

UNIVERSITY OF TARTU
School of Economics and Business Administration

Maria-Britta Sütt

DEVELOPING A TWO-STAGE PORTFOLIO CONSTRUCTION MODEL WITH
MEAN ABSOLUTE DEVIATION SCREENING AND CONDITIONAL
VALUE-AT-RISK-CONSTRAINED ENTROPY WEIGHTING

Master's Thesis

Supervisor: lecturer Mark Kantšukov (PhD)

Tartu 2026

This thesis is an independent body of work. Third-party articles and other literature, opinions, and data are credited to their respective authors.

Table of Contents

Abstract	4
Introduction	4
1. Literature Review	7
1.1. Evolution of Portfolio Construction Methods	7
1.2. Approaches to Risk Optimisation in Portfolio Construction	12
1.3. Asset Allocation Under Portfolio Cardinality Constraints	16
2. Methodology and Data	20
2.1. Methodology	20
2.2. Data	25
3. Results and Discussion	28
3.1. ETF Universe Simulation Results	28
3.2. Technology Universe Simulation Results	37
3.3. Discussion	44
Conclusions	50
References	53
Appendices	57
Appendix A: Coherence Axioms for Risk Measures	57
Appendix B: Sharpe, Sortino, and Information Ratio Definitions	58
Appendix C: Paired <i>t</i> -tests and Paired Wilcoxon Signed-Rank Test Results . .	59
Appendix D: MCE Portfolio Cardinality Sensitivity Tests	61
Appendix E: Assets Used in Backtesting	63
Appendix F: January 2026 Technology Universe Equal-Weight And MAD-CVaR- Entropy Portfolio Constituent Assets	68
Appendix G: GitHub Repository	69
Resümee	70
Non-exclusive License for Reproduction and Public Access	72

Abstract

This thesis develops and tests a two-stage portfolio construction model that combines mean return and mean absolute deviation-based screening with entropy maximisation under a conditional value-at-risk constraint. Findings indicate that portfolio performance depends materially on the chosen asset universe. The proposed model achieved its aim of a balanced outcome between return and risk, while being simpler to implement than mean-variance. The model was backtested out-of-sample against an equal-weighted portfolio and two mean-variance models on a broad U.S. ETF universe in 39 semi-annual holding periods, and on U.S. technology shares in 21 semi-annual holding periods. In the ETF universe, the proposed model achieved the highest cumulative return, 384.5%, among the compared portfolios, with the equal-weight portfolio coming in second best at 337.4%. Mean tail-loss risk, measured in conditional value-at-risk terms, was also lower for the developed model: 1.9% vs equal-weight portfolio's 2.4%. The model's mean turnover was lower at 22% vs equal-weight portfolio's 35%. In the technology universe, no model dominated across return, tail-risk, and turnover metrics. Equal weighting captured the strongest cumulative return at 697% but also the highest mean tail risk at 4.6%, while the proposed model remained more balanced overall with 315.5% cumulative return and 3% mean tail-loss risk. In the technology universe, when compared specifically with a mean-variance portfolio constructed from the same assets, the proposed model achieved more than twice the cumulative return (315.5% vs 153%), with almost identical mean tail-loss risk (3% for both) and lower mean turnover (67% vs 83%).

Introduction

Investment portfolio construction has been a well-studied field since the 1950s with the inception of Harry Markowitz's modern portfolio theory. Research has tracked technological progress in computing with increasingly complex lines of inquiry into data and methods. Seventy years later, investing remains fundamentally the same balancing act when deconstructed to first principles: expected returns are traded off against expected risks. Generally, portfolio construction entails two steps (Markowitz, 1959): selecting assets, and then selecting weights. Modern portfolio construction employs a

broad range of quantitative optimisation methods to find efficient portfolio allocations under constraints (Kim *et al.*, 2024). However, investors frequently rely on simple decision rules and summary indicators rather than complex mathematical models when making investment decisions, because implementation feasibility is not a trivial consideration in practice (Platanakis *et al.*, 2021).

Crucially, it has been seen that more complex models do not automatically produce better portfolios. The Markowitz mean-variance model (MV) has been shown to be highly sensitive to input errors, and adding practical constraints to optimisation can still introduce computational difficulties (Palomar, 2025). Solutions that have been shown to enhance MV's out-of-sample stability, like covariance shrinkage and factor models with varying levels of sophistication (De Nard *et al.*, 2021), and robust optimisation methods (Xidonas *et al.*, 2020), extend implementation complexity significantly. At the same time, these extensions have seen only limited improvements in realised performance, and most of these gains have been in risk metrics rather than in returns. All the while, naïve equal-weight (EW) portfolios remain difficult to beat in both robustness and returns (DeMiguel *et al.*, 2009). However, EW lacks risk controls as a baseline. Portfolio size is a real practical consideration but limiting the number of active positions is not necessarily straightforward. Cardinality restriction has generated research into increasingly advanced optimisation methods, but these add further implementation complexity and may require substantial specialist knowledge. During review, the author was not able to find any studies to address all of the aforementioned concerns about out-of-sample stability, return performance, and portfolio cardinality at the same time.

The aim of this thesis is to develop a portfolio construction model that relies on mean absolute deviation (MAD), entropy, and conditional value-at-risk (CVaR). The substantial implementational complexity of MV models with methods to improve out-of-sample stability has been seen to be disproportionately greater than the limited empirical improvements over EW, particularly in realised returns. The proposed model addresses these disadvantages in a two-stage approach. The purpose of the first stage is to eliminate the need for cardinality optimisation entirely by ranking and choosing an exact number of eligible assets for the final portfolio. The stage uses return and mean absolute deviation (MAD) as simple indicators for asset screening to capture recent his-

torical upside and some expected volatility risk. The second stage optimises portfolio weights between the selected assets and it does so without using covariances. Instead, allocations are diversified by maximising entropy under a capped conditional value-at-risk (CVaR) expected tail-loss risk control target. As an additional feature, the model applies exponential time weighting in both stages to give more weight to more recent observations.

The model is evaluated out-of-sample by using historical market data and comparing performance against three alternative portfolios: two mean-variance (MV) models with different cardinality, and EW with the same cardinality as the proposed model, but a different asset selection rule. S&P 500 is further included as a market proxy benchmark. Model performance is compared and contrasted on two different U.S.-domiciled asset universes: a broad-coverage set of 72 ETFs, and 371 individual mid-to-large-cap technology sector shares. Backtesting uses 12-month rolling estimation windows and semi-annual rebalancing.

The thesis is structured as follows. Chapter 1 reviews existing literature based on the relevancy of studies for the topic of this thesis, the quality of study design, and recency. The review focuses on three areas. It starts with the evolution of portfolio construction methods from the original Markowitz mean-variance model to later developments addressing MV's various limitations and the recent resurgence of interest in EW portfolios. The second area deals with approaches to risk optimisation, and the third is concerned with asset weighting and diversification within portfolios. The literature review concludes with identifying a gap in existing research that this thesis aims to address. Chapter 2 presents the methodology of the proposed model design and the backtesting protocol and introduces testing data. Chapter 3 reports empirical findings and discusses results. The final chapter of the thesis presents conclusions and study limitations, and gives suggestions for further research.

Keywords: portfolio optimisation, portfolio investments, portfolio theory, investment portfolio.

CERCS classification code: S181 financial science.

1. Literature Review

1.1. Evolution of Portfolio Construction Methods

Portfolio construction is a multidisciplinary area that spans finance, operations research, and computer science (Kim *et al.*, 2024). The process originates from the modern portfolio theory framework (Markowitz, 1952) and it broadly decomposes into two consecutive activities: (i) data modelling and (ii) optimisation (Palomar, 2025).

Markowitz (1952) suggests that rational investors choose between portfolios that trade off two objectives: expected returns and variance of returns. The fundamental principles have not changed in the last seventy years; expectations and trade-offs still form the backbone of portfolio construction (Palomar, 2025). The Markowitz mean-variance model (MV) uses a covariance matrix of asset returns to balance expected return against return variance (Markowitz, 1952). The two main caveats of this particular approach are that variance penalises upside and downside deviations equally, and estimating high-dimensional covariance matrices can be computationally infeasible in practice. MV has since been widely shown to have poor out-of-sample performance from estimation errors, with even small input changes causing large differences in portfolio allocations output by the optimiser (Palomar, 2025). Extensive research has been devoted to improving MV estimation stability and out-of-sample performance by incorporating increasingly more prior information into portfolio construction models (Palomar, 2025).

Sharpe (1963) achieved similar allocations to original MV but did not assess out-of-sample model performance. Ledoit and Wolf (2004a, 2004b, 2017) improved estimation, using shrinkage to transform the sample covariance matrix to a more stable regularised model, and managed to improve the out-of-sample stability of minimum-variance MV compared to market index and equal-weights (EW) portfolio benchmarks. Covariance shrinkage is empirically supported but methodically complex with no single known best solution; the application of these methods is not trivial (Palomar, 2025).

Factor models—starting with Sharpe’s (1963) single-factor approach—are a standard method for incorporating prior information, e.g. structural market and asset information, into covariance estimation, but doing this increases data modelling complexity even further (Palomar, 2025; Ledoit and Wolf, 2004b). De Nard *et al.* (2021)

found that including multiple factors into a model injects additional uncertainty and can thereby worsen performance outcomes instead of improving them. Furthermore, De Nard *et al.* (2021:243) remark that there is “little consensus in the literature about the nature and the number” of factors to be used. As Palomar (2025) notes, the practical implication is that practitioners typically either buy factors from data providers at a high cost, or use portfolio designs that do not require estimating expected returns at all.

One approach that emerged as a response to parameter uncertainty was robust optimisation (Palomar, 2025). Robust optimisation treats parameter inputs as uncertain and accounts for estimation errors by using uncertainty sets instead of fixed point-estimate forecasts (Xidonas *et al.*, 2020). However, there is a lack of studies on the empirical performance of robust portfolio optimisation methods, and Xidonas *et al.* (2020) suggest that investors should not expect robust optimisation to perform better than classic methods, especially if estimation errors have little impact. Compared to nominal versions, robust models have been shown to produce similar portfolio compositions (Georgantas *et al.*, 2024), while improving out-of-sample risk-adjusted performance (Sehgal and Mehra, 2021; Georgantas *et al.*, 2024). This has been shown to be the case especially in crisis periods (Georgantas *et al.*, 2024) and in reducing tail-losses (Kim *et al.*, 2018). These studies suggest that robust optimisation may add most value to the investor when risk aversion is the main concern in the risk-reward equation.

While input regularisation and robust optimisation have led to more stable risk-related outcomes in MV’s out-of-sample performance, from the perspective of returns, equal-weight portfolios still appear to be difficult to outperform (Palomar, 2025). EW divides the investable sum equally between all selected assets. EW does not rely on estimating asset returns and does not require any mathematical optimisation (DeMiguel *et al.*, 2009). DeMiguel *et al.* (2009) compare 14 MV variants with various estimation error reduction methods to an EW benchmark. They evaluate U.S. and global index and market portfolios constructed from 7 datasets ranging between 1963 and 2004, using rolling estimation windows and monthly rebalancing. In the study, none of the optimised strategies were found to be consistently better out-of-sample than EW in terms of Sharpe ratios, certainty-equivalent returns (CEQ), or portfolio turnover (DeMiguel *et al.*, 2009).

There is further evidence in support of EW's superior returns performance. Plyakha *et al.* (2017) compared S&P 500-constituent EW portfolios against value-weighted (VW), and price-weighted (PW) portfolios using monthly rebalancing cycles over 515 months, and found that EW beat both VW and PW on total returns and risk-adjusted returns (Sharpe, Sortino, Treynor ratios, CEQ), but EW had higher volatility and turnover. Plyakha *et al.* (2017) conclude that EW's edge in their study appears to come from a specific source: monthly rebalancing and high turnover. These characteristics expose the portfolio more to market, size, and value factors, and maintaining equal weights at rebalancing is essentially a contrarian strategy that benefits from mean reversion. As a caveat, Plyakha *et al.*'s (2017) portfolios ranged from 30-300 assets, which towards the higher end especially is likely to diversify out a lot of individual asset-level noise. However, large portfolios can be impractical for real investors to hold.

Kritzman *et al.* (2010) agree with the arguments on concentration avoidance, contrarian rebalancing and reversion effects in support of EW, but criticise foregoing optimisation entirely. Kritzman *et al.* (2010) show that the choice of expected returns in optimiser inputs can produce portfolios with better out-of-sample Sharpe ratios than unoptimised EW. This claim does not invalidate EW but it does suggest that EW's advantage is sensitive to how beliefs about returns are formed and imposed at the data modelling stage: essentially, that asset selection prior to portfolio construction is important. This is not a new claim. Markowitz (1959) clearly stated the importance of asset selection as a step that should precede mathematical optimisation.

This progression in the literature is summarised in Table 1. The table groups the main research streams discussed in this subchapter and outlines their core characteristics and principal empirical findings.

Table 1

Summary of research streams in the last seventy years of portfolio construction methods evolution

Research Stream	Core Characteristics	Main Empirical Findings
Modern portfolio theory. 1952 onwards (Markowitz, 1952).	Investment portfolio as a balanced, diversified whole that trades off expected return and variance under uncertainty. Investor selects assets, then positions are allocated using mean-variance (MV) optimisation. Optimisation allocates weights by combining expected returns with an estimated covariance matrix, and then solving under constraints of (a) full investment, and (b) no short selling.	Utility function-satisfying efficient portfolios can be constructed with mathematical optimisation methods (Markowitz, 1952). Early research acknowledges weaknesses of MV (Markowitz, 1952; Markowitz, 1959): variance penalises upside and downside deviations equally. Estimating high-dimensional covariance matrices can be computationally infeasible.
MV estimation-error reduction. Initiating studies: mid-1960s onwards (Sharpe, 1963), late 1970s onwards (DeMiguel <i>et al.</i> , 2009), early 2000s onwards (De Nard <i>et al.</i> , 2021).	MV found to be input-sensitive; estimation error compounds and drives weak out-of-sample performance. Estimation-error reduction proposed via covariance adjustments: • simplification via structure imposed through single- and multi-factor models, • shrinkage in order to reduce the impact of extreme values.	Earliest one-factor model by Sharpe (1963) improved computational burden of mathematical optimisation while achieving similar allocations to original MV efficient frontier portfolios. Shrinkage improves MV out-of-sample stability (Ledoit and Wolf, 2004a; Ledoit and Wolf, 2004b; Ledoit and Wolf, 2017). Choice of shrinkage factor matters greatly (Ledoit and Wolf, 2004a; Ledoit and Wolf, 2004b). No consensus on best number and nature of factors to include; additional factors can worsen results (De Nard <i>et al.</i> , 2021). Implementation of factor models is complex, expensive in practice (Palomar, 2025).
Robust optimisation Early 2000s onwards (Fabozzi <i>et al.</i> , 2010).	Adopted to portfolio construction from mathematical programming. Aims to improve MV out-of-sample performance. Unlike classic MV, treats optimisation inputs as non-deterministic and estimated with errors. Optimiser inputs are formulated as uncertainty sets instead of fixed point-estimates.	Uncertainty set definition requires judgement (Xidonas <i>et al.</i> , 2020). Implementation can be complex in practice (Moon and Yao, 2011). Out-of-sample improvement shown over MV in risk-adjusted metrics, especially in crisis periods (Georgantas <i>et al.</i> , 2024; Kim <i>et al.</i> , 2018; Sehgal and Mehra, 2021). Comparative literature on robust optimisation methods is still not comprehensive (Kim <i>et al.</i> , 2018; Georgantas <i>et al.</i> , 2024). Research focuses on risk rather than returns (Xidonas <i>et al.</i> , 2020).

Research Stream	Core Characteristics	Main Empirical Findings
Equal-weight portfolios. 2009 onwards (DeMiguel <i>et al.</i> , 2009).	EW portfolios do not use mathematical optimisation to allocate positions. Shown to regularly outperform MV out-of-sample in returns terms. Lack downside risk controls. Recent rise in research interest due to strong empirical return performance.	Seen to achieve better out-of-sample risk-adjusted returns, but volatility is higher than in MV portfolios (DeMiguel <i>et al.</i> , 2009; Plyakha <i>et al.</i> , 2017). Advantage seems to stem from exposure to market, size, and value factors, and maintaining equal asset weights at each rebalance appears to benefit from mean reversion (Plyakha <i>et al.</i> , 2017; Kritzman <i>et al.</i> , 2010). Outperformance attributed to poor MV optimisation inputs rather than model superiority (Kritzman <i>et al.</i> , 2010).

Author: compiled by the author.

In conclusion, portfolio construction models and the methods employed within these models have evolved from simple—the Talmudic equal-weights rule goes back to ancient Babylon (Palomar, 2025)—to increasingly complex ways of incorporating prior information with an aim of achieving return-risk profiles that are optimal for investors. In the 1950s and 60s, technology constraints forced research to focus on making optimisation problems computationally feasible, rather than evaluating the quality of solutions. In the 1950s and 1960s, the literature focused more on formulating tractable optimisation problems (Markowitz, 1971; Sharpe, 1963) than on the later estimation-error concerns that became an important stream in subsequent research.

Advances in data processing and solvers have improved the empirical stability of MV through sophisticated modelling and robust optimisation (Palomar, 2025), however, these improvements remain clustered around stability and risk performance, not returns. EW’s persistent out-of-sample strength continues to question the realised value of added modelling complexity (DeMiguel *et al.*, 2009; Plyakha *et al.*, 2017; Kritzman *et al.*, 2010). At the same time, optimisation is able to add a layer of methodical control over investor constraints like risk appetite (Kritzman *et al.*, 2010). Risk modelling research shows movement toward measures that capture downside losses more directly than variance (Kritzman *et al.*, 2010; Righi and Borenstein, 2018). The consideration of risk is of substantial practical importance since people in general tend to weigh losses more heavily than gains (Kahneman and Tversky, 1979). Thus, the next subchapter takes a deeper look at alternative approaches to risk optimisation in portfolio construction.

1.2. Approaches to Risk Optimisation in Portfolio Construction

Risk measures quantify the overall seriousness of possible losses below a chosen reference level (Rockafellar *et al.*, 2006). Risk optimisation requires complex modelling because small changes in optimiser inputs can lead to very large changes in portfolio asset weights.

Righi and Borenstein (2018) distinguish monetary loss measures (e.g. maximum loss) from deviation measures (e.g. standard deviation) and combine the two approaches into double-barrel loss-deviation measures. This creates a baseline loss measure plus a downside dispersion penalty that captures how spread-out outcomes are when they are worse than the chosen baseline. In their large simulation study, portfolios that were optimised using loss-deviation measures tended to perform better in risk terms than those using monetary loss alone, but no single dominating risk measure was identified. Righi and Borenstein (2018) conclude that no single risk measure dominates in all situations and therefore risk measure selection should be situational. This takeaway is important because it allows one to conclude that a risk measure should be evaluated and chosen based on the specific job that it is intended to do, rather than arbitrary criteria like how commonly-used it is. The next step is therefore to ask which measure is better suited to the empirical properties of asset returns.

Asset returns do not typically follow a normal distribution and they have a tendency to exhibit fat tails, which is a concern when an optimiser relies on variance and standard deviations (Palomar, 2025). Applications of mean absolute deviation (MAD) do not require the assumption of a Gaussian distribution and MAD remains well-defined under any distribution with a finite first moment (the mean). MAD is the average distance between each return observation x_i and the mean of all observations, $\bar{x}$ (Sauga, 2017), and as such can be computed for almost every empirical distribution:

$$MAD = \frac{\sum_{i=1}^n |x_i - \bar{x}|}{n} \quad (1)$$

In one of the earliest alternatives to MV's variance-based formulation, Konno and Yamazaki (1991) developed a MAD-based model, where a linear programming formulation replaces quadratic variance-based optimisation. Using three 60-month datasets of 224 NIKKEI 225 shares, they found that MAD generated portfolios and performance

broadly similar to MV, but with a much lower computational burden. They therefore present MAD as a practical alternative to variance. Moon and Yao (2011) developed a robust MAD portfolio optimisation model (RMAD). Backtesting on 100 U.S. shares selected randomly from NYSE, NASDAQ and AMEX exchanges, with datasets from 1996-2003, showed that RMAD generally outperformed MAD in out-of-sample realised cumulative returns, but underperformed in declining markets with long estimation horizons, in low-volatility portfolios, and in low beta portfolios (Moon and Yao, 2011). Explicit risk metrics were not reported, therefore loss-related evidence is largely indirect insofar as gains can be viewed as negative losses. Almeida *et al.* (2023) studied a returns-maximisation model using MAD thresholds to detect meaningful directional price changes. Their results were overall comparable with the performance of S&P 500, EW, and minimum-variance benchmarks, but depended on how model parameters were tuned. The findings of these three studies are consistent with Righi and Borenstein’s (2018) conclusions about risk optimisation being situational rather than there being one best way to measure and control risk.

