

The Maximum Geometric Mean Portfolio

Moshe Levy*

The Maximum Geometric Mean (MGM) portfolio has well-known advantages for long-run investors. When borrowing is realistically constrained, the MGM portfolio dominates the maximal Sharpe portfolio for all but the most extremely risk-averse investors, even for short horizons. While there is generally no analytical solution for the exact MGM, we derive a closed-form solution for the portfolio weights of the approximate MGM portfolio, based on the mean-variance approximation to the GM. This approach yields an excellent approximation to the MGM, does not require knowledge of the distribution of returns, and allows for the application of standard statistical shrinkage methods.

Keywords: geometric mean, Sharpe ratio, limited borrowing, mean-variance analysis, statistical shrinkage.

*The Hebrew University, Jerusalem 91905, Israel. mslm@huji.ac.il

1. Introduction

The portfolio with Maximal Geometric Mean (MGM) has well-known advantages for long-run investors: in an i.i.d. setting, this portfolio yields a higher terminal wealth than any other portfolio with a probability approaches 1 as the investment horizon increases. This property convinced many scholars that long-run investors should aim to maximize their portfolios' GM (Kelly 1956, Latane 1959, Breiman 1960, Thorp 1975, Bernstein 1976, Markowitz 1976, 2006, MacLean et al. 1992, Bali et al. 2009, Levy 2016, 2024a,b, Ziemba 2015, and Lo, Orr, and Zhang 2018). Samuelson (1971) and Merton and Samuelson (1974) object to this long-run objective, because investors with Constant Relative Risk Aversion (CRRA) utility and a RRA coefficient different than 1 are generally better-off with portfolios different than the MGM portfolio, and this holds for any investment horizon. M. Levy (2024) offers a possible reconciliation of these two opposing views: he shows that any RRA value substantially different than 1 leads to choices that no “reasonable” person would ever make, and thus concludes that GM maximization is a valid objective for any reasonable long-run investor. Bali et al. (2009) and Levy (2016) reach similar conclusions, based on the application of the almost stochastic dominance criterion (Leshno and Levy 2002).

While the virtues of the MGM portfolio for long-run investors have been extensively discussed, this portfolio has important additional advantages for short-run investors when borrowing is realistically limited. The top panel of Figure 1 illustrates this point. If borrowing at the risk-free rate is unrestricted, the portfolio with the maximal Sharpe ratio dominates all other portfolios by the mean-variance criterion: for any other portfolio with a lower Sharpe ratio, such as portfolio P , there exists a combination of the maximal Sharpe portfolio and the risk-free asset, portfolio P' ,

that has exactly the same standard deviation as P , but a higher expected return.¹ In contrast, if borrowing is limited, portfolio P' may be unattainable, and some investors may prefer portfolio P over the maximal Sharpe portfolio. In this setting, there is no single portfolio that is optimal for all investors. However, the MGM has two properties, discussed in the next section, that make it an attractive choice for many investors: i) it is on (or very close to) the MV efficient frontier; and ii) it typically has a much higher expected return (and standard deviation) than the maximal Sharpe portfolio. The bottom panel of Figure 1 shows the certainty equivalent (CE) return for CRRA investors choosing the optimal allocation between the maximal Sharpe portfolio and the risk-free asset as a function of their relative risk aversion (thin line). Leverage is assumed to be realistically restricted to 50% of the invested wealth.² The risk-free rate and stock return parameters are taken as their monthly sample values over the 2005-2024 period.³ The 100 largest U.S. stocks are included in the analysis. The bold line shows the CE return for CRRA investors choosing the optimal allocation between the MGM portfolio and the risk-free asset. The figure reveals that the MGM portfolio dominates the maximum Sharpe portfolio for all but the most risk-averse investors (who don't wish to borrow anyway).⁴ This result is consistent with the results of Levy (2017), who

¹ If return distributions are normal, P' dominated P not only by the mean-variance criterion, but also by First-order Stochastic Dominance, implying dominance not only for all risk-averse investors, but for all investors with non-decreasing utility functions, including, for example, Prospect Theory investors or investors with various aspiration levels (Levy and Levy 2004, Levy, De Giorgi, and Hens 2012).

² Regulation T caps borrowing at 50% of the investor's invested wealth. Fortune (2000) reports that most investors don't borrow at all to finance their equity investments. He documents that for investors using major brokers such as Merrill Lynch, Paine Webber, and Charles Schwab, the aggregate debt is below 2.5% of the aggregate assets. This very limited borrowing could be due to regulation additional to Reg T, to high interest rates for borrowing, to uncertainty about the future equity return parameters, to a psychological aversion to borrowing, or to some combination of these factors.

