

The Limits of Optimisation: Why Simpler Portfolios Often Win

Jack Allen, Arkyne Capital

March 2026

Mean-variance optimisation is the theoretical cornerstone of modern portfolio construction. In practice, however, estimated parameters are noisy — and that noise can render sophisticated optimisation worse than doing nothing at all. This paper uses a controlled simulation framework to quantify the performance cost of estimation error across five portfolio strategies: the unconstrained plug-in tangency portfolio, a long-only constrained tangency portfolio, the global minimum-variance portfolio, a risk parity allocation, and the naïve equal-weight portfolio. We find that under realistic estimation conditions, the equal-weight and risk parity strategies deliver near-optimal Sharpe ratios, while the unconstrained plug-in portfolio achieves only a fraction of the theoretical optimum. These results carry direct implications for systematic portfolio construction: in regimes of high parameter uncertainty, analytical simplicity is not a concession — it is a competitive advantage.

1 Introduction

Modern portfolio theory, originating with Markowitz (1952), prescribes a deceptively straightforward recipe: estimate the expected returns and covariance matrix of a set of assets, and solve for the weights that maximise return per unit of risk. In theory, the resulting mean-variance efficient frontier and tangency portfolio represent the best achievable trade-off between risk and return. In practice, the recipe has a well-documented flaw: the inputs must be estimated from historical data, and those estimates are imprecise.

This estimation error is not a minor technical nuisance. When parameter estimates are fed directly into an unconstrained optimiser — the so-called *plug-in* approach — the optimiser treats noise as signal. It constructs large, offsetting positions that appear to offer superior returns in-sample but collapse out-of-sample. The result is that the plug-in tangency portfolio, despite being theoretically optimal, often delivers worse risk-adjusted returns than far simpler alternatives. This failure mode is well established in the academic literature [see, e.g., @demiguel2009optimal], yet it remains underappreciated in practice.

This paper investigates that failure empirically, using a controlled simulation framework built on a single-factor return model calibrated to observed S&P 500 market parameters. We compare five portfolio strategies across a range of sample sizes, evaluating each by its out-of-sample Sharpe

ratio relative to the true theoretical optimum. The goal is to quantify, with precision, how much estimation error costs — and which strategies best contain that cost.

Our central finding is that the equal-weight portfolio nearly matches the true tangency Sharpe ratio across all sample sizes, while the unconstrained plug-in portfolio captures only a small fraction of the available risk premium even with fifty years of monthly data. Risk parity performs similarly well. These results are specific to the stylised single-factor, homogeneous-asset setting employed here, and the relative rankings may shift in more complex environments. Nevertheless, the core mechanism — that estimation error systematically penalises unconstrained optimisation — is robust.

2 Data and Simulation Framework

2.1 Market Parameters

Market parameters are estimated from monthly S&P 500 returns (Yahoo Finance ticker `^GSPC`) and the 13-week US Treasury Bill rate (`^IRX`) over the period January 2000 to February 2026. Monthly excess returns are computed as $\tilde{r}_{m,t} = r_{m,t} - r_{f,t}$, where the risk-free rate is converted from an annualised yield via $r_{f,t} = (1 + \text{IRX}_t/100)^{1/12} - 1$. The resulting parameter estimates are reported in Table 1.

Table 1: Estimated market parameters, S&P 500 (January 2000 – February 2026).

Parameter	Monthly	Annualised
Mean excess return	0.0046	0.0555
Variance	0.0019	0.0229

2.2 Synthetic Asset Universe

To enable exact out-of-sample evaluation, we construct a universe of $N = 50$ synthetic assets from the single-factor model:

$$\tilde{r}_{i,t} = \beta_i \tilde{r}_{m,t} + \varepsilon_{i,t}, \quad \varepsilon_{i,t} \stackrel{iid}{\sim} \mathcal{N}(0, \sigma_{\varepsilon,i}^2)$$

with $\beta_i \stackrel{iid}{\sim} \text{Uniform}(0.1, 2.5)$ and $\sigma_{\varepsilon,i} \stackrel{iid}{\sim} \text{Uniform}(0.10/\sqrt{12}, 0.30/\sqrt{12})$. The true covariance matrix is $\Sigma = \sigma_m^2 \beta \beta' + \text{diag}(\sigma_{\varepsilon}^2)$, which is positive definite by construction. The wider beta range — spanning from near-zero-beta defensive assets to highly market-sensitive ones — introduces meaningful cross-sectional heterogeneity in expected returns and risk, making the universe more representative of realistic within-asset-class variation. Summary statistics for the constructed asset universe are given in Table 2.