Another benefit of MAD appears to be that it has been shown to work with lower-quality price data, since it has lower sensitivity than variance, and also with missing observations (Chen *et al.*, 2023), which is often a concern with empirical data. However, while MAD can capture the entire returns distribution, it does not directly target tail-losses. Conditional value-at-risk (CVaR) has been proposed as an attractive measure for capturing risk, specifically tail-loss risk (Rockafellar and Uryasev, 2000). CVaR is an extension to value-at-risk (VaR) and it quantifies losses beyond a given VaR threshold. At a confidence level α , and with ξ as the loss of a portfolio over a period of time, a portfolio’s CVaR can be defined as follows (Palomar, 2025):

$$CVaR_\alpha = E[\xi \mid \xi \geq VaR_\alpha] \quad (2)$$

Rockafellar and Uryasev (2000) derived a CVaR formulation that works with discrete returns and all returns distributions, in addition to being effectively minimisable in optimisation problems. Rockafellar and Uryasev’s (2006) CVaR is a coherent risk measure in the Artzner *et al.* (1999) sense, meaning that it is translation invariant, subadditive, positively homogeneous, and monotonic (see Appendix A). Coherence is a defensible baseline property because it rules out risk measures that can behave per-

versely under basic portfolio operations like diversification or when scaling positions.

Silva *et al.* (2017) developed several Beta-CVaR optimisation models with MAD and CVaR as risk measures. They backtested on two samples of 2004-2013 Brazilian index constituents; however, without using rolling out-of-sample windows and rebalancing portfolios. Results between their models were heavily dependent on parameter-tuning. The return-maximisation variant consistently achieved the highest returns but also the highest CVaR values. These findings illustrate a trade-off in tail-risk optimisation, where adding a loss-prevention target to optimisation does not automatically reduce risk, and portfolio outcomes depend on how objectives and constraints are set up.

Lorimer *et al.* (2024) compared several non-variance risk measures, including MAD, semi-absolute deviation (SAD), and CVaR, in a model based on a gradient-descent numerical algorithm. They found that MAD, and especially SAD, achieved better outcomes than CVaR, in both annualised returns and Sharpe ratios, although CVaR still outperformed variance. Importantly, the less restrictive CVaR specification (at 90% confidence level) performed better than the more restrictive CVaR specification (at 95% confidence level). However, this ranking should be interpreted cautiously because the results were produced within a specific optimisation architecture that maximises a return-to-risk objective, and in the paper’s formulation, SAD is a monotonic transformation of MAD ($SAD = \sqrt{MAD}$) so part of the reported performance gap may reflect their specific optimisation setup rather than the intrinsic superiority of one risk measure over another.

Empirical evidence in literature is fragmented and direct comparisons between studies are confounded by differences in optimisation methods and objectives, chosen datasets, and evaluation protocols. This context-sensitivity can be argued to be a symptom of underlying fundamental uncertainty, which may make the search for a universal best risk measure a red herring problem. Righi and Borenstein’s (2018) conclusion about risk measures being situational is consistent with this perspective. A more defensible approach would seem to be to evaluate risk modelling within relevant context and look at how a measure behaves in the specific portfolio construction model in the specific asset universe where the model is being used.

Table 2 summarises the main risk measures discussed above. Risk measures are

presented together with their core characteristics and the main empirical findings reported in the reviewed literature.

Table 2

Summary of risk measures discussed in the context of portfolio construction

Risk Measure	Core Characteristics	Main Empirical Findings
Variance / standard deviation. 1952 onwards (Markowitz, 1952).	Symmetric dispersion around the mean. Classic risk measure in MV models, implemented through a returns covariance matrix. Does not exist for all distributions.	Advantages: Widely used in portfolio construction; supported by a large optimisation literature; improved covariance estimators improve out-of-sample performance. (Ledoit and Wolf, 2017). Disadvantages: upside and downside dispersion treated symmetrically; optimisation is highly sensitive to covariance estimation; does not exist for all distributions.
Mean absolute deviation. 1991 onwards (Konno and Yamazaki, 1991). (MAD)	Average absolute deviation from the mean. Computable for all distributions with a finite first moment.	Advantages: Well-defined for distributions with a finite first moment; can replace the variance-based Markowitz formulation with a linear program and produce broadly similar allocations and performance at much lower computational cost (Konno and Yamazaki, 1991); appears less sensitive to lower-quality data (Chen <i>et al.</i> , 2023); robust MAD can outperform nominal MAD, depending on market conditions (Moon and Yao, 2011). Disadvantages: Does not directly target tail-losses.
Conditional value-at-risk. 2000 onwards (Rockafellar and Uryasev, 2000). (CVaR)	Models tail-losses beyond VaR. Explicitly targets the expected severity of losses in the tail.	Advantages: Coherent in the sense of Artzner <i>et al.</i> (1999); targets estimated tail-losses directly; tractably optimisable (Rockafellar and Uryasev, 2000). Disadvantages: Empirical optimisation outcomes depend on the optimisation approach and the choice of α (da Silva <i>et al.</i> , 2017; Lorimer <i>et al.</i> , 2024).
Hybrid monetary loss & deviation loss risk measures	Combinations designed to capture both baseline loss and downside dispersion.	No single dominating risk measure has been identified; risk measures should be assessed in context (Righi and Borenstein, 2018). Portfolio outcomes also depend on the optimisation framework and parameter setup (Righi and Borenstein, 2018; Lorimer <i>et al.</i> , 2024).

Author: compiled by the author.

Overall, the literature does not identify a universally superior risk measure: what is optimal appears to depend on specific circumstances. It has been shown that MAD can result in broadly similar outcomes to variance at a lower computational cost, while CVaR can be used effectively to model potential tail-losses. However, choosing a risk measure to be used does not yet itself resolve the full portfolio construction prob-

lem. Risk can be reduced through broader diversification, as formalised by Markowitz (1959), but in practice, investors are unlikely to hold a large number of individual assets because portfolios with many positions are harder to trade and monitor, not to mention more expensive once transaction costs are considered. The problem therefore involves not only how risk is defined and controlled, but also how many assets are held at any given time, otherwise known as portfolio cardinality. Once a cardinality restriction is imposed, portfolio construction changes from a continuous weighting problem, like MV optimisation, into an asset-selection-and-weighting problem with particular computationally unattractive features. For this reason, the next subchapter considers the challenges with asset allocation under cardinality constraints, and the solutions and compromises that have been found in order to address these issues.

1.3. Asset Allocation Under Portfolio Cardinality Constraints

The role of optimisation in portfolio construction is to decide which assets are selected and how much weight each asset receives. A cardinality constraint limits the number of assets in a portfolio, therefore introducing a trade-off between optimal diversification and optimal practicality. When several optimisation targets and constraints (e.g. cardinality, forbidding short-selling, etc.) are included, closed-form analytical solutions can become difficult or intractable, especially in large asset universes. The optimisation problem is typically solved numerically with optimisation algorithms and solvers, instead of analytically (Kim *et al.*, 2024). The second implementation concern is that unconstrained optimisation can produce very high-cardinality portfolios with numerous non-zero weights, many of which can be very small, as seen with nominal MV models. This requires some form of cardinality management to be practical for investors who do not wish to hold a high number of very small positions (Konno and Yamazaki, 1991).

Jobst *et al.* (2001) study cardinality in the MV framework and show that once an exact asset-count restriction is imposed on the final portfolio, the optimisation problem becomes NP-hard; if fractional trading is disallowed, the problem becomes NP-complete. Jobst *et al.* (2001) test two approximations, an integer restart heuristic and a reoptimisation heuristic, and conclude that these approximations can perform well, although the trade-off is that these methods do not fully reproduce the true-solution discrete MV efficient frontier. Chang *et al.* (2000) reach a similar conclusion after test-

ing different heuristic approximations (genetic algorithm, tabu search, and simulated annealing). These studies show that approximations are a trade-off, but that they can be a sensible practical approach to cardinality constraints.

In practice, optimisation difficulties are often also addressed through staged portfolio construction (Platanakis *et al.*, 2021). Platanakis *et al.* (2021) test four combinations of MV and EW in a two-stage protocol that first allocates weights across ten broad U.S. industry sectors, and then allocates weights to individual shares within those sectors. They find that MV performs better at the sector weighting stage than for individual share allocation, mainly because sectors are composed of many individual assets and this has a smoothing effect in input data. EW in that study performs better than MV at individual share allocation stage, as measured by out-of-sample risk-adjusted returns (CEQ; Sharpe, Dowd, Omega ratios). Liagkouras *et al.* (2022) develop a two-stage design with an initial ESG-based screening procedure followed by MV optimisation over the selected assets under constraints on cardinality floor and ceiling. Their evaluation focuses on ESG-metrics and optimisation quality rather than return and risk measures, but the procedure itself supports separating the asset selection stage by investor-defined criteria from the mathematical weighting optimisation stage.

Overall, the literature does not establish a universally superior allocation optimisation model, but two-stage construction is used in practice and there is support for such construction methods in research as well. Two-stage design supports practical cardinality constraints and can simplify later weights optimisation.

The literature reviewed so far evidences the strong out-of-sample performance of simple unoptimised EW portfolios, but it also illustrates the trade-off between the higher returns and the higher risk of the approach, when compared to models that attempt to optimise risk. Entropy is a mechanism that offers a more structured, or less naïve way to spread portfolio weights. Entropy measures order vs disorder in a system: the less concentrated the portfolio, the higher the entropy. Maximum entropy value is achieved at equal asset weights (Usta and Kantar, 2011). Although there are various formulations (Zhou *et al.*, 2013), the focus here shall be on Shannon entropy, mainly because of the work done by Usta and Kantar (2011) in showing the diversification effectiveness of Shannon entropy in portfolio selection modelling. Usta and Kantar (2011) define Shannon entropy for a portfolio of assets as:

$$H(x) = - \sum_{i=1}^n x_i \ln x_i \quad (3)$$

where x_i is the weight of an asset, $\{x \in \mathbb{R}_+^n : \mathbf{1}^\top x = 1\}$ and $0 \ln(0) = 0$. Assets can only have non-negative weights by default since the natural logarithm of a negative value is undefined (Usta and Kantar, 2011).

Bera and Park (2008) developed a cross-entropy model (CE1), essentially a variant of MV that uses an entropy-based objective to pull portfolio allocation towards an equal weights target. Their backtesting benchmarks included original MV and EW, and they used eight MSCI international equity indices with monthly observations from 1969 to 2005, and rolling estimation windows of 24–120 months. In a 24-month backtesting window, CE1 produced more stable and less concentrated weights than MV while achieving out-of-sample Sharpe ratios and CEQ values minutely above EW and below MV.

Usta and Kantar (2011) included entropy maximisation as one of three objectives in a mean-variance-skewness-entropy (MVSEM) model. They backtested against several benchmark models on various markets using 120 and 240-month rolling windows. The U.S. dataset contained 180 monthly observations for 20 industry portfolios between 1993 and 2007. Focusing on the 120-month rolling window out-of-sample backtest, Sharpe ratio differences between MVSEM and EW were statistically insignificant, but MVSEM had better risk-adjusted returns with a 27% higher Sharpe ratio than MV. MVSEM portfolio turnover was about 70% lower than MV's, which, as Usta and Kantar (2011) note, has practical significance with respect to transaction costs. EW in the study had no turnover by design.

In the reviewed studies, including entropy in MV optimisation helped improve optimiser behaviour, leading to more stable asset weights and reduced turnover. However, the effects on risk-adjusted returns were minimal and EW portfolios remained among the strongest performers. In conclusion, entropy is not directly linked to superior returns but it is supported as an effective diversification measure that retains the attractive weighting and turnover behaviours of equal-weighting while being somewhat more flexible, since it allows for non-equal weight solutions as well.

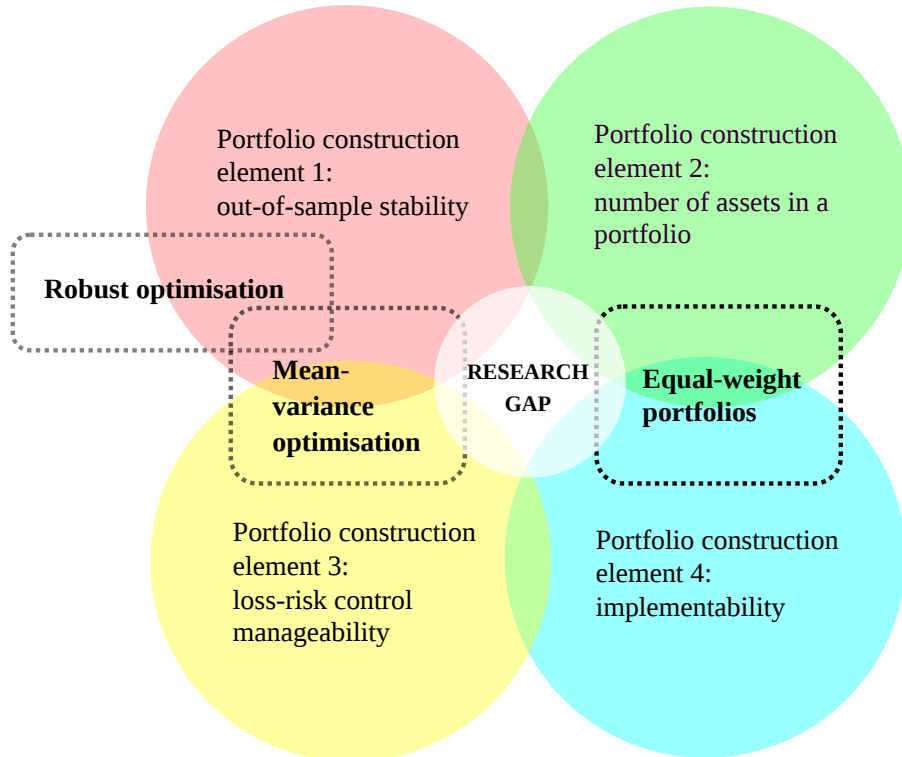
Portfolio design literature clearly evidences estimation-error induced optimisa-

tion limitations of MV with poor out-of-sample performance and the practical complexities of implementing improvements to mitigate this concern. At the same time, even improved MV variants' realistic risk-return outcomes do not appear to meaningfully exceed those of EW portfolios, with EW scoring an additional point for being able to bypass a complex cardinality constraints implementation.

This thesis addresses a middle-ground practitioner-oriented research gap between MV and EW approaches. Figure 1 illustrates four elements of practical portfolio construction: future expectations about portfolio stability; the number of assets that an investor is willing to include; the ability to manage the risk of expected losses, and finally, implementing the model in practice. It also shows where existing academic research overlaps with these construction elements, and where it does not, leaving a research gap that this thesis aims to address.

Figure 1

Portfolio construction research gap addressed in this thesis



Author: compiled by the author.

Note. Mean-variance optimisation here includes methodological advances made to increase MV's out-of-sample stability. The application of robust optimisation in literature and practice is considerably broader (e.g. computer science, engineering, operations) than within the context of MV and investment portfolios in general.

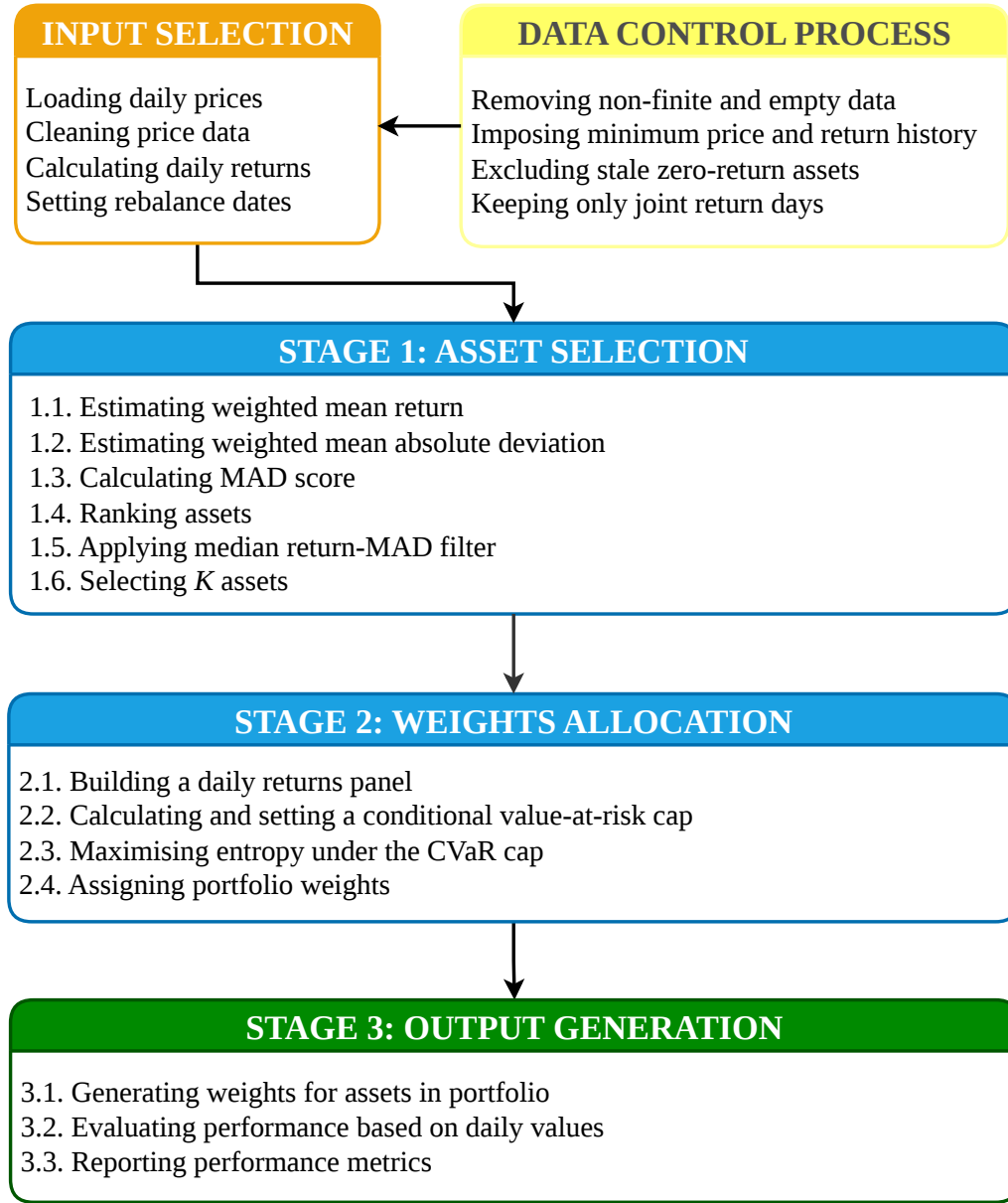
The author believes that implementation simplicity is an important consideration in practice and therefore excludes MV optimisation from the developed model. At the same time, the author also believes that the baseline EW approach is not prudent enough from a risk management perspective. The two-stage construction approach developed in this thesis bypasses the need to solve a cardinality optimisation problem, but it still imposes risk controls. The first stage chooses an exact number of assets with attractive expected return characteristics, using a simple and tractable mean return and MAD-based screening rule. This stage serves both as a returns-seeker and a loss risk-moderator. Then, the optimisation stage assigns weights under a CVaR constraint using entropy, thereby trading off between an equal asset weights optimisation target and tail-loss risk. Based on findings in reviewed literature, this approach is expected to perform better than MV in raw returns terms, but worse than EW and a broad market index from the addition of a tail-loss risk control that is expected to prune some high-yield but volatile assets from constructed portfolios.

2. Methodology and Data

2.1. Methodology

The MAD-CVaR-entropy (MCE) portfolio construction model is implemented in two stages. An overview of the construction process is visualised in Figure 2. The figure shows the steps in processing model inputs along with the application of data controls, and then lists the actions taken through Stage 1 and Stage 2 before outputs are generated.

Figure 2
MCE portfolio construction process



Author: compiled by the author.