³ This should not be taken to imply that we believe that these are the best estimates of the ex-ante parameters. To the contrary, we acknowledge that statistical shrinkage can be helpful, as discussed in Section 4. We employ the sample parameters here only as a "real world" set of parameters for the analysis.

⁴ Most studies empirically and experimentally estimate RRA to be in the range 0-2. See, for example, Arrow (1971) Friend and Blume (1975), Hansen and Singleton (1982), Szpiro (1986a,b), De Mel, McKenzie and Woodruff (2008), and Barro and Jin (2011). Figure 1 shows that the MGM portfolio dominates the maximum Sharpe portfolio for all RRA values up to 6.

analyzes alternative criteria for mutual fund selection, and finds that when borrowing is limited, funds' GMs are much more closely aligned with investors' expected utility than their Sharpe ratios.

Another advantage of GM maximization as an investment objective is that under i.i.d returns the ranking of portfolios by their GMs is invariant to the investment horizon (Levy 2017). This is in contrast to the ranking of portfolios by their Sharpe ratios or their alphas, rankings that depend on the investment horizon, even if returns are i.i.d. (Levy 1972, Levhari and Levy 1977, and Bessembinder, Cooper, and Zhang 2023). Thus, the MGM portfolio is optimal for all investors seeking to maximize their GM, regardless of their investment horizon.

The virtues of the MGM portfolio motivate us to look for a simple and efficient method for constructing this portfolio from a given set of assets. In principle, this can be done via direct numerical optimization. This direct approach encounters several problems. First, when the number of assets is large, the problem is high-dimensional and may involve multiple local maxima. Second, the solution depends on the assumption regarding the shape of the return distributions, which are unknown. Markowitz (2014) considers this point as one of the main advantages of MV analysis over direct expected utility maximization.⁵ Finally, in practical situations where one needs to maximize out-of-sample performance based on sample returns, statistical shrinkage is usually very helpful. While there is extensive literature about shrinkage of the means and the covariance matrix, these are not sufficient to characterize the multivariate return distribution in the general case. Thus, even if the multivariate return distribution is known, the way to apply shrinkage in the context of direct GM maximization is far from obvious.

⁵ Taking the sample returns themselves as representing the return distribution implicitly assumes an equally-likely discrete return distribution. Vander Weide, Peterson, and Maier (1977) and Maier, Peterson, and Vander Weide (1977) discuss finding the MGM in this case.

