 2000.
- Juliane Schäfer and Korbinian Strimmer. A shrinkage approach to large-scale covariance matrix estimation and implications for functional genomics. *Statistical Applications in Genetics and Molecular Biology*, 4(1), 2005.
- Bernd Scherer. Portfolio resampling: Review and critique. *Financial Analysts Journal*, 58(6):98–109, 2002. doi: 10.2469/faj.v58.n6.2489.
- Maria Grazia Scutellà and Raffaella Recchia. Robust portfolio asset allocation and risk measures. *Annals of Operations Research*, 204(1):145–169, 2013. doi: 10.1007/s10479-012-1266-3.
- Anish Shah. Short-term risk from long-term models. Presented at Northfield Information Services Research Conference in Key Largo FL, 10 2007. URL: <http://northinfo.com/documents/286.pdf>.
- Anish Shah. Mitigating estimation error in optimization. Presented at Northfield Information Services seminar in London, 11 2009. URL: <https://www.northinfo.com/documents/361.pdf>.

- Anish Shah. Uncertain risk parity. *Journal of Investment Strategies*, 10(1):29–41, 2021a. doi: 10.21314/JOIS.2021.009.
- Anish Shah. Risk decomposition under uncertainty and price movement. CFA Society VBA Netherlands, Feb 2021b.
- Anish Shah. Risk under uncertainty. Qwafaxnew NYC, Mar 2021c.
- Anish Shah. Risk under uncertainty and price movement. London Quant Group, Mar 2022.
- Anish R Shah. Short-term risk and adapting covariance models to current market conditions. *SSRN*, 2014. doi: 10.2139/ssrn.2501071.
- Anish R Shah. Uncertain covariance models and uncertainty-penalized portfolio optimization. *SSRN*, 2015. doi: 10.2139/ssrn.2616109.
- Anish R Shah. Portfolio risks under estimation uncertainty and price movement. *SSRN*, 2020. doi: 10.2139/ssrn.3774239.
- Annastiina Silvennoinen and Timo Teräsvirta. *Multivariate GARCH Models*, pages 201–229. Springer Berlin Heidelberg, Berlin, Heidelberg, 2009. ISBN 978-3-540-71297-8. doi: 10.1007/978-3-540-71297-8_9.
- Charles Stein. Inadmissibility of the usual estimator for the mean of a multivariate normal distribution. In *Proceedings of the Third Berkeley Symposium on Mathematical Statistics and Probability*, volume 1, pages 197–206, 1956.
- Nassim Nicholas Taleb. *Antifragile: Things that gain from disorder*, page 425. Random House Incorporated, 2012. ISBN 9780812979688. appendix II, table 12.
- Edward O. Thorp. Portfolio choice and the Kelly criterion. *Proceedings of the Business and Economics Section of the American Statistical Association*, pages 215–224, 1971.
- Ruey S Tsay. *Analysis of Financial Time Series*. John Wiley & Sons, 2010.

- Pauli Virtanen, Ralf Gommers, Travis E. Oliphant, Matt Haberland, Tyler Reddy, David Cournapeau, Evgeni Burovski, Pearu Peterson, Warren Weckesser, Jonathan Bright, Stéfan J. van der Walt, Matthew Brett, Joshua Wilson, K. Jarrod Millman, Nikolay Mayorov, Andrew R. J. Nelson, Eric Jones, Robert Kern, Eric Larson, C J Carey, İlhan Polat, Yu Feng, Eric W. Moore, Jake VanderPlas, Denis Laxalde, Josef Perktold, Robert Cimrman, Ian Henriksen, E. A. Quintero, Charles R. Harris, Anne M. Archibald, Antônio H. Ribeiro, Fabian Pedregosa, Paul van Mulbregt, and SciPy 1.0 Contributors. SciPy 1.0: Fundamental Algorithms for Scientific Computing in Python. *Nature Methods*, 17:261–272, 2020. doi: 10.1038/s41592-019-0686-2.
- Ioannis D Vrontos, Petros Dellaportas, and Dimitris N Politis. A full-factor multivariate GARCH model. *The Econometrics Journal*, 6(2):312–334, 2003.
- Roy E Welsch and Xinfeng Zhou. Application of robust statistics to asset allocation models. *REVSTAT–Statistical Journal*, 5(1):97–114, 2007.
- Dennis Yang and Qiang Zhang. Drift-independent volatility estimation based on high, low, open, and close prices. *The Journal of Business*, 73(3):477–492, 2000.
- Liusha Yang, Romain Couillet, and Matthew R McKay. A robust statistics approach to minimum variance portfolio optimization. *IEEE Transactions on Signal Processing*, 63(24):6684–6697, 2015. doi: 10.1109/TSP.2015.2474298.
- Chao Zhang, Yihuang Zhang, Mihai Cucuringu, and Zhongmin Qian. Volatility forecasting with machine learning and intraday commonality. *Journal of Financial Econometrics*, 03 2023. ISSN 1479-8409. doi: 10.1093/jjfinec/nbad005.