As Figure 2 illustrates, during model input selection, infinite values are replaced with missing values; if there are fully empty rows and columns, these are removed, and then daily simple returns are calculated from adjusted close prices only where valid consecutive prices are available. Minimum price and return history requirements are applied so that assets are only considered eligible when enough observations are available in the lookback window. Assets with excessive zero-return days are excluded because repeated unchanged prices may indicate stale or illiquid data. Only joint return days are retained for each selected set of assets, meaning that calculations use days

where valid returns are available for all relevant assets. The purpose is to avoid estimating portfolio risk or return from mismatched asset histories.

The number of assets a hypothetical investor wishes to hold is K . Stage 1 handles the selection of assets for the final portfolio using a rule-based heuristic, outputting a list of K assets for the Stage 2 optimiser. Stage 1 uses daily scenarios computed from historical prices. A scenario is one observed historical trading day in the lookback window, represented by the vector of asset returns on that day. Let $r_t \in \mathbb{R}^n$ denote the vector of asset returns on each day $t = 1, \dots, T$. Daily returns are calculated as simple returns from daily adjusted close prices: $r_t = \frac{P_t - P_{t-1}}{P_{t-1}}$. Exponential time-decay weights $\omega_t \geq 0$ are applied to down-weight older observations, following RiskMetrics™-style (J.P. Morgan and Reuters, 1996) half-life logic to give more recent observations greater weight. Time-weighting in the model is optional, it can be turned off, and default parameter values are based on the author's assumptions. Half-life default value is $h = 6$ months, and the decay factor $\lambda = \frac{\ln 2}{h}$. Let $a_t \in \{0, \dots, T-1\}$ be the scenario age, where $a_t = 0$ is the most recent observation. Scenario weights are then:

$$\omega_t = \frac{e^{-\lambda a_t}}{\sum_{s=1}^T e^{-\lambda a_s}}, \quad \text{where} \quad \sum_{t=1}^T \omega_t = 1. \quad (4)$$

Let j define an eligible asset. Time-weighted mean return over the lookback window for each j is defined as follows:

$$\mu_{w,j} = \sum_{t=1}^T \omega_t r_{t,j}. \quad (5)$$

Time-weighted MAD for each j over the lookback window is found from weighted return and weighted mean return:

$$MAD_{w,j} = \sum_{t=1}^T \omega_t |r_{t,j} - \mu_{w,j}|. \quad (6)$$

Then the mean-MAD-based asset ranking score for each j is calculated from the weighted mean return and weighted MAD as follows:

$$S_{w,j} = \frac{\mu_{w,j}}{MAD_{w,j}}. \quad (7)$$

The score in (7) is used as a practical screening heuristic. It deliberately gives priority

to assets with positive recent weighted mean returns relative to their absolute return dispersion. An asset with a small negative weighted mean return can be ranked lower than an asset with a small positive weighted mean return even if the former has lower MAD. This reflects the model's Stage 1 objective of selecting assets with recent upside characteristics. Risk minimisation is not the only purpose of the screening step.

The rule also means that extreme return observations can affect screening outcomes before Stage 2 is reached. MAD is less sensitive to extreme observations than variance, but it is not robust in the strict statistical sense, especially because the numerator also uses the weighted mean return. Large positive or negative return observations can increase MAD and lower an asset's ranking, or prevent it from passing the median-MAD gate. This is an intentional trade-off in the model. Stage 1 is designed to avoid assets whose recent returns were accompanied by large return swings, even when some of those swings were favourable. A median-return or median-absolute-deviation-based screen could encode a structure more robust to outliers, which could be interpreted as reflecting more conservative investor preferences but that would constitute an entire alternative screening protocol and therefore a different model. In MCE, assets are ranked by the $S_{w,j}$ score in descending order.

Let $m_{MAD} = \text{median}_j(MAD_{w,j})$ be the cross-sectional median of weighted MAD values across all eligible assets in the current lookback window. This threshold is calculated separately at each rebalance date and the model assumes that this ranking reflects desirable asset return characteristics that may be likely to hold in the next period. Stage 1 scans the ranked list and selects the first K assets that satisfy the $MAD_{w,j} \leq m_{MAD}$ returns volatility gate: assets that satisfy the inequality are in the lower-MAD half of the eligible asset universe. If fewer than K assets satisfy the condition, the remaining slots are filled by continuing down the ranked list until exactly K assets are selected. This means that some slots may get filled with assets where the $MAD_{w,j} \leq m_{MAD}$ condition is not satisfied. High-MAD assets can enter the final selected set only if they still satisfy this gate, or if the fallback rule is needed to maintain the fixed portfolio cardinality of K assets. The cardinality-related constraint is accepted as a practical trade-off.

Stage 2 determines weights for the selected K assets. Shannon entropy (3) is maximised subject to a time-weighted CVaR constraint on historical portfolio losses.

The formulation does not allow for negative positions, therefore excluding short selling and only allowing long positions by default. In (3), each term $(-x_i \ln x_i)$ is the contribution of each one asset to total portfolio entropy. Asset weights are treated analogously to probabilities because they are non-negative and sum to one. Total portfolio entropy is the sum of all individual entropy terms, and the terms can be interpreted as each asset i 's contribution to total portfolio diversification.

The CVaR cap is based on a Rockafellar and Uryasev (2000) formulation. Let $x \in \mathbb{R}^K$ be the weight vector over the selected K assets with $x \geq 0$ and $\mathbf{1}^\top x = 1$. x_i is therefore the portfolio weight of asset i . Separately, ω_t is the time-decay weight assigned to historical return scenario t , as defined in (4). Let scenario loss be defined as $L_t(x) = -r_t^\top x$. For confidence level $\alpha \in (0, 1)$, with $\eta \in \mathbb{R}$ being an auxiliary threshold variable corresponding to the portfolio's VaR level, and $z \in \mathbb{R}_+^T$, where each z_t represents the excess of scenario loss over η in period t :

$$CVaR_\alpha(x) = \min_{\eta \in \mathbb{R}, z_t \in \mathbb{R}_+^T} \left\{ \eta + \frac{1}{1 - \alpha} \sum_{t=1}^T \omega_t z_t \right\} \quad (8)$$

subject to $z_t \geq L_t(x) - \eta$, $t = 1, \dots, T$. CVaR is imposed via capped $CVaR_\alpha(x) \leq \bar{C}$, with $\bar{C}$ anchored to the computed CVaR of an equal-weight portfolio composed of all eligible assets in the lookback window. Essentially, the constructed portfolio must not be riskier than an EW portfolio made from all assets in the eligible universe in the same period. The anchoring choice is based on the author's subjective judgement of a reasonable target. A different basis for the CVaR cap target would constitute a different model. Shannon entropy achieves a maximum value under equal weights so if the CVaR cap is high enough not to bind, for example in a relatively calm lookback window, the resulting portfolio's assets are weighted at $1/K$. This is intentional and exploits the empirical high performance of EW portfolios seen in literature.

Turnover rate is calculated for each portfolio. Turnover considers position values at rebalance and shows the fraction of value traded in order to move from previous weights to new target weights:

$$turnover_t = \frac{1}{2} \sum_i \left| w_{i,t}^{\text{target}} - w_{i,t}^{\text{pre-rebalance}} \right| \quad (9)$$

Three comparison portfolios are constructed on the same rebalance schedule.

MV-K uses the same K asset names selected in Stage 1, and MV-All uses the full asset universe present in the lookback window. Portfolio weights for the MV variants are obtained from a long-only mean-variance optimisation using Ledoit-Wolf covariance shrinkage (Ledoit and Wolf, 2004a) as available in the Python package Scikit-learn. Optimisation is implemented as a target-return sweep of variance-minimising portfolios, and the selected solution is the one with the highest return-to-volatility ratio (measured by variance). If the tangency search is degenerate or infeasible (e.g. in a falling market with mostly negative returns in the lookback window), a minimum-variance fallback portfolio is constructed instead as a practical trade-off. The third comparison portfolio is EW over K assets. This is not the same subset of assets that are used in the MCE portfolio. EW uses a very simple ranking approach on purpose, based on the author’s aim to keep the EW construction model as simple as possible. Ranking is necessary because the asset universe is larger than K . At each rebalance date, the eligible asset universe is ranked by arithmetic mean return over the lookback window. From that ranked list, the top number of K assets are selected, and each receives a weight of $1/K$. Neither the MCE portfolio nor the alternative portfolios consider transaction costs. This narrows the conclusions that can be drawn, but avoids making additional assumptions and adding additional parameters that can confound results.

2.2. Data

The MCE, EW, MV-K, MV-All models were backtested on two U.S. datasets: 72 exchange-traded funds, and 371 individual technology-sector shares. Names and tickers of assets are listed in Appendix E. S&P 500 (SPY) was selected as a broad market-proxy benchmark.

Table 3

Datasets used in portfolio construction model backtesting

Dataset and source	N	Data coverage	Label
Broad-coverage U.S. exchange traded funds Source: author-chosen	72	03/01/2005– 02/01/2026	ETF universe
Individual U.S. large-cap technology shares Source: Stock Analysis website	371	02/01/2014– 02/01/2026	Tech universe

Author: compiled by the author.

ETF universe assets were selected manually by the author. The aim was to achieve broad cross-asset exposure, reduce direct overlap with SPY constituents, and

avoid an excessive concentration of short-duration or otherwise cash-like defensive funds. The technology universe was constructed using the Stock Analysis website by selecting all U.S.-domiciled companies in the technology sector with market capitalisation above USD 300 million.

Fixed-cardinality portfolios were constructed with $K = 20$ assets each as a pragmatic but subjective design choice based on an assumption about how many assets the average investor may consider holding in practice. MCE K -sensitivity analysis can be found in Appendix D. MV models were allowed to use the entire asset universe for portfolio construction.

Table 4

Backtested portfolio construction models, abbreviations used, and number of assets in portfolios

Model	Abbreviation	K
MAD-CVaR-entropy	MCE	= 20
Equal-weights	EW	= 20
Mean-variance full universe	MV-All	≤ eligible universe
Mean-variance K-assets	MV-K	≤ 20
Benchmark: S&P 500	SPY	not reported

Author: compiled by the author.

Periodic portfolio rebalancing was set to take place semi-annually instead of a monthly cycle more commonly seen in literature. The main considerations were practical transaction costs, which although not accounted for, would be considerably higher with monthly rebalancing than semi-annual, and the assumption of longer periods allowing for smoother data, albeit possibly missing out on momentum signals. Daily adjusted close prices were downloaded from Yahoo Finance using the Python package `yfinance`, from 03 January 2005 to 02 January 2026 for the ETF universe, and 02 January 2014 to 02 January 2026 for the technology universe. SPY was explicitly excluded from inclusion into any constructed portfolio. Risk-free rate used in performance metrics was the 3-month U.S. Treasury constant-maturity yield (DGS3MO) across the evaluation period. Final data used in backtesting simulations was downloaded on 19 March 2026. Survivorship bias in the downloaded data is very likely because the data only contains tickers and prices that exist in the Yahoo Finance database at the time the API connection for a download is made. There is no information available on which tickers may have been removed from the database in the past. If the delisted companies had poorer

returns than remaining listed companies, then the results may be biased upward. There is unfortunately no way of confirming or refuting this assumption using the data available from the given source.

Data quality controls were applied downstream:

- After loading the price panel, infinite values were replaced with missing values, and rows or columns that were entirely empty were removed before calculating daily simple returns.
- Asset eligibility was enforced within each lookback window by requiring a minimum number of observed prices and returns. Since this study used 12-month lookback windows, the requirement was at least 240 price observations and 239 return observations, since returns are calculated from changes in price.
- Under the minimum observations rule, the first MV-All estimation window in the technology universe included only 38% of the total 371 assets, although this share increased in later periods. This indicated limited usable history in early estimation windows and supported the decision not to look back further than 2014 in order to preserve a relatively large eligible asset count across the simulation.
- For the EW and MV-All portfolios, assets whose maximum number of exact zero-return days across rebalancing windows exceeded 10 days were excluded in Stage 1 in order to reduce the influence of stale or illiquid assets on Stage 2 outcomes.
- Optimisation scripts for the fixed-asset portfolios, MCE and MV-K, used strict intersection-of-days hygiene, meaning that only days with valid returns for all selected assets were retained. Rebalancing dates with insufficient joint coverage were skipped or, in the case of MCE, treated as a hard failure.
- Evaluation imposed strict consistency checks across files and dates: rebalance dates had to match the return index, selected tickers had to exist in the return panel, and portfolio return series were marked missing whenever any constituent asset return was missing on a valuation day.

Together, these controls limit the influence of missing, stale, or inconsistent data in portfolio construction and evaluation. Controls also improve the interpretability and comparability of results across the tested construction models. As a result, observed

differences in model performance are more likely to reflect modelling choices rather than distortion introduced by data-quality problems.

3. Results and Discussion

This chapter reports and discusses the empirical out-of-sample backtesting of the MCE portfolio construction model relative to MV-K, MV-All, EW, and the SPY benchmark. Backtesting was done on two materially different asset universes and therefore results are first presented separately and then brought together in a final discussion. The following metrics were used in the evaluation: (i) final cumulative return, (ii) mean semi-annual holding-period return, (iii) annualised Sharpe ratio, (iv) annualised Sortino ratio, (v) annualised information ratio (IR) using SPY as the benchmark, (vi) mean $CVaR_{0.95}$ across all holding periods, and (vii) mean turnover rate across all holding periods. See Appendix B for ratio definitions. Paired t -test and paired Wilcoxon signed-rank test results comparing MCE's semi-annual returns to other models are in Appendix C. The sensitivity of MCE results to cardinality choices is tested in Appendix D. All scripts, data files and audit logs are available on request from the author's GitHub repository (see Appendix G).

3.1. ETF Universe Simulation Results

The ETF universe backtesting simulation was executed over a total of 39 consecutive semi-annual holding periods using daily price series data from 2005 to 2026. At each January and July rebalance date, daily historical data was evaluated in a 12-month lookback window using 6-month half-life for time-weighting observations.

Table 5
ETF universe empirical portfolio performance metrics

Model	Periods	Cumulative return (%)	Mean semiannual return (%)	Annualised Sharpe ratio	Annualised Sortino ratio	Annualised information ratio	CVaR _{0.95} (%)	Mean turnover (%)
EW	39	337.358	4.519	0.407	0.381	-0.361	2.417	35.427
SPY	39	667.909	5.955	0.551	0.521	—	—	—
MCE	39	384.470	4.564	0.493	0.459	-0.431	1.919	22.146
MV-K	39	265.972	3.585	0.513	0.486	-0.358	1.403	62.862
MV-All	39	152.708	2.604	0.347	0.328	-0.487	1.346	55.152

Author: compiled by the author.

Note. Cumulative return is the final full-sample (backtest scenario) compounded return.

Mean semiannual return is calculated by compounding daily returns within each holding period and then averaging across these periods.

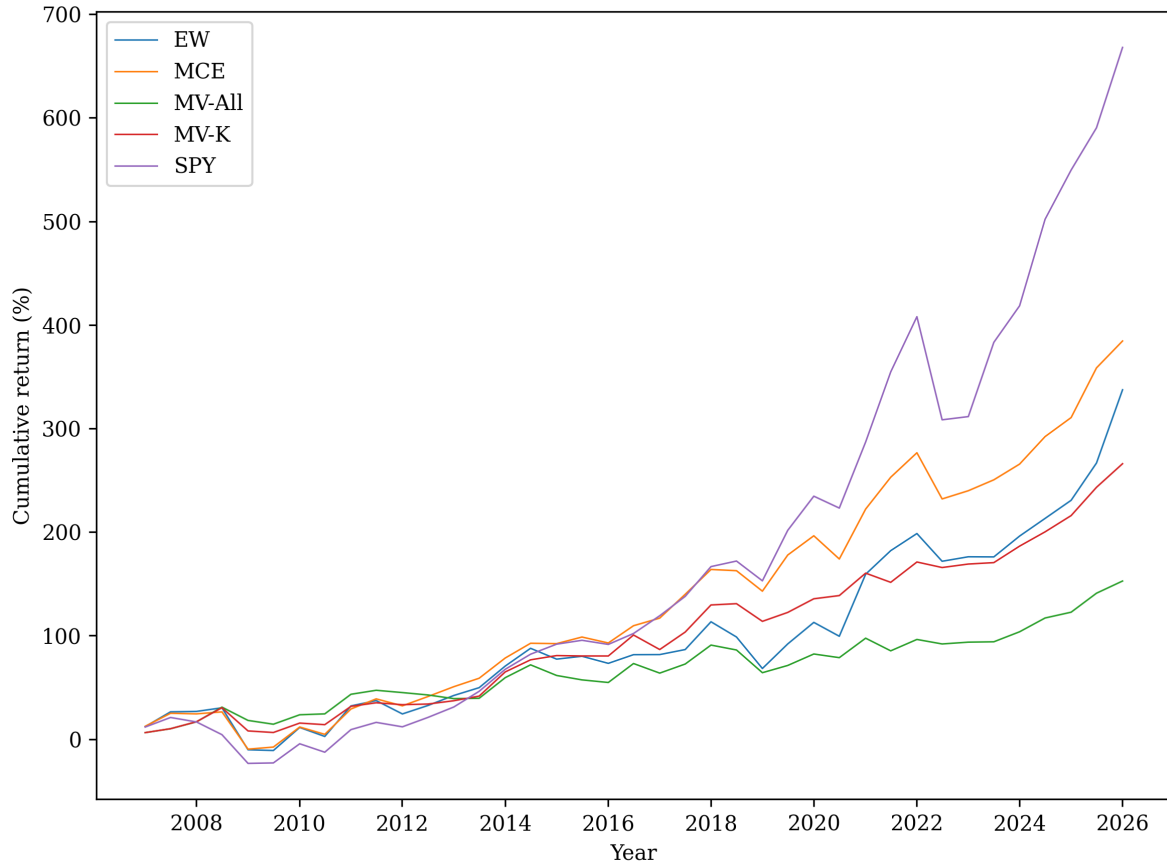
Annualised ratios are calculated from daily returns across the full sample.

Information ratio is calculated with S&P 500 (SPY) as a benchmark.

Mean CVaR is calculated from daily losses within each holding period and then averaged across all holding periods.

Mean turnover is averaged across rebalance dates. SPY is a single-asset benchmark; no CVaR or turnover is reported.

Out of all models, MCE achieved the highest cumulative return by the end of the simulation (384.47%), which is 14% higher than EW's outcome (337.36%) in relative terms. Both MV variants ended with lower cumulative returns than MCE and EW, but MV-K performed relatively better than MV-All by 74.09%, which is a significant performance gap between the models in this scenario. MV-All was the worst performer by far with a 152.71% cumulative return. However, no constructed portfolio came close to the 667.91% cumulative return achieved by SPY.

Figure 3
ETF universe cumulative returns of tested portfolios and SPY benchmark


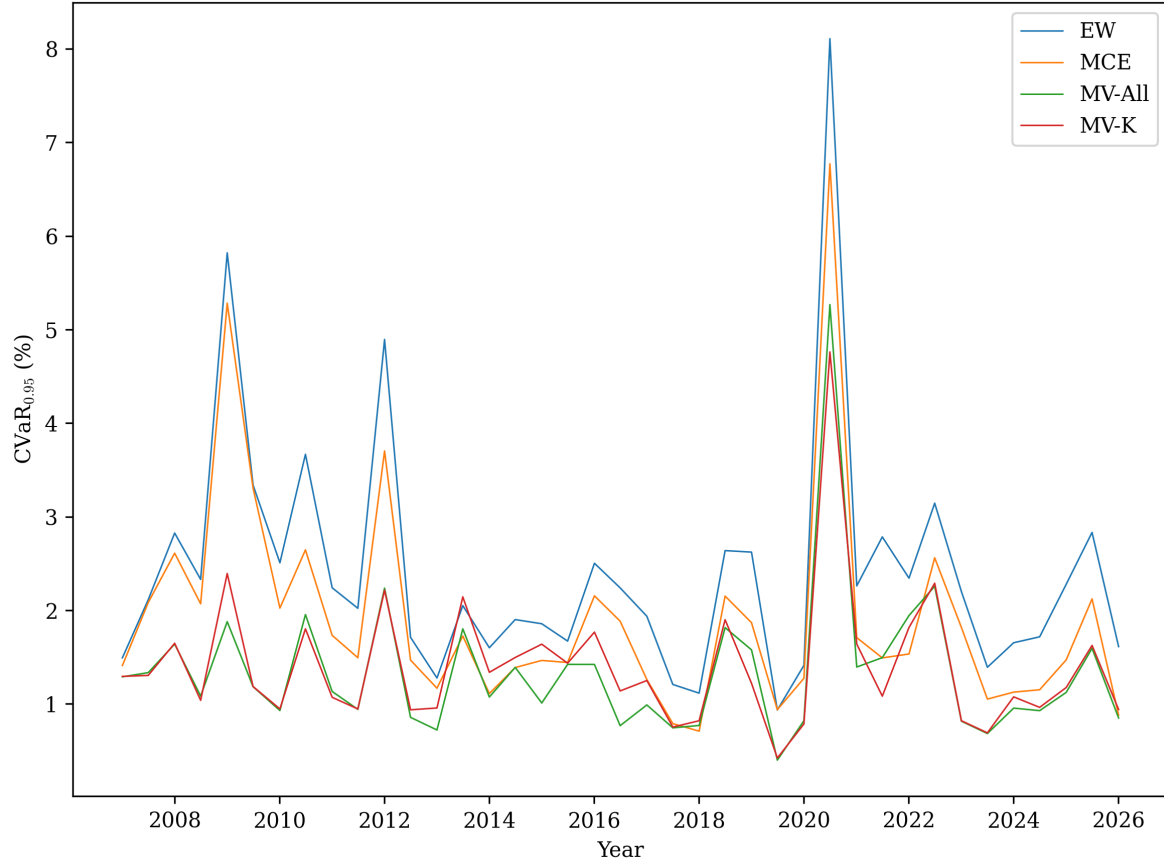
Author: compiled by the author.