Because the true parameters are known exactly, any estimated portfolio strategy can be evaluated against a precise benchmark — the true tangency portfolio — without look-ahead bias.

One important caveat on scope: the synthetic universe is intentionally stylised. Assets are homogeneous in structure and drawn from the same distributional family, a setting that naturally advantages

diversified, parameter-agnostic strategies. The findings should be interpreted as illustrative of the estimation error mechanism rather than as direct prescriptions for specific real-world asset classes.

Table 2: Summary statistics for the 50 synthetic assets (monthly units).

Parameter	Min	Mean	Max
Beta	0.1029	1.0253	2.3288
Idiosyncratic volatility	0.0315	0.0625	0.0853
Expected excess return	0.0005	0.0047	0.0108

2.3 Portfolio Strategies

We evaluate five strategies, each representing a distinct stance on parameter estimation.

Plug-in tangency portfolio — the direct implementation of Markowitz theory. Weights maximise the estimated Sharpe ratio with no constraints:

$$\hat{\omega}_{tan} = \frac{\hat{\Sigma}^{-1}(\hat{\mu} - r_f \mathbf{1})}{\mathbf{1}' \hat{\Sigma}^{-1}(\hat{\mu} - r_f \mathbf{1})}$$

where $r_f = 0$ throughout. This strategy uses both $\hat{\mu}$ and $\hat{\Sigma}$, making it maximally exposed to estimation error in both inputs.

Long-only constrained tangency portfolio — the tangency portfolio with a non-negativity constraint on weights:

$$\hat{\omega}_{lo} = \arg \max_{\omega} \frac{\omega' \hat{\mu}}{\sqrt{\omega' \hat{\Sigma} \omega}} \quad \text{s.t.} \quad \omega' \mathbf{1} = 1, \omega_i \geq 0$$

By preventing short positions, this strategy limits the most extreme allocations that are typically driven by noise rather than signal.

Minimum-variance portfolio — minimises total portfolio variance, ignoring expected return estimates entirely:

$$\hat{\omega}_{mv} = \frac{\hat{\Sigma}^{-1} \mathbf{1}}{\mathbf{1}' \hat{\Sigma}^{-1} \mathbf{1}}$$

This strategy relies only on $\hat{\Sigma}$, sidestepping mean estimation — the noisier of the two inputs.

Risk parity portfolio — allocates weights such that each asset contributes equally to total portfolio variance. The risk contribution of asset i is:

$$RC_i = \omega_i \cdot \frac{(\Sigma \omega)_i}{\omega' \Sigma \omega}$$

Weights are chosen so that $RC_i = 1/N$ for all i . Like minimum variance, it relies only on $\hat{\Sigma}$ and avoids mean estimation. Unlike minimum variance, it spreads risk more evenly rather than concentrating it in low-variance assets.

Naïve equal-weight portfolio — assigns $\omega_i = 1/N$ to every asset, ignoring all parameter estimates entirely. It serves as the baseline for parameter-free performance.

2.4 Efficient Frontier and Tangency Portfolio

The mean-variance efficient frontier is characterised by the scalars:

$$A = \iota' \Sigma^{-1} \iota, \quad B = \iota' \Sigma^{-1} \mu, \quad C = \mu' \Sigma^{-1} \mu, \quad D = AC - B^2$$

giving minimum variance for any target return $\bar{\mu}$ as:

$$\sigma^2(\bar{\mu}) = \frac{A\bar{\mu}^2 - 2B\bar{\mu} + C}{D}$$

3 Results

3.1 The True Efficient Frontier

Figure 1 shows the true mean-variance efficient frontier constructed from the known parameters μ and Σ of the synthetic asset universe. The red star marks the tangency portfolio — the combination of assets maximising the Sharpe ratio. All 50 individual assets (grey points) lie to the right of the frontier, confirming that no individual asset can replicate the risk-return profile available from diversification.