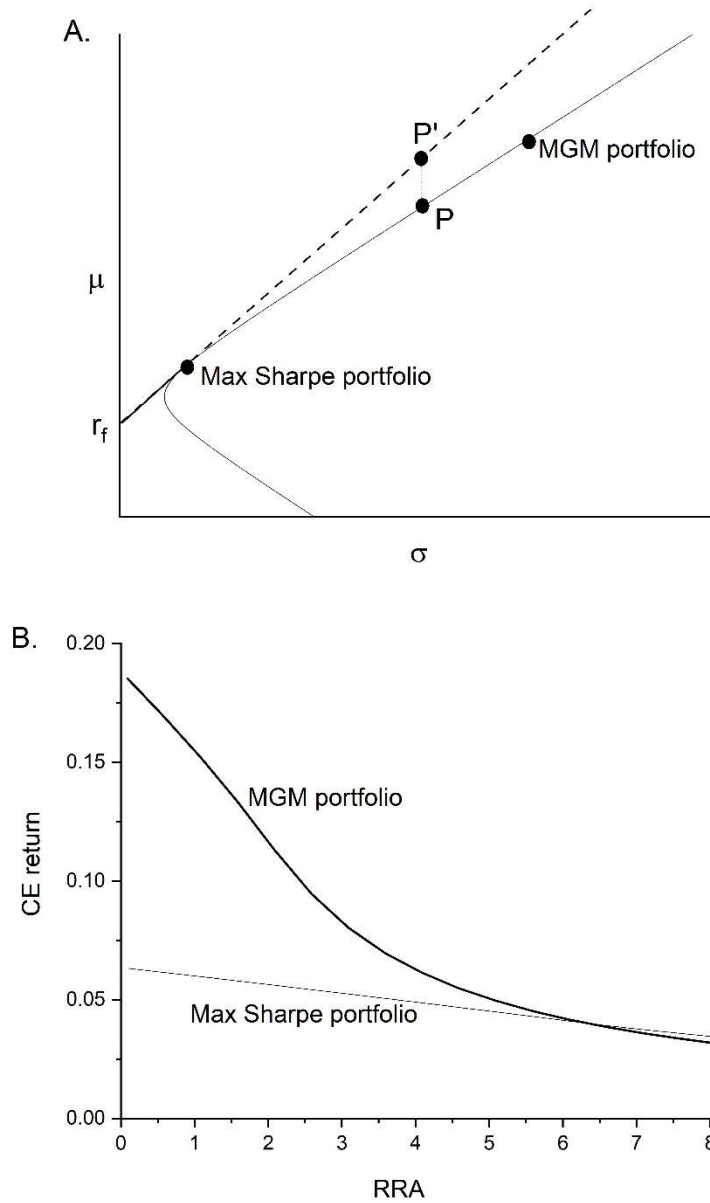


Figure 1. If borrowing is unlimited, the portfolio with the maximal Sharpe ratio dominates any other portfolio, as shown in Panel A: for any portfolio, such as portfolio P , there is a combination of the maximal Sharpe portfolio and the risk-free asset that dominates it, portfolio P' . However, if borrowing is realistically restricted, portfolio P' may be unattainable, and investors may prefer portfolio P over the maximal Sharpe portfolio. The MGM portfolio is typically on the efficient frontier (or very close to it), and has a high expected return (as discussed in Section 2). This implies that most investors, all but the most extremely risk-averse, prefer the MGM portfolio over the maximal Sharpe portfolio. Panel B shows the CE return of these two portfolios as a function of the relative risk aversion, when borrowing is limited to 50% of the investor's wealth.

The approach we take here is different. Rather than numerically looking for the exact MGM, we derive an analytical solution for an approximation of the MGM portfolio, based on the most simple mean-variance approximation to the GM:

$$GM \approx 1 + \mu - \frac{\sigma^2}{2}, \quad (1)$$

where μ is the expected rate of return, and σ^2 is the variance. While there are more elaborate and more precise approximations to the GM (Young and Trent 1976, and Markowitz 2012), we show that the simple approximation (1) yields an elegant closed-form solution for the approximate MGM portfolio that provides an excellent approximation to the exact MGM portfolio. The approximate MGM portfolio is very close to the exact MGM portfolio in mean-variance space, yielding almost the same GM. These results are robust to different assumptions regarding the shape of the return distribution. Finding the approximate MGM is equivalent to finding the mean-variance tangency portfolio, but with the risk-free rate being replaced by r^* , a parameter endogenously determined by the stocks' expected returns and covariances. Other than the elegance of a closed-form solution, the analytical solution of the approximate MGM has two important practical advantages. First, this solution depends only on the expected returns and the covariance matrix, and therefore it does not require making any assumptions about the shape of the multivariate return distribution. Second, and for the same reason, it enables using standard shrinkage methods in settings where one needs to estimate out-of-sample returns.

In the next section we analytically derive the approximate MGM portfolio. Section 3 shows that the approximate MGM portfolio is extremely close to the exact MGM portfolio, and that this result is robust to the return distribution assumed. In Section 4 we illustrate the application of standard shrinkage methods for maximizing the out-of-sample GM. Section 5 concludes.

2. The Approximate MGM Portfolio

Consider a set of N risky assets, with a vector of expected returns $\underline{\mu}$ and a covariance matrix Σ . The MGM portfolio is the combination of these risky assets that yields the maximal GM.⁶ In general, the MGM portfolio can only be found numerically. The *approximate* MGM portfolio is the portfolio maximizing $GM \approx 1 + \mu - \frac{\sigma^2}{2}$, or alternatively, the portfolio with the maximal $\mu - \frac{\sigma^2}{2}$. As this expression is a function of only the portfolio mean and variance, it is natural to employ the standard mean-variance (MV) framework (regardless of the shape of the return distributions). The approximate MGM portfolio is on the MV efficient frontier – for any interior portfolio, there is a portfolio on the efficient frontier directly above it, with the same σ but a higher μ , i.e. with a higher $\mu - \frac{\sigma^2}{2}$. Because the portfolio is on the frontier, finding it corresponds to finding the tangency portfolio for some hypothetical value of the risk-free rate, which we denote by r^* , see Figure 2.

Proposition 1: the weights of the approximate MGM portfolio are given by:

$$x_{aMGM} = \frac{\Sigma^{-1}(\underline{\mu} - r^*)}{\underline{1}' \Sigma^{-1} (\underline{\mu} - r^*)}, \quad (2)$$

where $\underline{1}$ is a vector of 1's, and r^* is given by:

$$r^* = \frac{\underline{1}' \Sigma^{-1} \underline{\mu} - 1}{\underline{1}' \Sigma^{-1} \underline{1}}. \quad (3)$$

See appendix for proof.