MCE's outperformance of EW was unexpected; however, this result should be interpreted within the specific ETF universe and backtesting protocol. Although both portfolios contain a K number of assets, the two models' asset selection protocol is substantially different, as covered in the Methodology subchapter. MCE and EW were essentially tied on mean semi-annual returns across the total simulation timeline, respectively at 4.56% and 4.52%; therefore, mean holding-period returns alone do not explain MCE's final higher cumulative return. Paired t -tests and Wilcoxon signed-rank tests on model returns support a cautious interpretation (see Appendix C). MCE's semi-annual return differences relative to EW and MV-K are statistically insignificant at the 5% significance level. Evidence relative to MV-All is mixed. The Wilcoxon signed-rank test is statistically significant at the 5% significance level but the paired t -test is not. The comparison with SPY shows the same mixed pattern in test results, but in the opposite direction.

Risk-adjusted return metrics show that out of all portfolios, MV-K achieved the highest Sharpe and Sortino ratios overall (0.513 and 0.486 respectively), followed closely by MCE (0.493 and 0.459), while MV-All was the weakest performer (0.347 and 0.328). No portfolio performed higher than SPY in risk-adjusted returns. MCE's annualised Sharpe and Sortino ratios were both about 21% higher than EW's, which is in line with MCE's higher final cumulative return even if the mean semi-annual return was practically identical for the portfolios. Annualised information ratios vs SPY show that all portfolios underperformed the benchmark index on a risk-adjusted basis, but substantive conclusions cannot be drawn between the portfolios themselves based on IR.

The tail-loss metric, mean $CVaR_{0.95}$, is the average severity of the worst 5% of daily losses within each holding period, which is then averaged across all 39 holding periods. MV-All had the lowest average daily tail-losses at mean $CVaR_{0.95}$ at 1.35%, followed closely by MV-K at 1.40%. MCE and EW achieved 1.92% and 2.42% respectively, but the 0.5% gain in the raw magnitude of daily tail-losses can be considered minor.

Figure 4 illustrates the differences in risk-handling between the models in this simulation. MV variants exhibit clearly lower average tail losses and track each other closely through rebalancing points. MCE and EW follow similar paths to each other, but also show that MCE improves on EW, which has no risk controls.

Figure 4
ETF universe portfolio $CVaR_{0.95}$ across semi-annual rebalancing windows


Author: compiled by the author.

The $CVaR_{0.95}$ metric spikes show an indication of market-stress clustering in the portfolios. The largest spikes after the 2008 credit crunch, the 2010 euro-area sovereign debt crisis, the 2011 U.S. sovereign debt downgrade, the 2020 start of the COVID-19 crisis, the 2022 start of the Russia-Ukraine war, and the 2025 U.S. tariff shock provide evidence of a lagged pattern. The more defensive models, MV-K and MV-All, were more protective in crisis periods, but MCE's CVaR control was able to temper the effects of realised tail-losses over EW by a relative 20% on average across the holding periods. Further differences in the outcomes between MV variants and MCE can be explained by portfolio allocations, since both used fundamentally different weighting protocols. EW allocation was of course a flat $1/K$ in all periods.

Both MCE and MV-K portfolios use the same asset universe input for weighting, but while the MCE optimisation target pushes allocations towards equal weights through entropy maximisation under a CVaR cap, MV-K uses a standard covariance-estimation process, which helps explain the lower number of practical active positions.

Table 6*ETF universe MV-K holdings count by rebalance date*

Rebalance date	Weights > 0	Weights above threshold	Rebalance date	Weights > 0	Weights above threshold
2006-07-03	17	7	2016-07-01	9	7
2007-01-03	20	6	2017-01-03	16	4
2007-07-02	18	6	2017-07-03	19	6
2008-01-02	16	7	2018-01-02	8	7
2008-07-01	15	15	2018-07-02	16	5
2009-01-02	14	10	2019-01-02	14	10
2009-07-01	16	9	2019-07-01	10	6
2010-01-04	18	6	2020-01-02	20	7
2010-07-01	15	6	2020-07-01	20	6
2011-01-03	17	6	2021-01-04	17	6
2011-07-01	10	7	2021-07-01	13	13
2012-01-03	19	19	2022-01-03	17	6
2012-07-02	4	4	2022-07-01	11	8
2013-01-02	15	8	2023-01-03	16	8
2013-07-01	19	5	2023-07-03	4	3
2014-01-02	16	5	2024-01-02	18	5
2014-07-01	14	6	2024-07-01	7	6
2015-01-02	20	6	2025-01-02	8	7
2015-07-01	6	6	2025-07-01	20	8
2016-01-04	4	4	2026-01-02	19	7

Author: compiled by the author.*Note.* Portfolio holdings counts are reported using two thresholds: weights greater than 0, and weights greater than 10^{-12} .

The eligible set of 20 assets for MV-K is produced by Stage 1. While most periods included over half of the eligible set into a portfolio, a relatively low number of assets were assigned non-tiny weights (above 10^{-12}). Table 6 shows that the July 2012, January 2016, and January 2017 portfolios contained only 4 non-tiny positions, and July 2023 had 3. Data shows that allocations for these periods were defensive, concentrating on defensive sectors, currencies, and gold.

Table 7*ETF universe MV-All holdings count by rebalance date*

Rebalance date	Weights > 0	Weights above threshold	Rebalance date	Weights > 0	Weights above threshold
2006-07-03	18	7	2016-07-01	43	43
2007-01-03	23	6	2017-01-03	28	8
2007-07-02	21	6	2017-07-03	32	14
2008-01-02	19	6	2018-01-02	39	39
2008-07-01	32	10	2018-07-02	44	10
2009-01-02	26	10	2019-01-02	36	12
2009-07-01	34	11	2019-07-01	22	8
2010-01-04	33	7	2020-01-02	47	47
2010-07-01	36	7	2020-07-01	48	8
2011-01-03	29	6	2021-01-04	48	6
2011-07-01	20	8	2021-07-01	37	37
2012-01-03	40	7	2022-01-03	31	11
2012-07-02	6	4	2022-07-01	43	10
2013-01-02	27	10	2023-01-03	35	10
2013-07-01	41	6	2023-07-03	33	33
2014-01-02	38	8	2024-01-02	49	10
2014-07-01	40	13	2024-07-01	46	46
2015-01-02	45	45	2025-01-02	47	47
2015-07-01	45	9	2025-07-01	32	32
2016-01-04	46	8	2026-01-02	44	14

Author: compiled by the author.*Note.* Portfolio holdings counts are reported using two thresholds: weights greater than 0, and weights greater than 10^{-12} .

MV-All showed similar behaviour to MV-K, although on a different scale. Audit logs show that the minimum-required observations data quality rule excluded some tickers from the 72-asset universe that had insufficient observations for the estimation window. 27 of the 72 ETFs were excluded at least once but by the last period, only 1 ETF remained ineligible. MV-All showed even stronger relative concentration than MV-K in the same periods: July 2012, January 2016, and January 2017. However, the July 2023 MV-All portfolio contained 47 assets vs MV-K's 3. Most of the weight in that portfolio was allocated to 10 names so it was still highly concentrated in practice despite having a higher raw asset count.

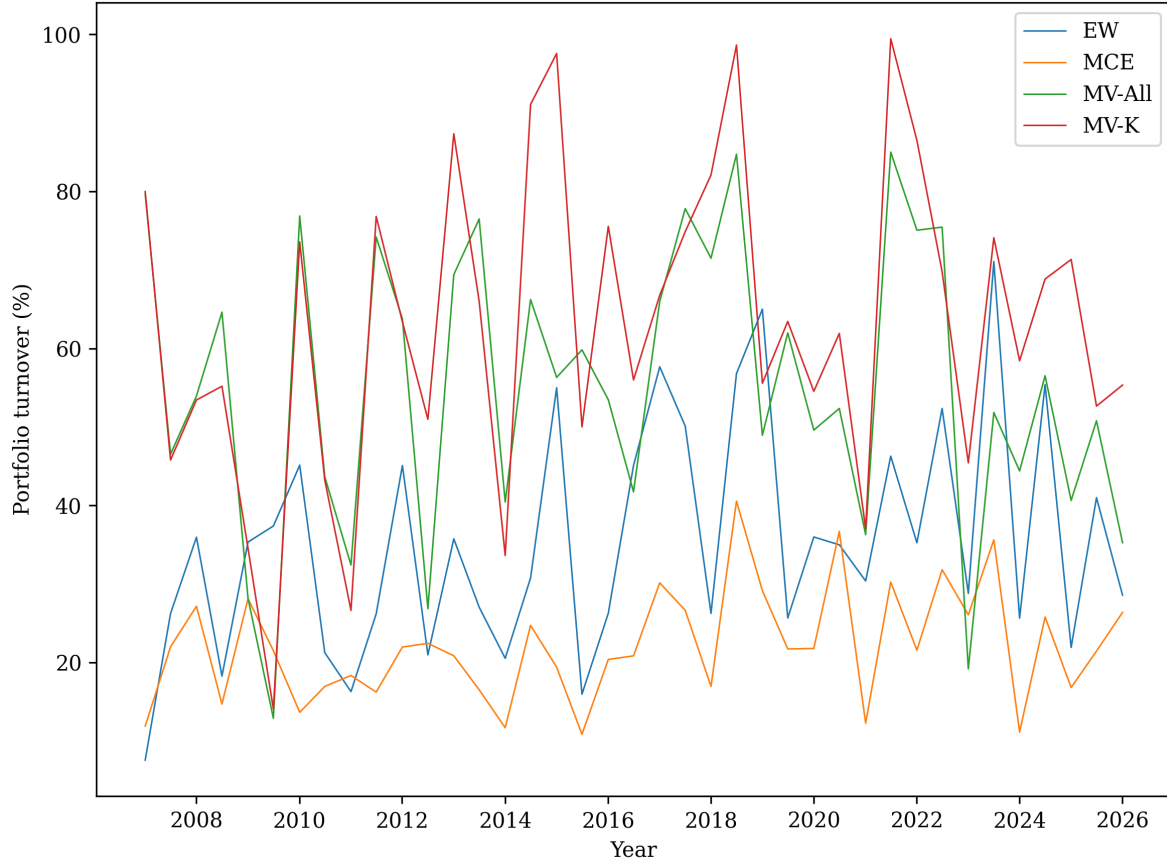
In the same observed periods, MCE portfolios were similarly defensive in asset composition, but contained all $K = 20$ assets due to strict hard cardinality enforced in the construction procedure. July 2012 and January 2016 portfolios were equally weighted, but the CVaR cap constraint fulfilled its intended purpose, binding in January 2017, so portfolio weights in that period were not uniform. The average overlap across all holding periods between assets in MCE and EW portfolios was 11.28 out of

20, or in percentage terms, 56.38%.

Turnover also showed significant differences between the portfolios. MCE achieved the lowest turnover rate out of all models. On average, per rebalance, MCE traded 22.15% of portfolio value and EW traded 35.43%. MV variants' turnover was notably higher with MV-K trading 62.86% and MV-All 55.15% of portfolio value on average. The turnover rate fluctuated for all portfolios, but Figure 5 illustrates that MCE was the most stable out of the models. MV-K and MV-All turnover behaviour is explained by how MV models estimate asset weights, which is further shown in position concentrations seen in Tables 6 and 7. Between the two portfolios with highest cumulative final returns, MCE's average turnover was in relative terms 37.49% lower than EW's. Transaction costs are not included in any of the tested models, and turnover is measured from changes in relative portfolio weights rather than the number of units traded, therefore the effects on realistic transaction costs cannot be directly inferred from the available data. However, lower turnover hypothetically implies lower trading-related expenses.

MCE performance sensitivity to different cardinality choices was tested and results are reported in Appendix D. Sensitivity testing was not performed for the other models. Lower-cardinality MCE portfolios produced higher cumulative and risk-adjusted returns in the ETF universe, but they also had higher turnover. The testing range was capped at $K = 22$ because only 22 assets passed the maximum zero-return days data quality control filter for the very first rebalance date. For this particular dataset, the baseline $K = 20$ specification can be argued to be a balanced trade-off between turnover and return.

Sensitivity tests further show that non-equal-weight allocations occurred more frequently at lower cardinalities and became rare as K approached 22. This appears to suggest that the CVaR constraint is more likely to influence asset weights when the portfolio is more concentrated. However, this result should be interpreted within the given asset universe and parameter values.

Figure 5*ETF universe portfolio turnover across semi-annual rebalancing windows**Author:* compiled by the author.

The comparison between MCE and MV-K is especially informative because both portfolios are evaluated on the exact same Stage 1-selected asset set, so the difference between them is distinctly attributable to weighting, not selection. MCE achieved a materially higher cumulative return than MV-K (384.47% versus 265.97%) and did so with much lower turnover (22.15% versus 62.86%). MV-K retained an advantage in Sharpe and Sortino ratios, and mean $CVaR_{0.95}$. This suggests that, on the same 20 assets, MV-K constructed a more defensive but notably more concentrated portfolio, whereas MCE extracted more terminal wealth from the selected basket and did so with a substantially lower rebalancing burden.

The ETF-universe results show that no active portfolio dominated on all metrics. Among the models, MCE emerged as the most balanced: it achieved the highest final cumulative return of the four, exceeding EW despite nearly identical mean semi-annual return, and did so with lower mean $CVaR$ and substantially lower turnover. MV-K performed best on Sharpe and Sortino ratios, and mean $CVaR$ among the port-

folios, but this was achieved with much higher turnover and, at several rebalance dates, with very sparse practical allocations. MV-All appeared overly defensive, producing the lowest cumulative return while not offering enough additional risk reduction to offset that weakness. Overall, in the ETF universe, MCE provided arguably the strongest practical compromise. The next subchapter examines how this pattern changed in the relatively much more narrow single-sector technology-share universe.

3.2. Technology Universe Simulation Results

The technology universe backtesting simulation was executed over a total of 21 consecutive semi-annual holding periods using daily price series data from 2014 to 2026. At each January and July rebalance date, daily historical data was evaluated in a 12-month lookback window using 6-month half-life for time-weighting observations.

Table 8

Tech universe empirical portfolio performance metrics

Model	Periods	Cumulative return (%)	Mean semiannual return (%)	Annualised Sharpe ratio	Annualised Sortino ratio	Annualised information ratio	Mean CVaR _{0.95} (%)	Mean turnover (%)
EW	21	697.135	12.720	0.665	0.641	0.464	4.639	64.686
SPY	21	292.808	7.176	0.696	0.656	—	—	—
MCE	21	315.496	7.747	0.611	0.572	0.140	3.001	67.275
MV-K	21	153.041	5.327	0.408	0.383	-0.235	3.027	82.792
MV-All	21	317.192	7.707	0.666	0.635	0.079	2.595	61.627

Author: compiled by the author.

Note. Cumulative return is the final full-sample (backtest scenario) compounded return.

Mean semiannual return is calculated by compounding daily returns within each holding period and then averaging across these periods.

Annualised ratios are calculated from daily returns across the full sample.

Information ratio is calculated with S&P 500 (SPY) as a benchmark.

Mean CVaR is calculated from daily losses within each holding period and then averaged across all holding periods.

Mean turnover is averaged across rebalance dates. SPY is a single-asset benchmark; no CVaR or turnover is reported.

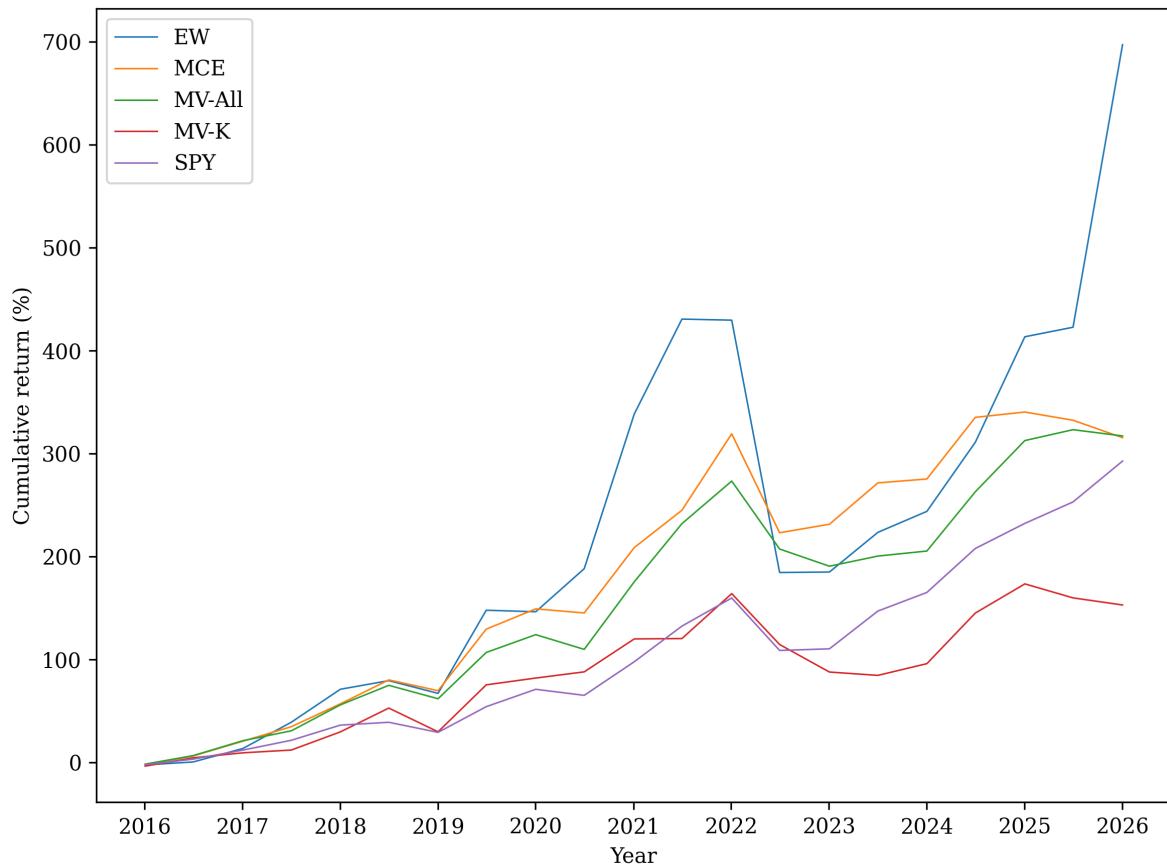
EW achieved the highest cumulative return (697.13%) out of all models: 119.78% more in relative terms than MV-All, which was the second highest earner with 317.19% cumulative return. MCE's cumulative return was 315.50%, or 0.05% lower than MV-All's in relative terms. These two models achieved almost identical terminal wealth despite using fundamentally different strategies. MV-K showed the worst cumulative return performance (153.04%) and SPY was second worst (292.80%). This is an interesting result for SPY compared to EW specifically. Figure 6 shows clearly that EW produced very high returns in the last holding period.