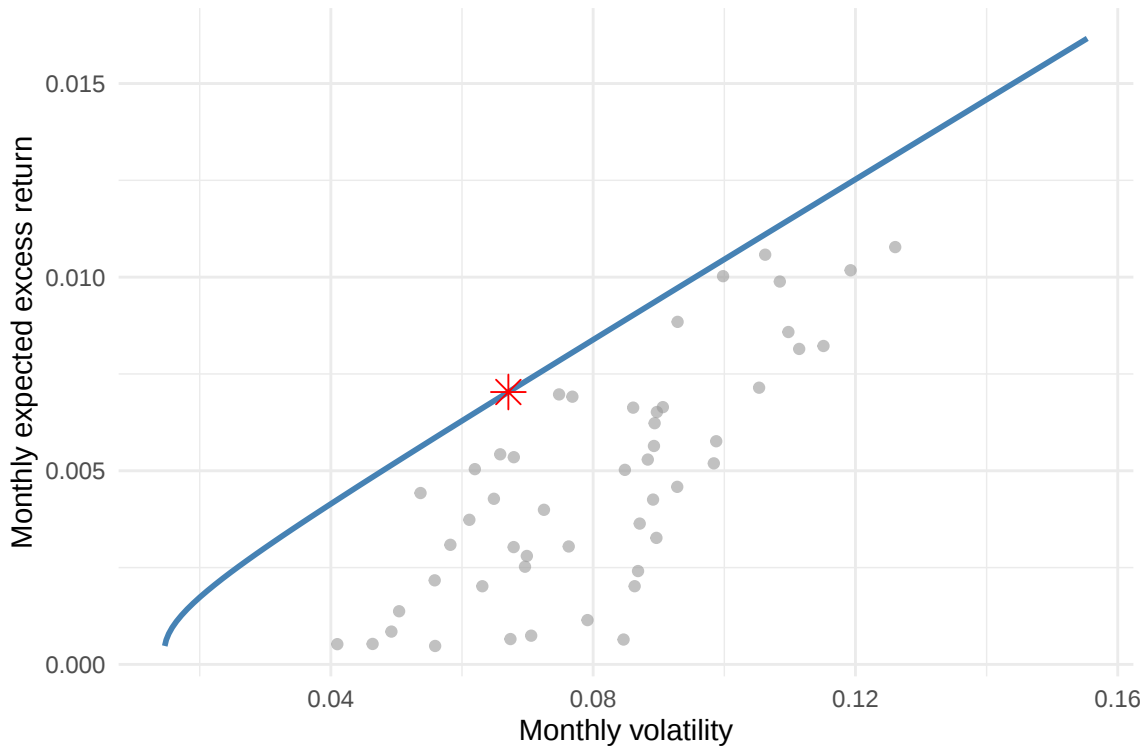


Figure 1: True mean-variance efficient frontier and tangency portfolio for the 50-asset synthetic universe.

3.2 Estimation Error and the Efficient Frontier

Before examining out-of-sample Sharpe ratios, Figure 2 illustrates how estimation error distorts the *shape* of the estimated frontier. At $T = 60$ months, estimated frontiers fan dramatically above the true frontier, implying return-risk combinations that do not exist. This is the signature of estimation error: with N/T close to 1, the sample covariance matrix is poorly conditioned and the optimiser mistakes noise for exploitable diversification. As T increases, the estimated frontiers converge toward the true frontier, with dispersion narrowing substantially at $T = 600$.