⁶ Adding the risk free asset to the analysis, with lending but no borrowing, does not typically lead to a higher GM: for typical parameters, the MGM portfolio is much higher on the efficient frontier than the maximal Sharpe portfolio (see analysis below, and Figure 3 in the next section). Thus, combinations of the MGM portfolio and the risk-free asset are interior portfolios, with lower GMs than the GM of the pure-stock MGM.

Equation (2) shows that finding the approximate MGM portfolio is equivalent to the standard method for finding the tangency portfolio, with the endogenously determined r^* replacing the actual risk-free rate. This closed-form solution for the approximate MGM portfolio is simple and independent of any assumptions about the shape of the return distributions.⁷ Additionally, it allows for the application of standard statistical shrinkage methods when Σ and $\underline{\mu}$ are unknown and need to be estimates from sample data.

The proof in the appendix shows that the slope of the MV frontier at the point of the approximate MGM portfolio is equal to the portfolio's standard deviation (see Figure 2). This implies that for typical return parameters, the approximate MGM portfolio is high on the frontier, with an expected return (and standard deviation) much larger than those of the maximum Sharpe portfolio. To see this, note that as one climbs up on the efficient frontier, the portfolio standard deviation increases, while the frontier slope decreases. The approximate MGM portfolio is located at the point where these two quantities are equal. In contrast, for the maximal Sharpe portfolio, the slope is typically much higher than the portfolio's standard deviation. For example, if we take the S&P 500 as a proxy for the maximal Sharpe portfolio, this portfolio has an excess monthly return of about 0.7% and a monthly standard deviation of 5.4%. Thus, this portfolio's slope (Sharpe ratio) is about 2.4 times larger than its standard deviation.⁸ Hence, one has to move considerably higher on the efficient frontier to get from the market portfolio to the approximate MGM portfolio.

⁷ Mean-variance analysis is often justified by the assumption of normal return distributions. This assumption is obviously unrealistic as it implies the possibility of rates of return lower than -100%. Levy and Markowitz (1979) show that mean-variance analysis is an excellent proxy for expected utility maximization even when the return distributions are not normal.

⁸ Based on the S&P500 monthly returns in the January 1926 – December 2024 sample period. $(\frac{0.007}{0.054})/0.054 = 2.4$.

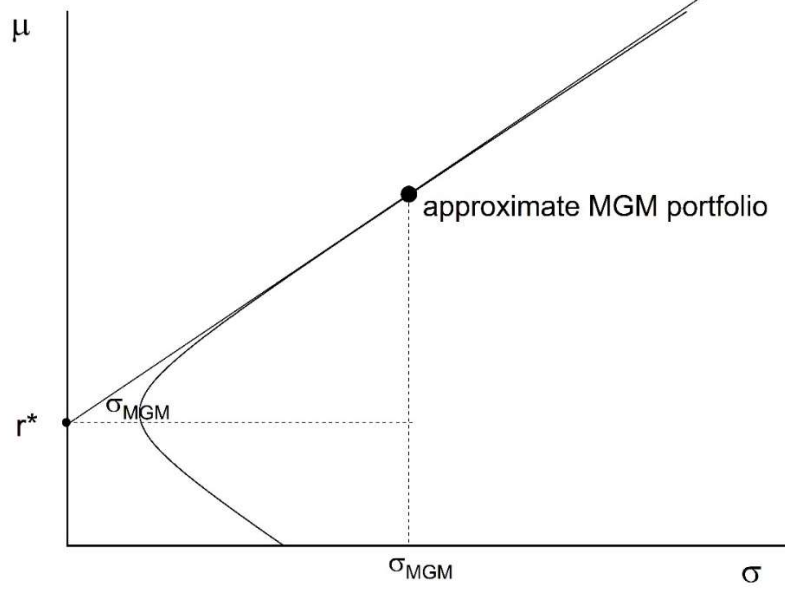


Figure 2. The approximate MGM portfolio is the mean-variance “tangency portfolio”, where the hypothetical r^* , given by eq.(3) replaces the risk-free rate. The slope of the tangency line is equal to the standard deviation of the approximate MGM portfolio (see appendix).