Although sample mean semi-annual returns differ across models, paired t -tests and Wilcoxon signed-rank tests (see Appendix C) do not show statistically significant paired return differences between MCE and any comparison portfolio at the 5% significance level. The results should be interpreted as descriptive and criterion-dependent, especially given the small sample of 21 holding periods.

Cardinality-sensitivity results are also descriptive rather than inferential (see Appendix D). Changing K changed which assets Stage 1 selected, but it did not change how Stage 2 behaved. The CVaR cap did not bind in any tested period, so the entropy objective always reached its maximum at equal weights. Therefore, performance differences between cardinality choices mainly reflect differences in the selected assets.

Figure 6

Tech universe cumulative returns of tested portfolios and SPY benchmark



Author: compiled by the author.

January 2026 portfolios deserve a deeper look into the data because they illustrate the nature of the U.S. technology sector in context with other periods. The constituent assets and share price changes for EW and MCE are provided in Appendix

F. Five assets in EW more than doubled in price with the top two gaining 462% and 322% respectively. Performance was focused on AI & defence-industry momentum. At the same time, the bottom two shares held by EW both dropped over 50% in price. MCE's biggest gainer appreciated by 74% in the same period, and none of its constituents lost value. Both portfolios were equally weighted in this period. The EW construction protocol picks the top gainers in the estimation window. MCE picks top gainers but penalises large price swings in either direction. SPY had a comparatively low final cumulative return compared to EW, MCE and MV-All, but it significantly outperformed all but EW in the final holding period ending in January 2026. EW gained 52.49% and SPY 11.24%, while MV-All ended with -1.43%, MV-K -2.63%, and MCE with -3.92%.

Similarly to cumulative returns, EW was the highest performer in mean semi-annual returns as well. MCE was second and MV-All third, with nearly identical results of 7.75% and 7.71% respectively. MV-K was the worst performer with 5.32%. EW's statistics were dominated by two very high-growth periods seen in Figure 6. The gains made between January 2020 and January 2021 were wiped out in July 2022, but the steep growth that began in January 2023 continued on until final evaluation in January 2026. EW's mean semi-annual return across the simulation was 12.72%, exceeding the second-best performing model, MCE, by a relative 64.20%.

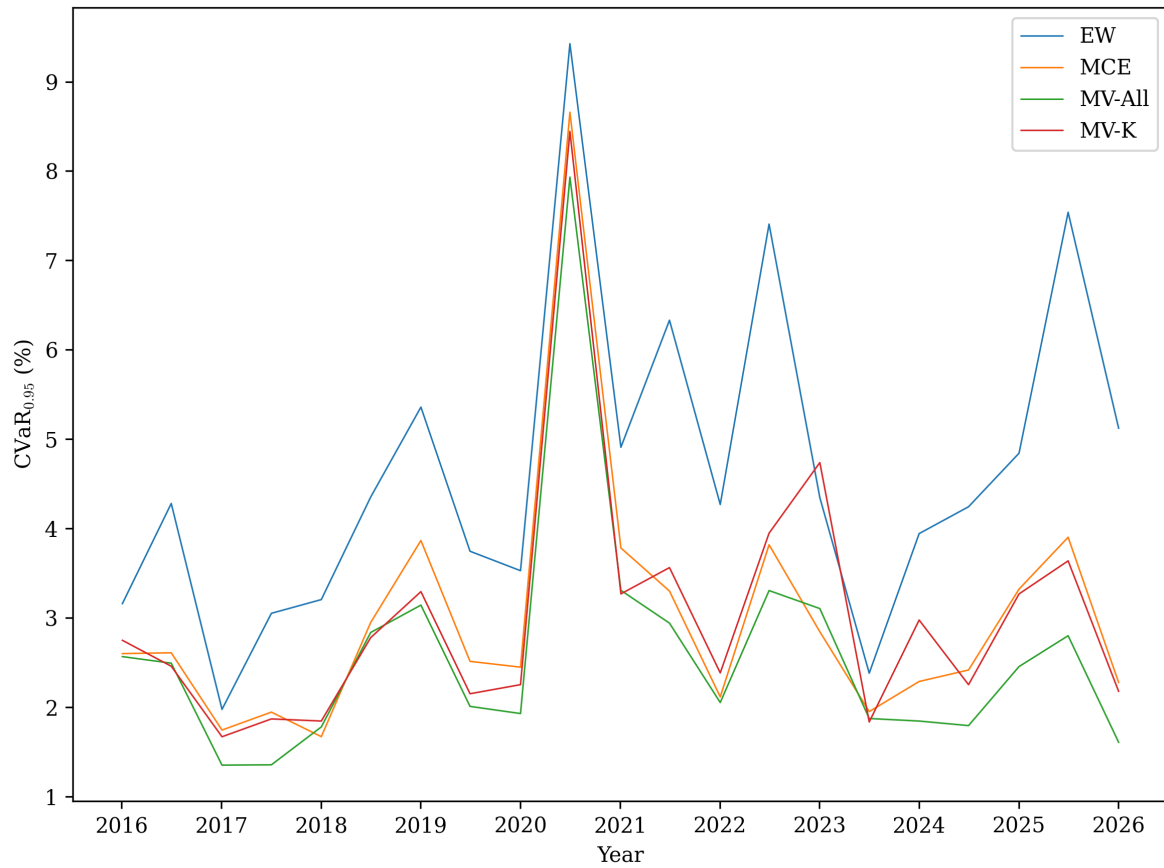
Risk-adjusted return metrics across all 21 holding periods show that the portfolios with the highest overall annualised Sharpe and Sortino ratios were EW (0.665 and 0.641) and MV-All (0.666 and 0.635). MCE ranked slightly lower with 0.611 and 0.572 respectively, and MV-K was the worst performer (0.408 and 0.383). SPY had the best risk-adjusted returns, achieving the highest annualised Sharpe and Sortino ratios (0.696 and 0.656). Across the entire simulation, although EW had the highest raw cumulative return performance, the SPY benchmark had the best risk-adjusted return performance. Annualised information ratios show that EW, MCE, and MV-All portfolios outperformed SPY on a risk-adjusted basis, but MV-K underperformed the benchmark.

The tail-loss metric, mean $CVaR_{0.95}$, provides further context for what's already been shown about EW's behaviour in this simulation. EW's mean $CVaR_{0.95}$ (4.63%) was nearly 54.57% higher than MCE's (3%) and 53.26% higher than MV-K's (3.03%) in relative terms. MV-All had the lowest mean $CVaR_{0.95}$ at 2.59%. EW's results il-

illustrate the strategy's trade-off between risk and return. This is evident in Figure 7, which shows that the three portfolios with explicit risk controls (MCE, MV-K, MV-All) remained relatively close to each other, while EW's mean $CVaR_{0.95}$ was higher in all periods, including in the first rebalance point after the start of the COVID-19 crisis, and also fluctuated comparatively more. Also, while MCE and MV-All use entirely different construction protocols, they achieved similar returns and risk performance. MCE and MV-K are constructed from the same list of assets, and their mean $CVaR_{0.95}$ values were practically the same, while MCE's cumulative return was more than twice that of MV-K. MCE was able to meaningfully exceed MV-K's outcomes using the same set of assets, meaning that the difference would have come from weights optimisation.

Figure 7

Tech universe portfolio $CVaR_{0.95}$ across semi-annual rebalancing windows



Author: compiled by the author.

As expected, MV models exhibited characteristic weights concentration, summarised in Tables 9 and 10. Since technology is a single sector, the “defensive” label would not be entirely accurate here as it would be in a broad market asset sample. It

would be more accurate to say that MV-K tilted toward relatively more stable assets with comparatively lower volatility (variance) and lower covariance.

Table 9

Tech universe MV-K holdings count by rebalance date

Rebalance date	Weights > 0	Weights above threshold	Rebalance date	Weights > 0	Weights above threshold
2015-07-01	16	10	2021-01-04	11	8
2016-01-04	12	11	2021-07-01	18	11
2016-07-01	16	11	2022-01-03	18	13
2017-01-03	14	12	2022-07-01	17	8
2017-07-03	17	14	2023-01-03	8	8
2018-01-02	20	15	2023-07-03	15	10
2018-07-02	15	12	2024-01-02	20	13
2019-01-02	10	10	2024-07-01	18	11
2019-07-01	20	13	2025-01-02	20	15
2020-01-02	18	9	2025-07-01	15	10
2020-07-01	13	9	2026-01-02	19	11

Author: compiled by the author.

Note. Portfolio holdings counts are reported using two thresholds: weights greater than 0, and weights greater than 10^{-12} .

On average, across all holding periods, the top three positions in MV-K made up approximately 57%, and the top five positions approximately 77% of portfolio weights. This indicates concentration characteristic of MV models in general. Since the MV-K portfolio is constructed from Stage 1-screened assets, MV-All results are more informative about the tech universe sample.

Table 10

Tech universe MV-All holdings count by rebalance date

Rebalance date	Weights > 0	Weights above threshold	Rebalance date	Weights > 0	Weights above threshold
2015-07-01	103	24	2021-01-04	113	16
2016-01-04	108	108	2021-07-01	143	24
2016-07-01	123	19	2022-01-03	143	26
2017-01-03	113	24	2022-07-01	161	16
2017-07-03	94	94	2023-01-03	216	16
2018-01-02	125	125	2023-07-03	140	22
2018-07-02	115	31	2024-01-02	108	30
2019-01-02	71	71	2024-07-01	192	35
2019-07-01	142	21	2025-01-02	230	34
2020-01-02	126	126	2025-07-01	236	27
2020-07-01	83	17	2026-01-02	222	22

Author: compiled by the author.

Note. Portfolio holdings counts are reported using two thresholds: weights greater than 0, and weights greater than 10^{-12} .

MV-All showed similar behaviour to MV-K, although on a larger scale. On average, a total of 135 assets were excluded by the missing observation rule across all periods, but later periods contained more eligible assets, rising to 250 in the end. 75 previously excluded assets in earlier periods were later admitted when these assets had sufficient observations again. On average across all holding periods, the top three positions in MV-All made up about 36%, and the top five about 51% of portfolio weights. MV-All was therefore broad in nominal terms, but still materially concentrated in economic terms.

Similarly to MV-All, data shows that MCE's Stage 1 screening chose relatively stable companies, like established enterprise software, business systems, networking, and infrastructure, while momentum-driven speculative names were largely not present. All observed MCE portfolios in the 21 periods were equally weighted because the CVaR cap optimisation constraint did not bind in any period. The main difference in MCE and EW results comes from asset selection protocol differences. The average overlap across all holding periods between assets in MCE and EW portfolios was 3.09 out of 20 or in percentage terms, 15.45%.

Mean portfolio turnover was high for all models, with MV-All trading 61.63%, EW 64.69%, MCE 67%, and MV-K 82.79% of portfolio value on average per rebalance (see Table 8). Figure 8 further illustrates how MV-K's turnover regularly exceeded 90% (in 8 periods out of total 21), and MV-All exhibited sharp turnover spikes, while MCE and EW tended to shuffle portfolios relatively less. If lower turnover is treated as a positive factor when transaction costs are accounted for, MCE's higher relative stability can be seen as an advantage over MV-All, especially considering that the final wealth was essentially the same for the two portfolios, albeit achieved by different paths.

Figure 8

Tech universe portfolio turnover across semi-annual rebalancing windows



Author: compiled by the author.

Comparing MCE and MV-K, which both construct a portfolio with the same exact 20 asset names, shows that the models produced markedly different outcomes. MCE delivered more than twice the cumulative return of MV-K (315.50% versus 153.04%), with higher mean semi-annual return, higher Sharpe and Sortino ratios, and lower turnover, while the two portfolios had almost identical mean $CVaR_{0.95}$. The evidence here suggests that the covariance-based MV-K allocation compressed the same basket of assets into overly-concentrated positions without buying a meaningful tail-loss risk advantage in return performance.

The tech universe results show that no portfolio dominated on all metrics in the sample. EW was the strongest realised performer by a wide margin, delivering the highest final cumulative return and mean semi-annual return, and the highest information ratio against SPY, but it did so with the highest mean $CVaR_{0.95}$ and without a turnover advantage. Among the other portfolios, MV-All was the most balanced: it achieved almost identical terminal wealth to MCE, while attaining a slightly higher

Sharpe ratio, the lowest mean $CVaR_{0.95}$, and the lowest turnover of all portfolio strategies. MCE remained competitive in return terms, but its risk-adjusted performance was weaker than that of EW and MV-All, and its turnover was higher than MV-All's. MV-K performed worst overall, with the lowest cumulative return, the weakest risk-adjusted performance, and the highest turnover. Overall, in the tech universe, EW captured the strongest upside, whereas MV-All provided the strongest practical compromise.

Taken together, the ETF- and technology-universe results indicate that at least with the particular set of models and data used in this study, portfolio performance appears to depend materially on the selection of asset universe. Stage 1 selection appears to have more of an effect than Stage 2 optimisation. The following discussion compares the results across the two universes to highlight performance differences.

3.3. Discussion

Four alternative portfolio construction models—MCE, EW, MV-K, MV-All—were tested and their performance analysed in two out-of-sample simulation scenarios on (i) a broad selection of U.S. ETFs, and (ii) single-sector U.S. technology shares. There were notable differences between model outcomes depending on the asset universe, and this subchapter will bring the results together to discuss where the most important differences and similarities were found, and what the possible causes might be. The discussion considers the returns performance, risk performance, and turnover of all four models.

Results indicate that model performance seems to depend materially on the underlying asset universe. This agrees with Markowitz's (1959) claim about the importance of asset selection in the portfolio construction process. An optimiser works only with the inputs provided to it. Out of the four portfolio models, MCE achieved the highest cumulative return in the ETF universe and second highest in the tech universe. Final returns are not directly comparable between the universes, and caution should be taken in attempting to generalise outside the study sample, but the processes that led to the results appear to be informative with regard to model performance.

In the ETF universe, MCE outperformed EW, which was somewhat unexpected given that there is broad empirical evidence that risk-return-optimised portfolios do not consistently outperform unoptimised EW out of sample (DeMiguel *et al.*, 2009). MCE

and EW have substantial differences in their respective portfolio construction protocols. MCE screens and ranks the asset universe across the estimation window using mean return over a MAD score, both time-weighted (7), and employs further mean-MAD gating, but EW ranks assets only by mean return in the same estimation window. This means that MCE by design prefers higher past returns but penalises large price swings, while EW only chooses by highest past returns. In cumulative-return terms, this appears to have benefitted MCE in the ETF universe, but penalised the model in the tech universe, where MCE's final wealth was less than half that of EW's (54.74% lower). The largest increase in EW portfolio value occurred in the last holding period, where there were a few very large gainers, but also two assets that lost half of their previous-period value. The MCE portfolio value in the last period was net positive.

The two MV variants underperformed EW in both universes in cumulative returns, which is in line with existing research on the conservative nature characteristic to MV (Plyakha *et al.*, 2017; DeMiguel *et al.*, 2009). However, in the tech universe, MV-All cumulative return outperformed MCE by 1.7% in absolute terms. In the tech universe, mean semi-annual returns were higher for all models than in the ETF universe. MCE had the highest mean semi-annual return in the ETF universe and second highest in the tech universe. MV-All produced markedly different outcomes in each, and MV-K performed the worst in both universes.

Risk-adjusted returns performance between asset universes also indicates structural differences in model input asset selection. MCE's mean annualised Sharpe and Sortino ratio ranking placed it third in both universes. MV-K was on top in the ETF universe but came last in the tech universe. MV-All had the lowest risk-adjusted returns in the ETF universe. However, in the tech universe it achieved second-highest risk-adjusted returns. The SPY benchmark underperformed all but the MV-K model in the tech universe, and outperformed all portfolios in the ETF universe.

Mean $CVaR_{0.95}$ values were rather informative about risk-handling differences between the models. MCE uses an explicit control. It sets the maximum allowable $CVaR_{0.95}$ value in the final portfolio based on the $CVaR_{0.95}$ of an equal-weight portfolio composed of all eligible assets in the lookback window. The result of this design choice is that the constructed portfolio cannot be riskier in-sample than an EW portfolio constructed from all assets in the eligible universe. Comparing MCE to MV mod-

els, it can be seen that the MV models concentrated on lower variance and covariance, and landed on more defensive or stable names, as should be expected from MV. MV variants had overall the lowest mean $CVaR_{0.95}$ in both universes. EW had the highest mean $CVaR_{0.95}$.

In both universes, MCE had consistently lower mean $CVaR_{0.95}$ than EW. This can be explained both by initial asset screening differences and by MCE's downstream risk controls. MCE and EW have similar weighting doctrines where EW portfolios are always equally weighted and MCE aims for equal weights, but only if the $CVaR_{0.95}$ cap allows it.

Across the entire simulation timeline, audit logs show that MCE was not practically (within a very small numerical tolerance) equally weighted in only seven periods in the ETF universe. MCE portfolios were practically equally weighted in every single tech universe holding period. If the $CVaR_{0.95}$ cap does not bind, entropy is able to spread out all asset weights equally. If the $CVaR_{0.95}$ binds, assets with higher $CVaR_{0.95}$ in the estimation window are given smaller weights, and assets with lower $CVaR_{0.95}$ get larger weights. Therefore, when the CVaR constraint is slack, MCE produces equal-weight portfolios using assets filtered by Stage 1. When the constraint is not binding, the CVaR component does not change portfolio weights actively but this is intentional: the purpose is to be a loss-risk backstop for when the possible losses of a selected basket of assets exceed that of the eligible universe in the evaluation window. In empirical results and conditional to the specific backtesting, MCE's risk controls in the asset selection and weighting stages can be seen to have tempered tail-loss risk sufficiently to reduce losses compared to EW, while still allowing for higher returns than MV models in all but one case (tech universe MV-All). Therefore, the inclusion of a CVaR-based tail-loss measure in the MCE portfolio construction model appears to be supported with the important caveat that no alternatives were tested. This does not disagree with Righi and Borenstein (2018)'s claim about risk measures being situational, however, no other risk measure was tested within the MCE context so further conclusions cannot be drawn. The last performance metric evaluated and compared was portfolio turnover. There were differences between universes here as well.

In the ETF universe, MCE had consistently lower turnover rate than the other models, while MV models had the highest and EW was in-between. In the tech uni-

verse, all models produced high-turnover portfolios, and MCE came in third after MV-All and EW, although turnover differences between the three models were not large. MV-K had the highest turnover in both universes. A turnover comparison between MV-K and MCE is apt, since MV-K works with the exact same list of asset names as MCE, but uses an entirely different weights optimisation strategy. In the ETF universe, MCE's average turnover was 22.15% per rebalance, while MV-K's was 62.86%. In the tech universe, turnover rates for MCE and MV-K were respectively 67.27% and 82.79%. Earlier discussion shows that MV-K's cumulative return was below MCE in both universes. This allows for a tentative conclusion that when given the same list of assets, MCE may be able to outperform a classic MV implementation in returns, but obviously not under all circumstances.

In the tech universe, between MCE and MV-All, MCE actually had a higher average turnover by a relative 8.39%. MCE's final wealth was also lower by a relative 1.69%. However, MCE's turnover was more stable from period-to-period, illustrated best in Figure 8. Given that MCE is essentially EW with added risk-return protocols, these findings are tentatively in line with De Nard *et al.* (2021), who found EW to have considerably lower turnover than MV-based factor models. DeMiguel *et al.* (2009) also report substantially higher turnover for models that use optimisation vs EW. If trading costs are a concern in practice, lower turnover can imply lower costs, although this claim was not directly tested. MCE performed relatively well in turnover terms, which offers support for its use in practice. Since all models had on average considerably higher turnover in the tech universe than in the ETF universe, portfolio compositions need further examination.

The ETF universe was intentionally assembled to provide broad cross-sector exposure, whereas the technology universe was a narrow single-sector sample of individual shares. This difference matters because a broader universe can provide more scope for reducing unsystematic risk, even with simple diversification protocols, while a narrow universe restricts such efforts. In the ETF universe, both MCE and EW appear to have benefitted from the more economically diverse opportunity set, since MCE had on balance the best performance out of the models, and EW was noticeably more stable out-of-sample than in the tech universe. Some of the differences in outcomes between the two can likely be explained by how portfolio compositions were affected by asset

selection.