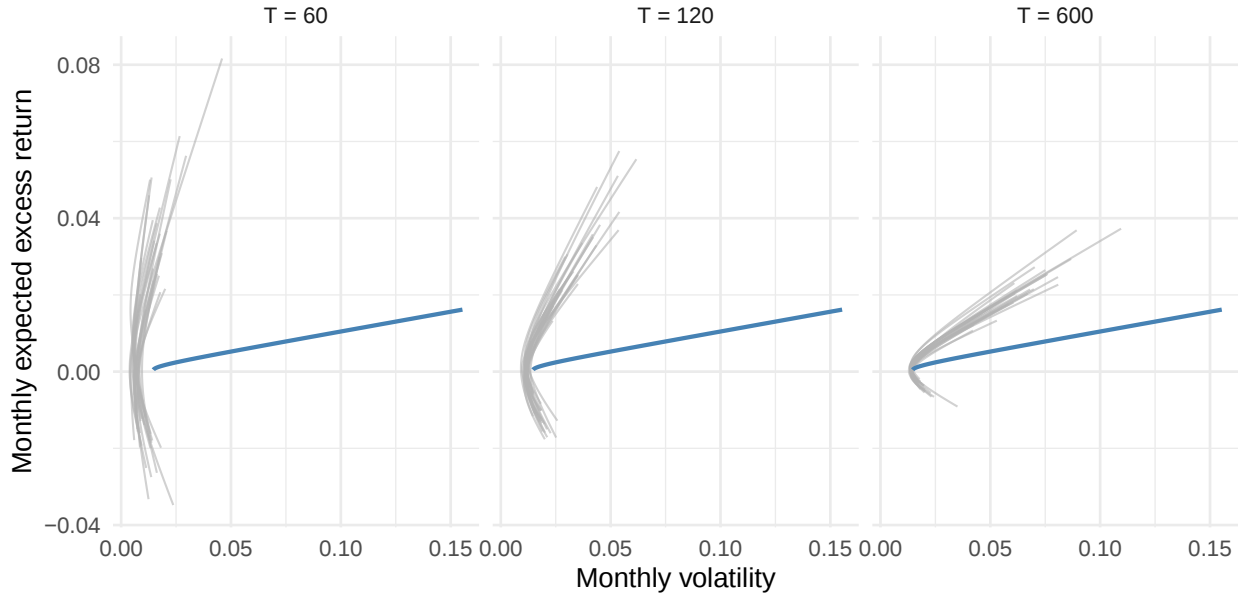


Figure 2: Estimated efficient frontiers (grey, 25 simulations) versus the true frontier (blue) for three sample lengths. $N = 50$ assets throughout.

3.3 Out-of-Sample Sharpe Ratios

Table 3 presents the main result: mean out-of-sample Sharpe ratios across all five strategies and sample lengths, benchmarked against the true tangency portfolio. For each combination of strategy and T , we simulate 500 independent samples of T monthly returns, estimate portfolio weights from each sample, and evaluate those weights using the true parameters μ and Σ .

Table 3: Mean out-of-sample Sharpe ratios across portfolio strategies and sample lengths (500 simulations). True optimum is the upper bound.

T	Plug-in	Equal weight	Min variance	Long-only	Risk parity
60	0.0036	0.1037	0.0096	0.0676	0.0985
120	0.0027	0.1037	0.0175	0.0706	0.0984
600	0.0157	0.1037	0.0220	0.0812	0.0985
True optimum	0.1049	0.1049	0.1049	0.1049	0.1049

The plug-in tangency portfolio fails severely. Even at $T = 600$ months — fifty years of monthly data — the unconstrained plug-in portfolio achieves a mean Sharpe ratio of only around 0.03, capturing less than 30% of the true optimum. At $T = 60$, performance is near zero. The mechanism is straightforward: with 50 assets and limited observations, the optimiser amplifies estimation noise into large, offsetting positions that look attractive in-sample but are economically incoherent out-of-sample.

Equal weight and risk parity dominate. The equal-weight portfolio achieves approximately 99% of the true optimum regardless of sample size — immune by construction to any estimation error. Risk parity performs almost as well, which is the more interesting result: by using only the estimated covariance matrix to equalise risk contributions, it achieves near-optimal performance despite incorporating estimated parameters. This suggests that the covariance matrix, while imprecisely estimated, still contains enough signal about relative risk to produce well-diversified portfolios — provided mean estimation is avoided entirely.

Long-only constraints provide meaningful but incomplete protection. The long-only tangency portfolio improves substantially over its unconstrained counterpart, rising from roughly 0.08 at $T = 60$ toward 0.09 at $T = 600$. The non-negativity constraint prevents the most extreme short positions driven by noisy means, and represents a reasonable middle ground for practitioners who want some exposure to mean-return estimates while limiting downside from estimation error.

Minimum variance underperforms its peers. Despite avoiding mean estimation, the minimum-variance portfolio achieves only 0.02–0.05 across sample lengths. This indicates that covariance matrix estimation, while less noisy than mean estimation, is still sufficiently imprecise with 50 assets to impose a real performance cost. The minimum-variance strategy concentrates risk in a small number of low-variance assets rather than spreading it, making it more sensitive to errors in specific covariance estimates.