3. Proximity of the Approximate MGM Portfolio to the Exact MGM Portfolio

To examine the proximity of the approximate MGM portfolio to the exact MGM portfolio we conduct a numerical analysis. In this analysis we assume that the return distributions are known – the issue of out-of-sample estimation is discussed in the next section. In order to obtain realistic return distributions, we employ the sample monthly returns over the 20-year period from January 2005 to December 2024. We first take the asset universe as the 100 largest U.S. stocks with complete monthly return records over this sample period (we later look at additional asset sets).⁹ Based on these returns we calculate the vector of sample average returns $\underline{\mu}^{sample}$, and the sample covariance matrix Σ^{sample} . It is well-known that the sample parameters are not the best estimates of the true parameters, and that better estimates can be obtained by applying statistical shrinkage

⁹ The survivorship bias introduced by this does not trouble us, because our only goal in the present analysis is to obtain a “reasonable” set of return parameters.

(Stein 1956, James and Stein 1961). There are different ways to “shrink” the sample parameters, as discussed in the next section, where we examine the shrinkage that produces the best out-of-sample performance. For the present purpose, consider the simple shrinkage:

$$\underline{\mu} = \gamma_{\mu} \underline{\mu}^{sample} + (1 - \gamma_{\mu}) \mu^{ave}, \text{ and:} \quad (4)$$

$$\Sigma = \gamma_{\sigma} \Sigma^{sample} + (1 - \gamma_{\sigma}) \Sigma^{target}, \quad (5)$$

where μ^{ave} is the cross-sectional mean average sample return, and the target covariance matrix Σ^{target} , has the mean sample variance on the diagonal, and the mean sample covariance on the off-diagonal (see Ledoit 1995, Ledoit and Wolf 2004b, and Wolf and Wunderli 2012). γ_{μ} and γ_{σ} are the shrinkage intensities. For now, we take typical values suggested in the literature, $\gamma_{\mu} = \gamma_{\sigma} = 0.2$ (see, for example, Pedersen, Babu and Levine 2021). In the next section we discuss more elaborate shrinkage formulations, and analyze the optimal shrinkage intensity for maximizing the out-of-sample GM.

The approximate MGM portfolio depends only on the estimates of the expected returns and covariance matrix, as evident from equations (2) and (3). In contrast, in order to numerically find the exact MGM, one must specify the shape of the return distributions, which are typically unknown (taking the sample returns themselves implicitly assumes an equally-likely discrete distribution, as analyzed by Vander Weide, Peterson, and Maier 1977). We examine three different cases of the multivariate return distributions: normal, lognormal, and the Student t distribution. For each of these cases, we draw 100,000 random returns for each stock from a multivariate distribution with means and a covariance matrix given by equations (4) and (5). We then

numerically search for the vector of weights x_{MGM} that maximizes the portfolio's GM.¹⁰ Note that in the case of the normal and Student t distributions negative total returns (i.e. rates of return below -100%) are theoretically possible, in which case the GM is undefined. In fact, for *any* portfolio with a non-zero weight in *any* of the stocks, these distributions imply the possibility of rates of return lower than -100%, which are realistically impossible.¹¹ This is a well-known problematic aspect of the normal (and Student t) distribution. To rule out negative total returns, we are formally considering distributions truncated at -100%. However, for realistic parameters this truncation is not a practical concern: even for a rather volatile stock with a monthly return standard deviation of 10%, a return below -100% represents a $10\text{-}\sigma$ event, an event for which the normal distribution implies a probability lower than 10^{-22} . In our numerical simulations the truncation is never needed.

The results are described in Figure 3. For all three distributions the approximate MGM portfolio and the exact MGM portfolio are very close to one another, and they are both located on the mean-variance frontier. In all three cases, the difference in the GMs of the two portfolios is very small: it is less than 0.0001% in the case of the normal and the lognormal distributions, and it is 0.013% in the case of the Student's t distribution with $\nu = 5$ degrees of freedom.¹² As discussed in the previous section, the MGM portfolios are much higher on the efficient frontier

¹⁰ We employ Matlab's *fmincon* optimization function. We employ different starting points in the optimization, and find that they lead to the same MGM portfolio. Thus, it seems that at least in this setting, the local maxima problem is not severe.

¹¹ Rates of return below -100% are in practice ruled out for individual stocks by limited liability. For portfolios, which may include short positions, such negative total returns are typically ruled out by margin requirements.