The average asset overlap across MCE and EW portfolios was higher in the ETF universe than it was in the tech universe. In the ETF universe, MCE and EW landed on the same assets more than half the time (56.38%), whereas tech portfolios were only barely similar (15.45%). As has been previously mentioned, MCE chooses assets with high recent returns relative to average absolute price fluctuations (MAD-based score), and then further prefers assets with lower price volatility (median MAD gate). EW, by contrast, only looks at past returns irrespective of price swings. This is the essence of naïve weighting, and it has been noted as a downside in literature. For example, Kritzman *et al.* (2010) do not argue against EW as a concept, but they warn against not including any investment knowledge, such as expectations about returns and risk in asset selection. Here, the tech universe results show that EW, without any expectations about potential losses, was noticeably more unstable in both returns and risk performance when compared to MCE. It was less evident in the ETF universe, but this only serves to highlight the differences between the universes and the overall importance of the pre-allocation asset selection process, than it does about the differences between the models.

Cardinality-sensitivity tests further illustrate the influence of asset selection (see Appendix D). In the ETF universe, lower cardinality increased returns but also increased turnover, and the CVaR constraint was binding more frequently the more concentrated the portfolio was. In the tech universe, Stage 2 constructed equal-weight portfolios for every single tested cardinality. Overall, in the combination of data and parameters selected, MCE appeared to be primarily driven by Stage 1 screening, with a conditional CVaR-based tail-risk backstop.

Platanakis *et al.*'s (2021) conclusions about EW being more suitable for asset class selection also appear to be supported by backtest results. ETFs are pooled assets, which can smooth return variance. Between the two universes, the mean $CVaR_{0.95}$ gaps between MCE and EW were different: 1.92% and 2.42% in the ETF universe, versus 3% and 4.63% in the tech universe.

MCE versus MV-K results were interesting in how they cleanly illustrated the difference in the models' weighting methodology, since both optimise within the same list of asset names. Across both universes, MCE made better practical use of the 20 as-

sets: in the ETF universe it traded some downside defensiveness for higher cumulative return and much lower turnover, while in the technology universe it dominated MV-K on return, risk-adjusted return, and turnover with almost no sacrifice in mean tail-loss control. Once Stage 1 had selected a fixed candidate set, the entropy-based MCE weighting rule appears to have preserved diversification and implementability better than covariance-based MV-K.

As has been discussed previously, MV models in practice have evolved to be much more complex than the original Markowitz (1952, 1959) design, incorporating ever more prior information, and yet they still tend to perform poorly out-of-sample compared to EW (Palomar, 2025). The backtesting results here place MCE as a strong practical compromise between the models considered, in the context of the historical data used. MCE achieved the highest cumulative return in the ETF universe and a relatively high one in the tech universe, compared to MV variants. Cumulative return in the tech universe was however much lower compared to EW, but it improved on EW's outcomes in other metrics. In both universes, MCE had lower mean tail-loss than EW. MCE also had a lower turnover rate than any other model in the ETF universe, and a relatively comparable turnover to EW and MV-All. Risk-reward trade-offs can be seen across all results. Overall, MCE did what it aimed to do: it is simple to implement, and empirical testing put it in the middle ground between naïve equal weighting and covariance-based MV optimisation. However, the results should realistically only be interpreted within the context of the particular empirical backtesting protocol and data sample.

Conclusions

This thesis developed and tested a two-stage portfolio construction model that combines mean return and mean absolute deviation-based screening with entropy maximisation under a conditional value-at-risk constraint. The model's aim was to address a practical gap in the portfolio construction literature: many optimisation-based models improve stability or risk control but do so at the expense of substantial methodological and implementation complexity, while simple equal-weight portfolios remain difficult to outperform in realised out-of-sample returns. The proposed model seeks a middle ground by separating asset selection from weight optimisation, enforcing an exact portfolio size without cardinality optimisation, and incorporating explicit tail-loss risk control.

Empirical findings show that the model's performance in backtests depended materially on the asset universe. In the broad ETF universe, MCE provided the strongest practical compromise among the models tested. It achieved the highest cumulative return and did so with lower mean $CVaR_{0.95}$ and substantially lower turnover. In the technology-share universe, the same pattern did not hold. There, EW captured the strongest upside by a wide margin, while MV-All produced almost identical terminal wealth to MCE with lower mean $CVaR_{0.95}$ and lower turnover. These results indicate that the usefulness of MCE may be conditional rather than universal. Advantages were more visible in the broader and more heterogeneous ETF universe than in the narrower and more momentum-driven tech universe.

The findings also show that portfolio construction outcomes were driven not only by the choice of weighting rule, but by the interaction between asset selection, risk control, and the structure of the asset universe. In the ETF universe, MCE and EW frequently selected overlapping baskets, but MCE's screening method produced a more stable and lower-risk subset without sacrificing final wealth. In the tech universe, overlap between MCE and EW was much lower, and the MCE screen systematically filtered out many of the most extreme return swings that later ended up driving the strongest realised gains in EW. This helps explain why the same model could outperform EW in one universe while materially lagging behind in another. The thesis therefore does not identify a universally superior portfolio construction rule. Rather, it shows how the

specific tractable two-stage design behaves under two materially different empirical settings, and under what conditions its trade-offs may be attractive to investors.

The motivation for developing the MCE model came from a gap found at the intersection of academic research and practical portfolio construction. The main outcome of this thesis is not a new best portfolio construction model, but rather the conclusion of trying to see if the learnings from the last seventy years of modern portfolio theory can be applied in a way that is within the reach of practically-oriented investors and analysts.

Because most of the literature drawn on by the model is relatively recent, this thesis makes a modest incremental contribution to the understanding of how mean and mean absolute deviation-based screening, conditional value-at-risk, and entropy can be combined under a strict cardinality rule in portfolio construction. Compared to mean-variance optimisation, the model simplifies computation significantly, gives greater emphasis to realised return performance alongside risk control, and makes it straightforward to construct a portfolio with an exact number of assets.

From a practical perspective, the results also indicate where the model appears more useful and where its limitations become more visible. The MCE model performed best in the relatively more heterogeneous ETF universe, where it achieved the strongest overall practical compromise between return, tail-loss risk, and turnover. In the technology universe, return performance was considerably weaker relative to EW, although MCE remained competitive with, or superior to, MV alternatives. Given the limited scope of empirical backtesting, these conclusions should be taken as conditional and tentative.

This study is mainly exploratory and has several limitations. The analysis is based on two specific U.S. asset universes, so the findings should be interpreted as conditional on the datasets and design choices used. The data likely contains survivorship bias because historical availability depends on what remains present in Yahoo Finance at the time when the data is fetched. Transaction costs were not modelled, although turnover was reported as relevant to implementation. The empirical comparisons also depend on the chosen parameter values, including lookback window length, time-decay half-life used in weighting, rebalance frequency, CVaR confidence level, and fixed cardinality K choice. In addition, the study did not test alternative covariance estimators,

alternative Stage 1 screening or scoring protocols, or alternative asset universes. Formal statistical inference is limited to exploratory paired tests of semi-annual holding-period return differences between the models. No inference is performed on cumulative wealth, CVaR, turnover, Sharpe ratios, Sortino ratios, or information ratios. These limitations do not invalidate the findings, but they do narrow the scope of conclusions that can be drawn.

Further research could extend the developed MCE portfolio construction model in several directions. The model could be tested on additional asset universes to assess the sensitivity of results to inputs. Statistical inference could be extended to risk-adjusted metrics, tail-risk measures, turnover, and transaction-cost-adjusted performance. Cardinality-sensitivity testing could be extended, especially for larger asset universes. The main design parameters could be varied, especially portfolio size, rebalance frequency, lookback window, time-decay half-life, and CVaR confidence level, in order to determine whether the observed trade-offs are robust. The first-stage selection rule could be compared against alternative simple screening rules, for example unweighted MAD ratios, downside-only measures, median-return, or alternative risk-adjusted ranking statistics. Explicit transaction-cost modelling would be valuable in order to improve the practical interpretation of turnover differences between portfolio construction models.

References

1. Almeida, R. J., Basturk, N., & Rodrigues, P. (2023). Portfolio Return Maximization using Robust Optimization and Directional Changes. *2023 IEEE Symposium Series on Computational Intelligence, SSCI 2023*. DOI: 10.1109/SSCI52147.2023.10371969
2. Artzner, P., Delbaen, F., Eber, J.-M., & Heath, D. (1999). Coherent measures of risk. *Mathematical Finance*, 9(3)
3. Bera, A. K., & Park, S. Y. (2008, January 1). Optimal Portfolio Diversification Using the Maximum Entropy Principle. *ECONOMETRIC REVIEWS*, 27(4/6). DOI: 10.1080/07474930801960394
4. Chang, T. J., Meade, N., Beasley, J. E., & Sharaiha, Y. M. (2000). Heuristics for cardinality constrained portfolio optimisation. *COMPUTERS AND OPERATIONS RESEARCH*, 27(13). DOI: 10.1016/S0305-0548(99)00074-X
5. Chen, D., Wu, Y., Li, J., Ding, X., & Chen, C. (2023). Distributionally robust mean-absolute deviation portfolio optimization using Wasserstein metric. *Journal of Global Optimization*, 87(2). DOI: 10.1007/s10898-022-01171-x
6. DeMiguel, V., Garlappi, L., & Uppal, R. (2009). Optimal Versus Naive Diversification: How Inefficient is the 1/N Portfolio Strategy? *REVIEW OF FINANCIAL STUDIES*, 22(5). DOI: 10.1093/rfs/hhm075
7. De Nard, G., Ledoit, O., & Wolf, M. (2021). Factor Models for Portfolio Selection in Large Dimensions: The Good, the Better and the Ugly. *JOURNAL OF FINANCIAL ECONOMETRICS*, 19(2). DOI: 10.1093/jjfinec/nby033
8. Fabozzi, F. J., Huang, D., & Zhou, G. (2010). Robust portfolios: contributions from operations research and finance. *ANNALS OF OPERATIONS RESEARCH*, 176(1). DOI: 10.1007/s10479-020-03630-8
9. Georgantas, A., Doumpos, M., & Zopounidis, C. (2024). Robust optimization approaches for portfolio selection: a comparative analysis. *Annals of Operations Research*, 339(3). DOI: 10.1007/s10479-021-04177-y

10. Goodwin, T. H. (1998), The Information Ratio. *Financial Analysts Journal*, 54(4). DOI: 10.2469/faj.v54.n4.2196
11. Jobst, N. J., Horniman, M. D., Lucas, C. A., Mitra, G. (2001). Computational aspects of alternative portfolio selection models in the presence of discrete asset choice constraints. *Taylor & Francis Journals, Quantitative Finance*, 1(5). DOI: 10.1088/1469-7688/1/5/301
12. Kahneman, D., & Tversky, A. (1979). Prospect Theory: An Analysis of Decision under Risk. *Econometrica*, 47(2). DOI: 10.2307/1914185
13. Kim, J. H., Kim, W. C., Kwon, D.-G., & Fabozzi, F. J. (2018). Robust equity portfolio performance. *Annals of Operations Research*, 266(1/2). DOI: 10.1007/s10479-017-2739-1
14. Kim, J., Lee, Y., Kim, W., Kang, T., & Fabozzi, F. (2024). An Overview of Optimization Models for Portfolio Management. DOI: 10.3905/jpm.2024.1.643
15. Konno, H. & Yamazaki, H. (1991). Mean-Absolute Deviation Portfolio Optimization model and Its Applications to Tokyo Stock Market. *Management Science*, 37(5). DOI: 10.1287/mnsc.37.5.519
16. Kritzman, M., Page, S., & Turkington, D. (2010). In defense of optimization: The fallacy of $1/N$. *Financial Analysts Journal*, 66(2). DOI: 10.2469/faj.v66.n2.6
17. Ledoit, O., & Wolf, M. (2017). Nonlinear Shrinkage of the Covariance Matrix for Portfolio Selection: Markowitz Meets Goldilocks. *Review of Financial Studies*, 30(12). DOI: 10.1093/rfs/hhx052
18. Ledoit, O., & Wolf, M. (2004a). A well-conditioned estimator for large-dimensional covariance matrices. *Journal of Multivariate Analysis*, 88(2). DOI: 10.1016/S0047-259X(03)00096-4
19. Ledoit, O., & Wolf, M. (2004b). Honey, I Shrunk the Sample Covariance Matrix. *JOURNAL OF PORTFOLIO MANAGEMENT*, 30(4). Accessed 01.03.2026, <http://www.ledoit.net/honey.pdf>

20. Liagkouras, K., Metaxiotis, K., & Tsihrintzis, G. (2022). Incorporating environmental and social considerations into the portfolio optimization process. *Annals of Operations Research*, 316(2). DOI: 10.1007/s10479-020-03554-3
21. Lorimer, D. A., van Schalkwyk, C. H., & Szczygielski, J. J. (2024). Portfolio optimisation using alternative risk measures. *Finance Research Letters*, 67. DOI: 10.1016/j.frl.2024.105758
22. Markowitz, H. (1952). Portfolio Selection. *The Journal of Finance*, 7(1), 77-91. DOI: 10.2307/2975974
23. Markowitz, H. M. (1971). *Portfolio Selection: Efficient Diversification of Investments*. New Haven: Yale University Press. Accessed 03.10.2025, <http://ebookcentral.proquest.com/lib/tartu-ebooks/detail.action?docID=3421285>
24. Moon, Y. & Yao, T. (2011). A robust mean absolute deviation model for portfolio optimization. *Computers & Operations Research*, 38(9). DOI: 10.1016/j.cor.2010.10.020
25. J.P. Morgan and Reuters (1996). RiskMetrics™: Technical Document (Fourth Edition). Accessed 21.10.2025, <https://www.msci.com/documents/10199/5915b101-4206-4ba0-aee2-3449d5c7e95a>
26. Palomar, D. P. (2025). *Portfolio Optimization: Theory and Application*. Cambridge University Press. Accessed 30.01.2026, <https://portfoliooptimizationbook.com/portfolio-optimization-book.pdf>
27. Platanakis, E., Sutcliffe, C., & Ye, X. (2021). Horses for courses: Mean-variance for asset allocation and 1/N for stock selection. *European Journal of Operational Research*, 288(1). DOI: 10.1016/j.ejor.2020.05.043
28. Plyakha, Y., Uppal, R., Vilkov, G. (2017). Why Does an Equal-Weighted Portfolio Outperform Value- and Price-Weighted Portfolios? *Econometric Modeling: Capital Markets*. DOI: 10.2139/ssrn.2724535
29. Righi, M. B., & Borenstein, D. (2018). A simulation comparison of risk measures for portfolio optimization. *Finance Research Letters*, 24. DOI: 10.1016/j.frl.2017.07.013

30. Rockafellar, R. T., Uryasev, S. (2000). Optimization of conditional value-at-risk. *JOURNAL OF RISK*. 2(3),
Accessed 02.10.2025, <https://sites.math.washington.edu/~rtr/papers/rtr179-CVaR1.pdf>
31. Rockafellar, R., Uryasev, S., & Zabarankin, M. (2006). Generalized deviations in risk analysis. *Finance & Stochastics*, 10(1). DOI: 10.1007/s00780-005-0165-8
32. Sauga, A. (2017). *Statistika õpik majanduseriala üliõpilastele*. Tallinn: TTÜ Kirjastus. Accessed 08.09.2025, <https://digikogu.taltech.ee/et/Download/eea31a2d-01a2-43a0-8ee2-86bb293a6210>
33. Sehgal, R., & Mehra, A. (2021). Robust reward-risk ratio portfolio optimization. *Int. Trans. Oper. Res.*, 28. DOI: 10.1111/itor.12652
34. Sharpe, W. F. (1963). A Simplified Model for Portfolio Analysis. *Management Science*, 9(2). Accessed 28.02.2026, <https://www-jstor-org.ezproxy.utlib.ut.ee/stable/2627407>
35. Silva, L. P. da, Alem, D., & Carvalho, F. L. de. (2017). Portfolio optimization using Mean Absolute Deviation (MAD) and Conditional Value-at-Risk (CVaR). *Production*, 27. DOI: 10.1590/0103-6513.208816
36. Usta, I., & Kantar, Y. M. (2011). Mean-variance-skewness-entropy measures: A multi-objective approach for portfolio selection. *Entropy*. 13.
DOI: 10.3390/e13010117
37. Xidonas, P., Steuer, R., & Hassapis, C. (2020). Robust portfolio optimization: a categorized bibliographic review. *Annals of Operations Research*.
DOI: 10.1007/s10479-020-03630-8
38. Zhou, R., Cai, R., & Tong, G. (2013). Applications of Entropy in Finance: A Review. *Entropy*, 15(11). DOI: 10.3390/e15114909

Appendices

Appendix A: Coherence Axioms for Risk Measures

Artzner *et al.* (1999) formalise minimal conditions for a risk measure to be adequate for investors or regulators. Artzner *et al.* (1999) call such risk measures coherent and define the following axioms:

Let X and Y be random future portfolio values, and $\rho(X)$ and $\rho(Y)$ respectively the risk measures of that portfolio. A risk measure $\rho(X)$ is coherent if it satisfies the following four conditions (Artzner *et al.*, 1999):

1. Translation invariance: $\rho(X + a) = \rho(X) - a$, for every real constant a . Adding a risk-free asset a (e.g. cash) to the portfolio decreases the investor's overall risk precisely by a .
2. Subadditivity: $\rho(X + Y) \leq \rho(X) + \rho(Y)$. The total risk of a portfolio does not exceed the sum of the individual risks of its components. Diversification does not increase the portfolio's total risk.
3. Positive homogeneity: $\rho(\lambda X) = \lambda\rho(X)$, where $\lambda \geq 0$. Increasing a position by any non-negative real factor λ increases the risk of the portfolio by the same λ .
4. Monotonicity: if for every possible X and Y , $X \leq Y$, then $\rho(Y) \leq \rho(X)$. If portfolio Y has a future value that is at least as high as X , then portfolio Y should not be assessed as riskier than portfolio X .

Without translation invariance, adding risk-free assets to the portfolio would not reduce the portfolio's total risk. Without subadditivity diversification could increase a portfolio's risk instead of reducing it. Without positive homogeneity, positions could be multiplied without multiplying risks at the same time. And without monotonicity, a portfolio that is never worse and is sometimes better than another, could still be assigned a higher risk value. Therefore, the four coherence axioms describe a minimal set of properties that an adequate risk measure should reasonably have, according to Artzner *et al.* (1999).

Table 11
Coherence of commonly used risk measures

Risk measure	Coherent?	Translation invariance	Subadditivity	Positive homogeneity	Monotonicity
Standard deviation σ	No	No	Yes	Yes	No; it is possible that $X \geq Y$ but $\sigma(X) > \sigma(Y)$.
Variance σ^2	No	No	Not guaranteed	No; it scales as λ^2	No; it is possible that $X \geq Y$ but $\sigma^2(X) > \sigma^2(Y)$.
MAD	No	No	Yes	Yes	No; it is possible that $X \geq Y$ but $MAD(X) > MAD(Y)$.
VaR	No	Yes	No, in general	Yes	Yes
CVaR	Yes	Yes	Yes	Yes	Yes

Author: compiled by the author.

Note. Rockafellar *et al.* (2006) for standard deviation, variance, and MAD; Artzner *et al.* (1999) for VaR and CVaR.

Appendix B: Sharpe, Sortino, and Information Ratio Definitions

Let $r_{p,t}$ denote the daily return of portfolio p on day t , $r_{f,t}$ the daily risk-free rate, and $r_{b,t}$ the daily return of the benchmark portfolio. In this thesis, SPY is used as the benchmark portfolio. Let $D = 252$ be the number of trading days in a year.

Final cumulative return measures the compounded return over the full evaluation period:

$$R_p^{\text{cum}} = \prod_{t=1}^T (1 + r_{p,t}) - 1.$$

Mean semi-annual return is calculated by first compounding daily portfolio returns within each semi-annual holding period and then averaging those holding-period returns:

$$\bar{R}_p^{\text{semi}} = \frac{1}{H} \sum_{h=1}^H \left[\prod_{t \in h} (1 + r_{p,t}) - 1 \right],$$

where H is the number of semi-annual holding periods.

Annualised Sharpe ratio measures mean excess return over the risk-free rate per unit of total excess-return volatility (Plyakha *et al.*, 2017):

$$SR_p = \sqrt{D} \frac{\bar{e}_p}{\sigma(e_p)},$$

where $e_{p,t} = r_{p,t} - r_{f,t}$ is the portfolio's daily excess return, $\bar{e}_p$ is the mean daily excess return, and $\sigma(e_p)$ is the standard deviation of daily excess returns.