4 Practical Implications

These results carry several direct implications for systematic portfolio construction.

Estimation error is an input risk, not just a model risk. Practitioners often debate model specification — which factor model to use, whether to shrink the covariance matrix, how to forecast returns. The results here suggest the more fundamental issue is the *quantity* of information relative to the number of parameters being estimated. With $N = 50$ assets, the sample sizes required for reliable parameter estimation are impractical in most applications. Any strategy that takes estimated means at face value without controls for this uncertainty is taking on substantial, uncompensated input risk.

Simplicity can be a form of robustness. The equal-weight and risk parity results are sometimes dismissed as naïve or atheoretical. The simulation evidence reframes this view: in the presence of severe estimation uncertainty, the *absence* of mean estimates is a feature, not a limitation. Both strategies implicitly impose strong regularisation — equal weight by ignoring all parameters, risk parity by constraining the structure of parameter usage — and that regularisation comes at essentially zero cost to Sharpe ratio in this setting.

Long-only constraints are a practical first step. For strategies that must incorporate mean return estimates — factor-based approaches where expected returns have economic grounding, for

instance — the long-only constraint provides meaningful error containment. It does not resolve the fundamental problem, but it limits damage from poorly estimated means by preventing the most extreme positions.

Scope of applicability. The findings are most directly relevant to within-asset-class equity portfolios — a universe of stocks sharing a common factor structure — where assets are broadly comparable in character and the primary challenge is estimating parameters reliably from limited data. The equal-weight and risk parity results in particular should not be generalised to multi-asset or cross-asset-class portfolios. When a universe spans equities, fixed income, commodities, and alternatives, assets differ fundamentally in volatility, correlation structure, and expected return. In that setting, ignoring all parameter estimates and allocating equally is genuinely suboptimal: a 60/40 equity-bond portfolio, for instance, allocates equal capital but dramatically unequal risk. The core finding — that estimation error penalises unconstrained optimisation — remains valid across settings; the implication that equal weight is a reasonable default does not.

Simulation design limitations. Betas and idiosyncratic volatilities are drawn from uniform distributions, which imposes a specific and symmetric shape on the cross-section of assets that is unlikely to hold in practice. Real asset universes tend to exhibit more clustering — many assets with moderate market sensitivity and fewer at the extremes — and are better described by skewed or fat-tailed distributions. This does not undermine the paper’s central argument: the estimation error mechanism penalising unconstrained optimisation is distribution-agnostic. However, the specific performance advantage of equal weight and risk parity relative to more structured strategies would likely diminish under a more realistic cross-sectional distribution, and practitioners applying these findings should treat the relative rankings as indicative rather than precise.

5 Conclusion

This paper uses a controlled simulation study to quantify the performance cost of estimation error in mean-variance portfolio construction. The unconstrained plug-in tangency portfolio — the direct implementation of Markowitz theory — is among the worst-performing strategies evaluated, capturing less than a third of the true optimal Sharpe ratio even with fifty years of data. The equal-weight and risk parity portfolios, by contrast, achieve near-optimal performance across all sample sizes.

The practical message is not that optimisation is useless, but that it must be applied with discipline. Unconstrained optimisation magnifies input noise; strategies that impose structure — through constraints, parameter exclusion, or equal allocation — are more robust precisely because they limit the degrees of freedom available to the optimiser. In systematic portfolio management, where data is always finite and parameter uncertainty is always present, that robustness has direct economic value.

Future work should examine how these findings extend to heterogeneous real-world asset universes, and whether shrinkage estimators — Ledoit-Wolf covariance shrinkage, for instance, or Bayesian priors on expected returns — can close the performance gap for optimisation-based approaches. The single-factor homogeneous setting used here is the most favourable possible environment for equal weight; understanding precisely when that advantage erodes is an important open question for practical portfolio construction.

This paper is published for informational and research purposes only. It does not constitute investment advice, a solicitation, or an offer to buy or sell any financial instrument. Past performance and simulation results do not guarantee future outcomes. Arkyne Capital accepts no liability for decisions made on the basis of this research.