¹² This result should not be surprising: Levy and Markowitz (1979) show that the expected utility implied by a portfolio can be very closely approximated by a function of the portfolio's mean and variance. The portfolio's GM is equal to the exponent of the expected utility of a Bernoulli investor with log utility: $EU = \sum_{i=1}^n p_i \log(R_i) = \log(\prod_{i=1}^n R_i^{p_i}) = \log(GM)$, where the R_i 's are the possible total returns on the portfolio (1+rate of returns), and the p_i 's are the corresponding probabilities.

than the maximum Sharpe portfolio (calculated with the average monthly risk-free rate in our sample period, which is 0.13%).

These results are quite general. When instead of the 100 largest stocks we randomly draw 100 sets of 100 stocks (out of the set of all stocks with complete return records over the sample period) and calculate the exact MGM and approximate MGM portfolios for each set, we find that the differences in the GMs of these two portfolios is very small. For the case of the normal distribution, the average difference between the GM of the exact and the approximate MGM portfolios (across the all stock sets) is $3.7 \cdot 10^{-5}\%$. The median value is $3.1 \cdot 10^{-5}\%$, and the maximal value obtained in all of the random stocks sets is $1.4 \cdot 10^{-4}\%$. Similar results are obtained for the lognormal distribution, with a mean difference value of $9.2 \cdot 10^{-5} \%$, a median value of $8.6 \cdot 10^{-5}\%$ and a maximal value of $2.1 \cdot 10^{-4}\%$. The differences are larger for the Student's t distribution ($\nu = 5$), but they are still very small: the mean difference is 0.012%, the median value is 0.011%, and the maximal value is 0.034%. Thus, the cost of using the approximate MGM instead of the exact MGM, in terms of the GM, is economically negligible.

The approximation $GM \approx 1 + \mu - \frac{\sigma^2}{2}$, which is employed in the derivation of the approximate MGM, breaks down when returns are large in absolute value. Thus, we expect the performance of the approximate MGM to deteriorate for portfolios with extreme returns. What is the range of parameters for which the approximate MGM still performs well? In order to examine this question, we increase the γ 's in eq.(4-5), reducing the shrinkage intensities. It is well-known that when the number of stocks is large, the raw sample parameters (i.e. without shrinkage) lead to mean-variance efficient portfolios with very extreme portfolio weights and portfolio returns (Best and Grauer 1991, Black and Litterman 1992, Jagannathan and Ma 2003, Levy and Ritov 2011). When