Annualised Sortino ratio replaces total volatility with downside deviation from the risk-free rate (Plyakha *et al.*, 2017):

$$SoR_p = \sqrt{D} \frac{\bar{e}_p}{DR_p},$$

where downside deviation is calculated from negative excess-return observations:

$$DR_p = \sqrt{\frac{1}{N_p^-} \sum_{t:e_{p,t} < 0} e_{p,t}^2},$$

and N_p^- is the number of days on which $e_{p,t} < 0$. Sortino ratio penalises only downside deviations from the chosen risk-free rate.

Annualised information ratio measures active return relative to the chosen benchmark per unit of active-return volatility, or tracking error (Goodwin, 1998):

$$IR_p = \sqrt{D} \frac{\bar{a}_p}{\sigma(a_p)},$$

where $a_{p,t} = r_{p,t} - r_{b,t}$ is the portfolio's daily active return relative to the benchmark, $\bar{a}_p$ is the mean daily active return, and $\sigma(a_p)$ is the standard deviation of daily active returns.

Appendix C: Paired *t*-tests and Paired Wilcoxon Signed-Rank Test Results

Table 12

Paired tests on semi-annual holding-period returns in the ETF universe

Comparison	<i>n</i>	Mean diff. (p.p.)	Median diff. (p.p.)	<i>t</i> -test <i>p</i> -value	Wilcoxon <i>p</i> -value
MCE – EW	39	0.045	0.324	0.949	0.576
MCE – MV-K	39	0.979	0.900	0.321	0.343
MCE – MV-All	39	1.960	2.107	0.079	0.033
MCE – SPY	39	-1.391	-1.589	0.070	0.039

Note. Differences are calculated as MCE minus the comparison portfolio using semi-annual holding-period returns. Mean and median differences are reported in percentage points.

Author: compiled by the author.

Paired *t*-tests and paired Wilcoxon signed-rank tests provide only limited support for return differences between MCE and the comparison portfolios in the ETF

universe. The return difference between MCE and EW is economically small and statistically indistinguishable from zero in both tests. The same is true for MCE relative to MV-K. For MCE relative to MV-All, the mean semi-annual return difference is positive at 1.96 percentage points, with the paired t -test not below the 5% significance level ($p = 0.07878$), while the Wilcoxon signed-rank test is below that threshold ($p = 0.03344$). This suggests some evidence that MCE more frequently achieved higher semi-annual returns than MV-All, but the evidence is not uniform across the test specifications.

The comparison with SPY goes in the opposite direction. The mean semi-annual return difference is negative at -1.39 percentage points. As with the MV-All comparison, the paired t -test is not below the 5% significance level ($p = 0.06963$), while the Wilcoxon signed-rank test is below it ($p = 0.03858$). This indicates that SPY tended to outperform MCE across holding periods in this test, even though the main performance table also considers other criteria such as CVaR, turnover, and risk-adjusted ratios.

These results should be interpreted cautiously. The tests are exploratory, use a limited number of non-overlapping semi-annual holding periods, and evaluate only paired holding-period return differences. They do not test final cumulative wealth, CVaR, turnover, Sharpe ratio, Sortino ratio, or full model superiority. The results therefore support the thesis's cautious interpretation: MCE's empirical performance in the ETF universe is conditional on the selected universe, model specification, and evaluation criteria, rather than evidence of unconditional dominance over the alternatives.

Table 13

Paired tests on semi-annual holding-period returns in the tech universe

Comparison	n	Mean diff. (p.p.)	Median diff. (p.p.)	t -test p -value	Wilcoxon p -value
MCE – EW	21	-4.974	-3.566	0.202	0.216
MCE – MV-K	21	2.420	1.390	0.144	0.179
MCE – MV-All	21	0.039	-0.116	0.975	0.973
MCE – SPY	21	0.570	1.211	0.725	0.892

Note. Differences are calculated as MCE minus the comparison portfolio using semi-annual holding-period returns. Mean and median differences are reported in percentage points.

Author: compiled by the author.

The paired tests provide no statistical evidence of systematic semi-annual holding-period return differences between MCE and the comparison portfolios in the technology universe. The mean difference between MCE and EW is negative at -4.97 percentage

points, but neither the paired t -test nor the Wilcoxon signed-rank test is statistically significant at the 5% significance level. MCE has a positive mean difference relative to MV-K, SPY, and a near-zero mean difference relative to MV-All, but these differences are also not statistically significant.

Tech universe results provide weaker support for return differences than ETF universe tests. These results should be interpreted cautiously because the tests use only 21 non-overlapping semi-annual holding periods and evaluate only paired holding-period returns, not cumulative wealth, CVaR, turnover, Sharpe ratios, Sortino ratios, or overall model performance.

Appendix D: MCE Portfolio Cardinality Sensitivity Tests

Table 14

MCE test portfolio ETF universe cardinality sensitivity tests

Cardinality	Periods	Cumulative return (%)	Mean Annualised semiannual return (%)	Sharpe ratio	Annualised Sortino ratio	Annualised information ratio	Mean CVaR _{0.95} (%)	Mean turnover (%)
5	39	534.461	5.236	0.616	0.576	-0.139	1.844	68.975
10	39	520.850	5.139	0.603	0.568	-0.206	1.837	50.600
15	39	440.467	4.801	0.534	0.500	-0.337	1.906	34.918
20	39	384.470	4.564	0.493	0.459	-0.431	1.919	22.146
22	39	332.090	4.270	0.452	0.422	-0.504	1.919	17.875

Author: compiled by the author.

Note. Only 22 assets passed zero-days data quality control in the first period, therefore 25 and higher cardinality could not be tested.

Cumulative return is the final full-sample (backtest scenario) compounded return.

Mean semiannual return is calculated by compounding daily returns within each holding period and then averaging across these periods.

Annualised ratios are calculated from daily returns across the full sample.

Information ratio is calculated with S&P 500 (SPY) as a benchmark.

Mean CVaR is calculated from daily losses within each holding period and then averaged across all holding periods.

Mean turnover is averaged across rebalance dates. SPY is a single-asset benchmark; no CVaR or turnover is reported.

The frequency with which MCE portfolio asset allocation differs from equal-weight distribution declines as cardinality increases. For $K = 20$, only 7 of 39 realised holding-period allocations are materially non-equal-weight, while for $K = 22$ this occurs in only 2 of 39 periods. By contrast, the lower-cardinality specifications $K = 5$, $K = 10$, and $K = 15$ produce non-equal-weight allocations in approximately 28–33% of periods. This indicates that the CVaR constraint is more often binding when the selected portfolio is more concentrated, whereas higher-cardinality portfolios more frequently reduce to equal-weight allocations over the selected assets. This can be seen in Table 15 below.

Table 15

MCE test portfolio ETF universe equal-weight and CVaR-bound allocation frequencies by cardinality

Cardinality	Periods	Equal-weight periods	CVaR-bound periods	Equal-weight share (%)	CVaR-bound share (%)
5	39	28	11	71.795	28.205
10	39	26	13	66.667	33.333
15	39	27	12	69.231	30.769
20	39	32	7	82.051	17.949
22	39	37	2	94.872	5.128

Author: compiled by the author.

Note. Equal-weight periods are periods in which portfolio weights are equal within numerical tolerance.

CVaR-bound periods are periods in which portfolio weights materially differ from equal weights.

Since the Stage 2 objective maximises entropy, equal-weight allocation occurs when feasible under the CVaR constraint.

Equal-weight allocation indicates that the CVaR constraint did not bind at rebalance.

Table 16

MCE test portfolio tech universe cardinality sensitivity tests

Cardinality	Periods	Cumulative return (%)	Mean semiannual return (%)	Annualised Sharpe ratio	Annualised Sortino ratio	Annualised information ratio	Mean CVaR _{0.95} (%)	Mean turnover (%)
5	21	266.174	7.533	0.533	0.495	0.054	3.462	90.476
10	21	379.021	8.572	0.660	0.624	0.227	3.087	78.604
15	21	280.034	7.334	0.568	0.531	0.069	3.112	73.089
20	21	315.496	7.747	0.611	0.572	0.140	3.001	67.275
25	21	362.076	8.222	0.662	0.620	0.236	2.962	64.329
30	21	340.194	7.961	0.646	0.607	0.199	2.913	60.358
35	21	320.546	7.762	0.630	0.589	0.156	2.892	56.221

Author: compiled by the author.

Note. Cumulative return is the final full-sample (backtest scenario) compounded return.

Mean semiannual return is calculated by compounding daily returns within each holding period and then averaging across these periods.

Annualised ratios are calculated from daily returns across the full sample.

Information ratio is calculated with S&P 500 (SPY) as a benchmark.

Mean CVaR is calculated from daily losses within each holding period and then averaged across all holding periods.

Mean turnover is averaged across rebalance dates. SPY is a single-asset benchmark; no CVaR or turnover is reported.

Tech universe cardinality-sensitivity results do not show the same monotonic pattern as the ETF universe. The highest cumulative return was achieved at $K = 10$, while the highest annualised Sharpe ratio and information ratio were achieved at $K = 25$. Lower cardinality did not consistently improve performance here. $K = 5$ had the highest turnover and the weakest risk-adjusted results. $K = 10$ and $K = 25$ performed better across most return and risk-adjusted metrics.

Mean turnover declined as cardinality increased, and mean CVaR_{0.95} generally declined as well, indicating that higher-cardinality portfolios were less exposed to tail-loss concentration and required less rebalancing. The baseline $K = 20$ was not the ex post best-performing cardinality in this dataset, but it remained in the middle of

the tested range and achieved a comparatively balanced outcome between return, tail-loss risk, and turnover. Because the CVaR constraint did not bind in any tested tech universe K -sensitivity specification, these differences reflect Stage 1 asset selection.

Table 17

MCE test portfolio tech universe equal-weight and CVaR-bound allocation frequencies by cardinality

Cardinality	Periods	Equal-weight periods	CVaR-bound periods	Equal-weight share (%)	CVaR-bound share (%)
5	21	21	0	100.000	0.000
10	21	21	0	100.000	0.000
15	21	21	0	100.000	0.000
20	21	21	0	100.000	0.000
25	21	21	0	100.000	0.000
30	21	21	0	100.000	0.000
35	21	21	0	100.000	0.000

Author: compiled by the author.

Note. Equal-weight periods are periods in which portfolio weights are equal within numerical tolerance.

CVaR-bound periods are periods in which portfolio weights materially differ from equal weights.

Since the Stage 2 objective maximises entropy, equal-weight allocation occurs when feasible under the CVaR constraint.

Equal-weight allocation indicates that the CVaR constraint did not bind at rebalance.

In the tech universe, MCE allocation reduced to equal weights in every holding period for all tested cardinalities from $K = 5$ to $K = 35$. This indicates that, conditional on the selected assets and the applied CVaR cap, the CVaR constraint did not bind or alter the entropy-maximisation allocation. These results therefore illustrate the effects of the Stage 1 asset selection protocol design.

Appendix E: Assets Used in Backtesting

Table 18

ETF universe (73, including SPY as benchmark)

Ticker	Name	Detail
SPY	SPDR S&P 500 Trust	Broad large-cap equity
IJH	iShares Core S&P Mid-Cap	Mid-cap equity
IJR	iShares Core S&P Small-Cap	Small-cap equity
VBR	Vanguard Small-Cap Value	Small-cap value
VBK	Vanguard Small-Cap Growth	Small-cap growth
QUAL	iShares MSCI USA Quality Factor	Quality factor
VLUE	iShares MSCI USA Value Factor	Value factor
MTUM	iShares MSCI USA Momentum Factor	Momentum factor
USMV	iShares MSCI USA Min Vol Factor	Minimum-volatility factor
SPLV	Invesco S&P 500 Low Volatility	Low-volatility large cap
SCHD	Schwab U.S. Dividend Equity	Dividend/quality income
XLB	Materials Select Sector SPDR Fund	Sector equity, materials
XLC	Communication Services Select Sector SPDR Fund	Sector equity, communication services
XLE	Energy Select Sector SPDR Fund	Sector equity, energy

Ticker	Name	Detail
XLF	Financial Select Sector SPDR Fund	Sector equity, financials
XLI	Industrial Select Sector SPDR Fund	Sector equity, industrials
XLK	Technology Select Sector SPDR Fund	Sector equity, information technology
XLP	Consumer Staples Select Sector SPDR Fund	Sector equity, consumer staples
XLU	Utilities Select Sector SPDR Fund	Sector equity, utilities
XLV	Health Care Select Sector SPDR Fund	Sector equity, health care
XLY	Consumer Discretionary Select Sector SPDR Fund	Sector equity, consumer discretionary
XLRE	Real Estate Select Sector SPDR Fund	Sector equity, real estate
SMH	VanEck Semiconductor	Thematic equity, semiconductors
XBI	SPDR S&P Biotech	Industry equity, biotechnology
IHI	iShares U.S. Medical Devices	Industry equity, medical devices
XAR	SPDR S&P Aerospace & Defense	Industry equity, aerospace & defense
XTN	SPDR S&P Transportation	Industry equity, transportation
PHO	Invesco Water Resources	Thematic equity, water/water infrastructure
REMX	VanEck Rare Earth and Strategic Metals	Thematic equity, rare earths/strategic metals
TAN	Invesco Solar	Thematic equity, solar energy
LIT	Global X Lithium & Battery Tech	Thematic equity, lithium/battery technology
KRE	SPDR S&P Regional Banking	Industry equity, regional banks
VXUS	Vanguard Total International Stock	International equity, broad global ex-U.S. equity
VEA	Vanguard FTSE Developed Markets	International equity, developed ex-U.S. equity
VWO	Vanguard FTSE Emerging Markets	International equity, emerging-markets equity
IEUR	iShares Core MSCI Europe	Regional equity, Europe
EWJ	iShares MSCI Japan	Single-country equity, Japan
EWU	iShares MSCI United Kingdom	Single-country equity, United Kingdom
EWG	iShares MSCI Germany	Single-country equity, Germany
EWQ	iShares MSCI France	Single-country equity, France
EWA	iShares MSCI Australia	Single-country equity, Australia
EWC	iShares MSCI Canada	Single-country equity, Canada
EWY	iShares MSCI South Korea	Single-country equity, South Korea
EWT	iShares MSCI Taiwan	Single-country equity, Taiwan
INDA	iShares MSCI India	Single-country equity, India
EWZ	iShares MSCI Brazil	Single-country equity, Brazil
EWX	iShares MSCI Mexico	Single-country equity, Mexico
EIDO	iShares MSCI Indonesia	Single-country equity, Indonesia
THD	iShares MSCI Thailand	Single-country equity, Thailand
TUR	iShares MSCI Turkey	Single-country equity, Turkey
GLD	SPDR Gold Shares	Commodity trust, gold
DBC	Invesco DB Commodity Index Tracking Fund	Commodity, broad commodities
DBA	Invesco DB Agriculture Fund	Commodity, agriculture
DBB	Invesco DB Base Metals Fund	Commodity, industrial/base metals
USO	United States Oil Fund, LP	Commodity fund, crude oil
UNG	United States Natural Gas Fund, LP	Commodity fund, natural gas
GDX	VanEck Gold Miners	Equity commodity-theme, gold miners
GDXJ	VanEck Junior Gold Miners	Equity commodity-theme, junior gold miners
VNQ	Vanguard Real Estate	Real estate equity, REITs/real estate
UUP	Invesco DB US Dollar Index Bullish Fund	Currency, long U.S. dollar basket
FXE	Invesco CurrencyShares Euro Trust	Currency trust, euro
FXJ	Invesco CurrencyShares Japanese Yen Trust	Currency trust, Japanese yen
FXF	Invesco CurrencyShares Swiss Franc Trust	Currency trust, Swiss franc
DBMF	iMGP DBi Managed Futures Strategy	Alternative, managed futures/trend-following

Ticker	Name	Detail
KMLM	KraneShares Mount Lucas Managed Futures Index Strategy	Alternative, managed futures/multi-asset trend-following
BTAL	AGF U.S. Market Neutral Anti-Beta Fund	Alternative, market-neutral anti-beta/equity hedge
TAIL	Cambria Tail Risk	Alternative, tail-risk hedge/crash protection
SPTS	State Street SPDR Portfolio Short Term Treasury	Bond, short-term U.S. Treasuries
STIP	iShares 0–5 Year TIPS Bond	Bond, short-term U.S. TIPS
IGSB	iShares 1–5 Year Investment Grade Corporate Bond	Bond, short-duration investment-grade corporates
SPSB	State Street SPDR Portfolio Short Term Corporate Bond	Bond, short-term U.S. corporate bonds
SHYG	iShares 0–5 Year High Yield Corporate Bond	Bond, short-duration high-yield corporates
SUB	iShares Short-Term National Muni Bond	Bond, short-term municipal bonds

Author: compiled by the author.

Table 19

Technology universe (371 shares)

Ticker	Name	Ticker	Name
CACI	CACI International Inc.	GDDY	GoDaddy Inc.
IONQ	IonQ, Inc.	FIG	Figma, Inc.
AMKR	Amkor Technology, Inc.	JKHY	Jack Henry & Associates, Inc.
RMBS	Rambus Inc.	NTNX	Nutanix, Inc.
IT	Gartner, Inc.	ONTO	Onto Innovation Inc.
SITM	SiTime Corporation	CFLT	Confluent, Inc.
DT	Dynatrace, Inc.	RBRK	Rubrik, Inc.
BSY	Bentley Systems, Incorporated	GWRE	Guidewire Software, Inc.
IDCC	InterDigital, Inc.	TTMI	TTM Technologies, Inc.
CGNX	Cognex Corporation	SWKS	Skyworks Solutions, Inc.
DOCU	DocuSign, Inc.	EPAM	EPAM Systems, Inc.
AUR	Aurora Innovation, Inc.	LFUS	Littelfuse, Inc.
SAIL	SailPoint, Inc.	MANH	Manhattan Associates, Inc.
SMTC	Semtech Corporation	PCOR	Procore Technologies, Inc.
SANM	Sanmina Corporation	U	Unity Software Inc.
QRVO	Qorvo, Inc.	ARW	Arrow Electronics, Inc.
ALGM	Allegro MicroSystems, Inc.	PEGA	Pegasystems Inc.
FORM	FormFactor, Inc.	CRUS	Cirrus Logic, Inc.
CHYM	Chime Financial, Inc.	DOX	Amdocs Limited
VICR	Vicor Corporation	QBTS	D-Wave Quantum Inc.
ESE	ESCO Technologies Inc.	SLAB	Silicon Laboratories Inc.
CWAN	Clearwater Analytics Holdings, Inc.	PSN	Parsons Corporation
PAYC	Paycom Software, Inc.	VSAT	Viasat, Inc.
DBX	Dropbox, Inc.	DOCN	DigitalOcean Holdings, Inc.
NVDA	NVIDIA Corporation	AAPL	Apple Inc.
MSFT	Microsoft Corporation	AVGO	Broadcom Inc.
MU	Micron Technology, Inc.	ORCL	Oracle Corporation
AMD	Advanced Micro Devices, Inc.	PLTR	Palantir Technologies Inc.
CSCO	Cisco Systems, Inc.	LRCX	Lam Research Corporation
AMAT	Applied Materials, Inc.	IBM	International Business Machines Corporation
INTC	Intel Corporation	TXN	Texas Instruments Incorporated
KLAC	KLA Corporation	APH	Amphenol Corporation
ANET	Arista Networks, Inc.	CRM	Salesforce, Inc.
ADI	Analog Devices, Inc.	QCOM	QUALCOMM Incorporated