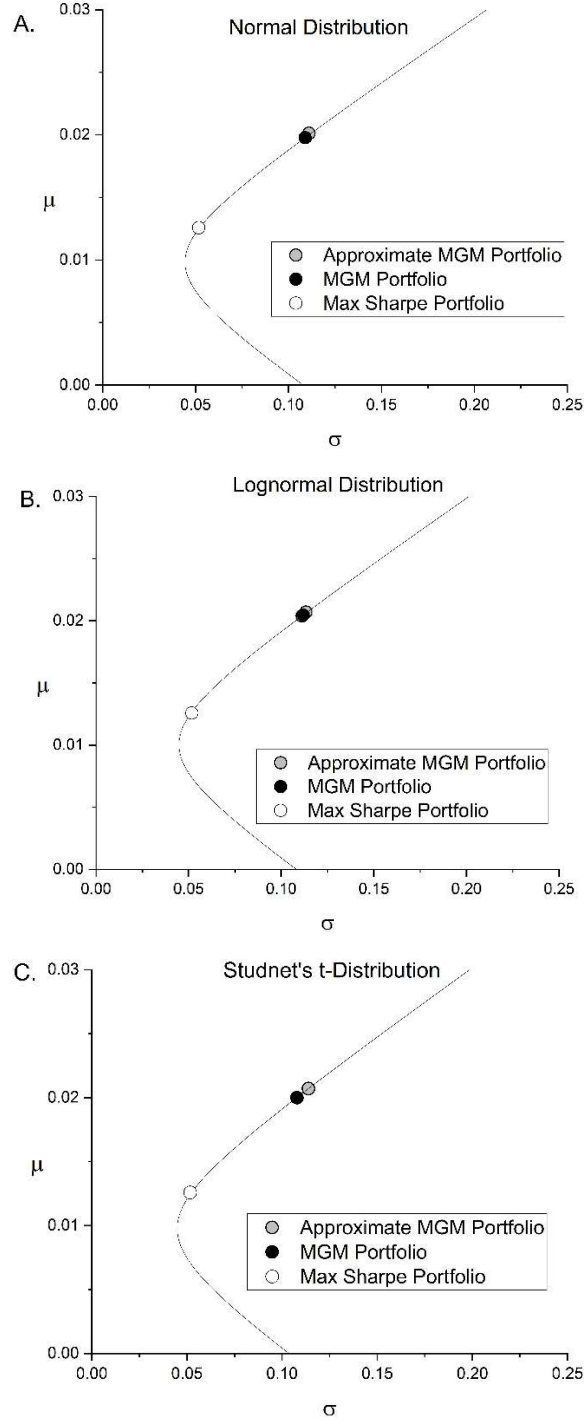


Figure 3. The exact and approximate MGM portfolios under different assumptions regarding the multivariate return distribution: normal (Panel A), lognormal (Panel B), and the Student t distribution with $\nu = 5$ degrees of freedom (Panel C). Monthly returns are estimated based on the January 2005 to December 2024 sample period, and formulas (4-5). The asset set is comprised of the largest 100 U.S. stocks, by December 2024 values. The difference in the GMs of the MGM and approximate MGM portfolios is 0.013% in the case of the Student's t distribution, and less than 0.001% in the cases of the normal and the lognormal. The MGM portfolios have substantially higher means and variances relative to the maximum Sharpe portfolio.

we reduce the shrinkage intensity, the portfolio weights and returns become extreme, and the performance of the approximate MGM portfolio deteriorates, as expected. Figure 4 reports the difference between the GMs of the exact MGM portfolio and the approximate MGM portfolio as a function of the root mean square portfolio weight, defined as:

$$RMS = \sqrt{\frac{1}{N} \sum_{i=1}^N x_i^2}, \quad (6)$$

where N is the number of stocks and x_i denotes the weight of stock i in the MGM portfolio. The results are based on 100 sets of 100 randomly drawn stocks. For each of these sets, the MGM portfolio and the GM differences are based on drawing 100,000 returns from a multivariate normal distribution with expected returns and covariances given by eqs.(4-5).

When the portfolio weights (and therefore portfolio returns) become more extreme, the difference between the GMs of the MGM and the approximate MGM portfolios increases. The dramatic increase occurs at a RMS value of about 0.1. This value implies rather extreme portfolio weights – the typical holding in each stock (in absolute value) is about 0.1 (or 10% of the total wealth). For comparison, the RMS value of the S&P500 index, which is considered to have rather skewed portfolio weights, is only 0.006 (indicated in the figure by the vertical dashed line).¹³ Hence, the deterioration in the performance of the approximate MGM portfolio occurs only for portfolios with extreme weights, which are practically not very relevant. For reasonably well-diversified portfolios, the approximate MGM portfolio yields almost the same GM as the exact MGM portfolio.

¹³ As of December 2024, see, for example, <https://finance.yahoo.com/quote/SPY/holdings/>.

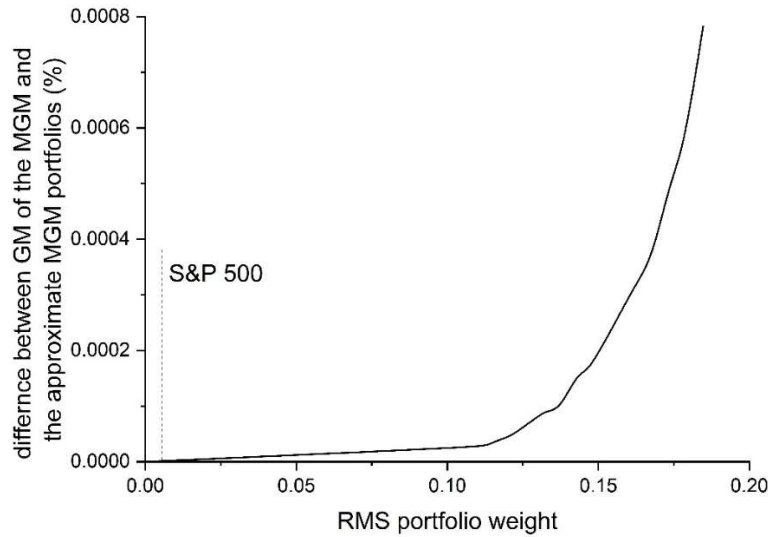


Figure 4. The difference between the GM of the exact MGM portfolio and the GM of the approximate MGM portfolio as a function of the Root Mean Square portfolio weight of the MGM portfolio. The difference is very small up to a RMS of about 0.1. This value indicates rather extreme portfolio weights – for comparison, the RMS of the S&P500 index is only 0.006 (indicated by the vertical dashed line).