Ticker	Name	Ticker	Name
UBER	Uber Technologies, Inc.	PANW	Palo Alto Networks, Inc.
GLW	Corning Incorporated	NOW	ServiceNow, Inc.
ADBE	Adobe Inc.	INTU	Intuit Inc.
CRWD	CrowdStrike Holdings, Inc.	WDC	Western Digital Corporation
SNDK	Sandisk Corporation	ADP	Automatic Data Processing, Inc.
SNPS	Synopsys, Inc.	DELL	Dell Technologies Inc.
CDNS	Cadence Design Systems, Inc.	MSI	Motorola Solutions, Inc.
MRVL	Marvell Technology, Inc.	NET	Cloudflare, Inc.
FTNT	Fortinet, Inc.	SNOW	Snowflake Inc.
MPWR	Monolithic Power Systems, Inc.	TER	Teradyne, Inc.
CRWV	CoreWeave, Inc.	ADSK	Autodesk, Inc.
LITE	Lumentum Holdings Inc.	DDOG	Datadog, Inc.
MSTR	Strategy Inc.	CIEN	Ciena Corporation
MCHP	Microchip Technology Incorporated	UI	Ubiquiti Inc.
COHR	Coherent Corp.	KEYS	Keysight Technologies, Inc.
WDAY	Workday, Inc.	FISV	Fiserv, Inc.
ROP	Roper Technologies, Inc.	PAYX	Paychex, Inc.
FICO	Fair Isaac Corporation	ASTS	AST SpaceMobile, Inc.
CTSH	Cognizant Technology Solutions Corpora- tion	TDY	Teledyne Technologies Incorporated
XYZ	Block, Inc.	HPE	Hewlett Packard Enterprise Company
MDB	MongoDB, Inc.	ON	ON Semiconductor Corporation
JBL	Jabil Inc.	ZS	Zscaler, Inc.
GFS	GLOBALFOUNDRIES Inc.	ZM	Zoom Communications, Inc.
FIS	Fidelity National Information Services, Inc.	FSLR	First Solar, Inc.
PSTG	Everpure, Inc.	CPAY	Corpay, Inc.
FLEX	Flex Ltd.	Q	Qnity Electronics, Inc.
ALAB	Astera Labs, Inc.	LDOS	Leidos Holdings, Inc.
BR	Broadridge Financial Solutions, Inc.	NTAP	NetApp, Inc.
VRSN	VeriSign, Inc.	ENTG	Entegris, Inc.
MTSI	MACOM Technology Solutions Holdings, Inc.	SMCI	Super Micro Computer, Inc.
PTC	PTC Inc.	NXT	Nextpower Inc.
MKSI	MKS Inc.	FTV	Fortive Corporation
SSNC	SS&C Technologies Holdings, Inc.	HPQ	HP Inc.
TOST	Toast, Inc.	TWLO	Twilio Inc.
CDW	CDW Corporation	TRMB	Trimble Inc.
AKAM	Akamai Technologies, Inc.	FFIV	F5, Inc.
IOT	Samsara Inc.	OKTA	Okta, Inc.
GEN	Gen Digital Inc.	TYL	Tyler Technologies, Inc.
LSCC	Lattice Semiconductor Corporation	HUBS	HubSpot, Inc.
ZBRA	Zebra Technologies Corporation	SNX	TD SYNNEX Corporation
VIAV	Viavi Solutions Inc.	APPF	AppFolio, Inc.
ENPH	Enphase Energy, Inc.	BDC	Belden Inc.
VNT	Vontier Corporation	PATH	UiPath, Inc.
OS	OneStream, Inc.	PCTY	Paylocity Holding Corporation
OLED	Universal Display Corporation	TTAN	ServiceTitan, Inc.
IPGP	IPG Photonics Corporation	FROG	JFrog Ltd.
CORZ	Core Scientific, Inc.	ST	Sensata Technologies Holding plc
RGTI	Rigetti Computing, Inc.	AVT	Avnet, Inc.
LYFT	Lyft, Inc.	PLXS	Plexus Corp.
WEX	WEX Inc.	KVYO	Klaviyo, Inc.
NOVT	Novanta Inc.	DUOL	Duolingo, Inc.
INGM	Ingram Micro Holding Corporation	FOUR	Shift4 Payments, Inc.
RAL	Ralliant Corporation	GTLB	GitLab Inc.
EXLS	ExlService Holdings, Inc.	BMI	Badger Meter, Inc.

Ticker	Name	Ticker	Name
ONDS	Ondas Inc.	S	SentinelOne, Inc.
RUN	Sunrun Inc.	BILL	BILL Holdings, Inc.
NTSK	Netskope, Inc.	OSIS	OSI Systems, Inc.
YOU	Clear Secure, Inc.	ITRI	Itron, Inc.
VISN	Vistance Networks, Inc.	ACMR	ACM Research, Inc.
ACIW	ACI Worldwide, Inc.	LIF	Life360, Inc.
PI	Impinj, Inc.	SAIC	Science Applications International Corporation
CVLT	Commvault Systems, Inc.	ZETA	Zeta Global Holdings Corp.
QLYS	Qualys, Inc.	CALX	Calix, Inc.
SYNA	Synaptics Incorporated	QTWO	Q2 Holdings, Inc.
NIQ	NIQ Global Intelligence plc	BOX	Box, Inc.
NATL	NCR Atleos Corporation	DIOD	Diodes Incorporated
SOUN	SoundHound AI, Inc.	AAOI	Applied Optoelectronics, Inc.
KD	Kyndryl Holdings, Inc.	TDC	Teradata Corporation
BULL	Webull Corporation	EEFT	Euronet Worldwide, Inc.
PAY	Paymentus Holdings, Inc.	DBD	Diebold Nixdorf, Incorporated
VRNS	Varonis Systems, Inc.	ACLS	Axcelis Technologies, Inc.
VRRM	Verra Mobility Corporation	AMBA	Ambarella, Inc.
LASR	nLIGHT, Inc.	BELFA	Bel Fuse Inc.
RELY	Remitly Global, Inc.	FSLY	Fastly, Inc.
TENB	Tenable Holdings, Inc.	NSIT	Insight Enterprises, Inc.
POWI	Power Integrations, Inc.	VSH	Vishay Intertechnology, Inc.
NAVN	Navan, Inc.	RNG	RingCentral, Inc.
UCTT	Ultra Clean Holdings, Inc.	BELFB	Bel Fuse Inc.
DAVE	Dave Inc.	KN	Knowles Corporation
SPSC	SPS Commerce, Inc.	PLAB	Photronics, Inc.
DXC	DXC Technology Company	ALRM	Alarm.com Holdings, Inc.
AGYS	Agilysys, Inc.	AVPT	AvePoint, Inc.
BL	BlackLine, Inc.	BLKB	Blackbaud, Inc.
NN	NextNav Inc.	PLUS	ePlus inc.
CSGS	CSG Systems International, Inc.	NTCT	NetScout Systems, Inc.
BHE	Benchmark Electronics, Inc.	ADEA	Adeia Inc.
FRSH	Freshworks Inc.	VECO	Veeco Instruments Inc.
VERX	Vertex, Inc.	CNXC	Concentrix Corporation
EXTR	Extreme Networks, Inc.	INTA	Intapp, Inc.
GTM	ZoomInfo Technologies Inc.	NVTS	Navitas Semiconductor Corporation
BRZE	Braze, Inc.	PAYO	Payoneer Global Inc.
GRND	Grindr Inc.	SONO	Sonos, Inc.
PTRN	Pattern Group Inc.	ROG	Rogers Corporation
NCNO	nCino, Inc.	QUBT	Quantum Computing Inc.
DGII	Digi International Inc.	SEMR	Semrush Holdings, Inc.
SHLS	Shoals Technologies Group, Inc.	ASGN	ASGN Incorporated
MQ	Marqeta, Inc.	EVCM	EverCommerce Inc.
ASAN	Asana, Inc.	BBAI	BigBear.ai Holdings, Inc.
ARRY	Array Technologies, Inc.	MXL	MaxLinear, Inc.
APPN	Appian Corporation	EVTC	EVERTEC, Inc.
ALKT	Alkami Technology, Inc.	CTS	CTS Corporation
ICHR	Ichor Holdings, Ltd.	CNXN	PC Connection, Inc.
RAMP	LiveRamp Holdings, Inc.	KDK	Kodiak AI, Inc.
PRGS	Progress Software Corporation	AI	C3.ai, Inc.
ATEN	A10 Networks, Inc.	SKYT	SkyWater Technology, Inc.
COHU	Cohu, Inc.	CXM	Sprinklr, Inc.
VIA	Via Transportation, Inc.	INOD	Innodata Inc.
DAKT	Daktronics, Inc.	PDFS	PDF Solutions, Inc.
FLYW	Flywire Corporation	VYX	NCR Voyix Corporation
GCT	GigaCloud Technology Inc.	FIVN	Five9, Inc.

Ticker	Name	Ticker	Name
AXTI	AXT, Inc.	HLIT	Harmonic Inc.
DFIN	Donnelley Financial Solutions, Inc.	ALNT	Allient Inc.
OUST	Ouster, Inc.	PENG	Penguin Solutions, Inc.
PGY	Pagaya Technologies Ltd.	CRCT	Cricut, Inc.
NABL	N-able, Inc.	TASK	TaskUs, Inc.
PRCH	Porch Group, Inc.	SSYS	Stratasys Ltd.
AEHR	Aehr Test Systems, Inc.	WOLF	Wolfspeed, Inc.
PAR	PAR Technology Corporation	INDI	indie Semiconductor, Inc.
AMPL	Amplitude, Inc.	ADTN	ADTRAN Holdings, Inc.
SCSC	ScanSource, Inc.	MNTN	MNTN, Inc.
AEVA	Aeva Technologies, Inc.	CTLP	Cantaloupe, Inc.
WYFI	WhiteFiber, Inc.	RSKD	Riskified Ltd.
AOSL	Alpha and Omega Semiconductor Limited	ALIT	Alight, Inc.
YEXT	Yext, Inc.	LYTS	LSI Industries Inc.
IIIV	i3 Verticals, Inc.	LPTH	LightPath Technologies, Inc.
AIP	Arteris, Inc.	DJCO	Daily Journal Corporation
AMBQ	Ambiq Micro, Inc.	PD	PagerDuty, Inc.
MITK	Mitek Systems, Inc.	VPG	Vishay Precision Group, Inc.
NTGR	NETGEAR, Inc.	CRSR	Corsair Gaming, Inc.
RDVT	Red Violet, Inc.	DSP	Viant Technology Inc.
CEVA	CEVA, Inc.	CCSI	Consensus Cloud Solutions, Inc.
GDYN	Grid Dynamics Holdings, Inc.	IBTA	Ibotta, Inc.
CLMB	Climb Global Solutions, Inc.	AIOT	PowerFleet, Inc.
MLAB	Mesa Laboratories, Inc.	APPS	Digital Turbine, Inc.
BLND	Blend Labs, Inc.	FEIM	Frequency Electronics, Inc.
RPD	Rapid7, Inc.	IMXI	International Money Express, Inc.
UMAC	Unusual Machines, Inc.	PRTH	Priority Technology Holdings, Inc.
CLFD	Clearfield, Inc.	DVLT	Datavault AI Inc.
GLOO	Gloo Holdings, Inc.	API	Agora, Inc.
NNDM	Nano Dimension Ltd.	OSPN	OneSpan Inc.
SPT	Sprout Social, Inc.	IBEX	IBEX Limited
KOPN	Kopin Corporation	BAND	Bandwidth Inc.
CRNC	Cerence Inc.	RBBN	Ribbon Communications Inc.
SABR	Sabre Corporation	TTGT	TechTarget, Inc.
EGHT	8x8, Inc.	HCKT	The Hackett Group, Inc.
BKKT	Bakkt, Inc.	ONTF	ON24, Inc.
AVNW	Aviat Networks, Inc.	MEI	Methode Electronics, Inc.
NVEC	NVE Corporation	OOMA	Ooma, Inc.
SMRT	SmartRent, Inc.	BKTI	BK Technologies Corporation
TLS	Telos Corporation	DDD	3D Systems Corporation
EXOD	Exodus Movement, Inc.		

Author: compiled by the author.

Appendix F: January 2026 Technology Universe Equal-Weight And MAD-CVaR-Entropy Portfolio Constituent Assets

Table 20

EW portfolio constituent technology universe shares on January 2026 with per-share returns

Ticker	Company	Industry	+/- %
ONDS	Ondas	Drones, defence wireless systems	462.20
LITE	Lumentum	Optical networking, datacenter optics	322.00
VISN	Vistance Networks	Connectivity, network infrastructure	122.70

Ticker	Company	Industry	+/- %
ALAB	Astera Labs	AI connectivity, datacenter semiconductors	102.50
LASR	nLIGHT	Industrial & defence lasers	100.50
SITM	SiTime	Timing semiconductors	79.50
UMAC	Unusual Machines	Drones, aerospace	70.30
AAOI	Applied Optoelectronics	Optical networking, datacenter optics	56.20
UI	Ubiquiti	Networking hardware	38.90
PLTR	Palantir	AI, government & enterprise software	28.50
CRNC	Cerence	Automotive voice AI, mobility software	21.70
TWLO	Twilio	Cloud communications software	18.00
NET	Cloudflare	Cloud networking, cybersecurity	6.00
RDVT	Red Violet	Identity, data analytics	4.80
SOUN	SoundHound AI	Conversational AI	1.10
LIF	Life360	Consumer safety & location platform	0.00
CORZ	Core Scientific	Bitcoin mining, digital infrastructure	-7.30
RBRK	Rubrik	Cyber resilience, backup security	-11.40
DUOL	Duolingo	Consumer language learning app	-56.20
MSTR	Strategy (MicroStrategy)	Bitcoin proxy, treasury software	-57.90

Author: compiled by the author.

Table 21

MCE portfolio constituent technology universe shares on January 2026 with per-share returns

Ticker	Company	Industry	+/- %
GLW	Corning Incorporated	Specialty glass, optical fibre, connectivity components	73.70
APH	Amphenol Corporation	Electronic connectors, sensors	43.90
KLAC	KLA Corporation	Semiconductor equipment, process control	42.30
AAPL	Apple Inc.	Consumer electronics, platform	30.70
TDC	Teradata Corporation	Data analytics, data warehousing	28.30
KEYS	Keysight Technologies, Inc.	Electronic test & measurement	25.60
IBEX	IBEX Limited	Customer-experience outsourcing, digital services	25.20
PLUS	ePlus inc.	IT solutions, value-added reseller	22.00
HPE	Hewlett Packard Enterprise Company	Enterprise servers, networking, hybrid cloud	19.40
CSGS	CSG Systems International, Inc.	Billing, customer-engagement software	17.40
ADI	Analog Devices, Inc.	Analog semiconductors	14.60
LDOS	Leidos Holdings, Inc.	Government IT, defence, engineering	14.50
SNX	TD SYNNEX Corporation	IT distribution, technology services	13.60
OSIS	OSI Systems, Inc.	Security, inspection systems, healthcare devices	12.70
DBD	Diebold Nixdorf, Incorporated	Banking & retail technology, ATMs	12.40
CSCO	Cisco Systems, Inc.	Networking hardware, enterprise infrastructure	12.00
JBL	Jabil Inc.	Electronics manufacturing services	11.30
CACI	CACI International Inc.	Government IT, defence tech	11.20
RDVT	Red Violet, Inc.	Identity, data analytics	4.80
IBM	International Business Machines Corporation	Enterprise IT, software, consulting	1.30

Author: compiled by the author.

Appendix G: GitHub Repository

<https://github.com/mariasutt/MCE>

Resümee

KAHEETAPILISE KESKMISEL TOOTLUSE JA KESKMISEL ABSOLUUTHÄLBEI PÕHINEVA EELVALIKU NING TINGIMUSLIKU RISKIVÄÄRTUSE PIIRANGU ALL ENTROOPIAT MAKSIMEERVA INVESTEERIMISPORTFELLI KOOSTAMISE MUDELI VÄLJA TÖÖTAMINE

Maria-Britta Sütt

Käesolevas magistritöös töötati välja kaheetapiline investeerimisportfelli koostamise mudel, "MCE". Mudel kasutab aktive keskmist tootlust, aktive keskmist absoluuthälvet, tingimuslikku riskiväärtust (CVaR) ning Shannoni entroopiat. Autori hinnangul on praktikutele portfelli koostamisel tähtsad neli elementi: 1) võrdlemisi stabiilsed valimivälised tulemused, 2) portfellis olevate aktive arv, 3) kontroll portfelli väärtuse kaotuse riski üle ja 4) mudeli reaalse rakendamise lihtsus. Tööd motiveeris kirjanduses tuvastatud uurimislünk nende nelja elemendi vahel. Aktivate keskmisel tootluse ja dispersioonil põhinevad mudelid (MV) on praeguseks leidnud erinevaid meetodilisi lahendusi valimiväliste tulemuste stabiilsuse probleemile, eriti kombineerituna robustse optimeerimisega. MV on kaotuseriski haldamisel pigem ülemäära konservatiivne, mis ei pruugi kõikidele investoritele sobida. Samuti ei ole MV mudelid, eriti tänapäevased variandid, mis rakendavad faktoreid, masinõpet ja muid keerukaid meetodeid varasemate andmete kaasamiseks, praktikutele lihtsad kasutada. Lisaks ei võimalda MV üldiselt määrata portfelli jaoks täpset soovitud aktive arvu. Võrdsete kaaludega portfellid (EW) on samas triviaalsed rakendada ning võimaldavad hoida täpselt nii palju aktiivaid kui investor tahab. Samas ei arvesta EW riskidega ega sisalda valimivälisest stabiilsusest kontrollivaid mehhanisme. Töö jaoks välja töötatud mudel, MCE, sisaldab kõike nelja praktilist elementi. Mudeli eesmärk saavutatakse jagades portfelli koostamine kaheks etapiks. Kirjandusest selgub, et praktikas on kaheetapiline portfellide koostamine sageli rakendatav lähenemine; kasutatakse erinevaid mõõdikuid ja valikukriteeriume. MCE on suures osas automaatne ning nõuab sisendina vähemalt aktive andmeid ning portfelli kaasatavate aktive arvu. MCE esimene etapp teeb etteantud aktive hulgast keskmise tootluse ja keskmise absoluuthälbe alusel eelvaliku. Teine etapp võtab tehtud valiku ette ning maksimeerib aktivele portfellis kaalude määramiseks tingimusliku riskiväärtuse piirangu all entroopia, et saavutada maksimaalne võimalik portfelli

hajutatus. MCE mudelit testiti kahes mitmeperioodilises simulatsioonis: väiksemal valimil laiapõhistest USA ETF-idest 39 pooleaastase perioodi jooksul, ning suuremal valimil USA tehnoloogiasektori aktsiatest 21 pooleaastase perioodi jooksul. Testides võrreldi MCE tulemusi kolme võrdleva mudeliga: laia ja kitsa MV mudeliga ning võrdsete kaaludega portfelli mudeliga. MCE saavutas ETF-idel testides teistega võrreldes parimaid tulemusi. Mudeli kumulatiivne tootlus oli kõrgeim, 384,5%, võrreldes EW 337,4%-ga. MCE kaotuse risk, mõõdetuna keskmise CVaR väärtusena, oli 1,9% võrrelduna EW 2,4%-ga. MCE keskmine aktive vahetuvus portfellis oli 22%, võrrelduna EW 35%-ga. Tehnoloogiaaktsiatel tehtud testides ei domineerinud tootluses, kaotuseriskis, ega aktive vahetuvuses ükski mudel. EW saavutas kõige kõrgema kumulatiivse tootluse, 697%, aga samas oli EW portfelli keskmine CVaR 4,6%. MCE oli võrreldavalt tasakaalustatum. Selle kumulatiivne tootlus oli 315,5% ning CVaR 3%. Tehnoloogiaaktsiate puhul saavutas kitsam MV variant, mis kasutab optimeerimissisendina samu aktiivide nagu MCE, MCE-st poole väiksema kumulatiivse tootluse, 153% vs 315,5%. Kaotuse risk oli mõlemal sisuliselt sama, 3%, aga aktive vahetuvus oli MCE-l jälle tunduvalt väiksem, 67% vs 83%. Kokkuvõttes võib MCE mudeli tulemusi antud valimi piires, testides kasutatud sisenditega ning võrrelduna teiste testitud mudelitega, hinnata eesmärgipäraselt tasakaalustatuteks. Samuti sai autori hinnangul täidetud eesmärk, et mudel oleks praktikutele võrdlemisi lihtne kasutada.

Non-exclusive License for Reproduction and Public Access

I, Maria-Britta Sütt, hereby grant the University of Tartu permission (non-exclusive license) to freely reproduce my thesis:

DEVELOPING A TWO-STAGE PORTFOLIO CONSTRUCTION MODEL
WITH MEAN ABSOLUTE DEVIATION SCREENING AND CONDITIONAL VALUE-
AT-RISK-CONSTRAINED ENTROPY WEIGHTING

Supervised by Mark Kantšukov (PhD), for the purposes of preservation, including digital archival under Creative Commons license CC BY NC ND 4.0, until copyright expiry. The license allows reproduction, distribution and public communication of the work when credited to the author, and forbids any commercial use until copyright expiry.

I am aware that I retain the aforementioned rights as the author and confirm that granting a non-exclusive licence does not infringe the intellectual property rights of any other parties or any rights arising from relevant data protection acts.

Maria-Britta Sütt

14.05.2026