4. Statistical Shrinkage

The sample value of a stochastic variable is typically not the best estimate of its true value. Instead, the best estimate is obtained by “shrinking” the sample value towards the cross-sectional average value, or another prior (Stein 1956, James and Stein 1961). One of the advantages of the approximate MGM portfolio is that it depends only on the mean returns and the covariance matrix, and there are standard methods for applying shrinkage to these. In contrast, a numerical solution of the exact MGM portfolio requires the set of all returns on all stocks. In this case, it is not obvious how one is to apply shrinkage. Suppose, for example, that we wish to find the MGM portfolio given a matrix of T sample return observations on N stocks. It is straightforward to apply shrinkage to the expected returns: we can add/subtract a constant from all of the sample returns of a given stock, in order to adjust its average return to any desired value. However, how should one alter the

matrix of T by N returns to achieve the desired covariance matrix? This is far from obvious.¹⁴ The approximate MGM approach circumvents this problem. This section illustrates the application of shrinkage in the context of GM maximization. Subsection 4.1 provides a brief background on statistical shrinkage. Section 4.2 describes the shrinkage formulas that we apply, and in Section 4.3 we empirically analyze the shrinkage intensities that work best, in the sense of producing the highest out-of-sample portfolio GM.

4.1 Background

The intuition for shrinkage can be illustrated with CAPM betas. Suppose that we observe a stock with a sample beta of 2. It is more likely that the stock's true beta is 1.8 and random estimation noise "pushed the estimate up" to 2, than it is likely that the true beta is 2.2 and noise "pushed the estimate down" to 2 – simply because there are more stocks in the population with a beta of 1.8 than there are with a beta of 2.2. Thus, our best estimate of this stock's beta will be somewhat lower than the sample value of 2. The opposite holds for a stock with a sample beta of 0.3: in this case it is more likely that the true beta is 0.5 than it is 0.1, because there are more stocks in the population with a beta of 0.5 than there are with a beta of 0.1. In this case, our best estimate will be somewhat higher than the sample value of 0.3. Thus, in general, our best estimate is obtained by "shrinking" the sample value towards the cross-sectional average value (which is 1 in the case of betas). The application of shrinkage to sample betas is discussed in Vasicek 1973, Levi and Welch 2017, and Welch 2022, and is employed in Bloomberg's beta estimates.¹⁵ Similarly,

¹⁴ If the shape of the multivariate distribution is known, and it depends only on the expected returns and the covariance matrix, one can draw random observations from this distribution with the shrunk parameters. This is what we did in the numerical analysis in Section 3. However, this requires knowledge of the return distribution, and requires that the distribution is completely described by the means and the covariance matrix. The approximate MGM approach, in contrast, does not depend on these assumptions.

¹⁵ See, for example, <https://guides.lib.byu.edu/c.php?g=216390&p=1428678>.

shrinkage should also be applied to expected returns (Jorion 1986, 1991), and to the covariance matrix (Ledoit and Wolf 2004a,b, 2017, 2022), as described below.

4.2 Shrinkage Formulas

Consider an investor who observes T returns of stock i , with a sample average $\hat{\mu}_i$ (which is shorthand notation for μ_i^{sample}). The investor forms his posterior belief regarding the expected return based on the sample returns and the investor's prior. For normally distributed returns, and a normal prior with mean μ and standard deviation σ_μ , the investor's posterior expected return for stock i is given by:

$$E(\mu_i | \hat{\mu}_i) = \frac{\frac{T}{\sigma_i^2} \hat{\mu}_i + \frac{1}{\sigma_\mu^2} \mu}{\left(\frac{T}{\sigma_i^2} + \frac{1}{\sigma_\mu^2} \right)}, \quad (7)$$

where σ_i is the standard deviation of stock i 's returns (see, for example, Lee 2012, p. 46). Eq.(7) can also be written as:

$$E(\mu_i | \hat{\mu}_i) = \gamma_i \hat{\mu}_i + (1 - \gamma_i) \mu, \quad (8)$$

where γ_i is given by:

$$\gamma_i \equiv \frac{\frac{T}{\sigma_i^2}}{\frac{T}{\sigma_i^2} + \frac{1}{\sigma_\mu^2}} = \frac{1}{1 + \frac{\sigma_i^2}{T \sigma_\mu^2}}. \quad (9)$$

γ_i determines how much the sample average return is shrunk towards the prior. The shrinkage formula in eq.(9) can be intuitively interpreted as follows: the ex-ante expected return of stock i is estimated as a weighted average of the stock's sample average return, and the prior μ . The higher the standard deviation of the stock, the more "noisy" its sample estimate, and therefore the less

weight is attached to its sample return (i.e. a high σ_i implies a low γ_i , see eq.(9)).¹⁶ $\gamma_i = 1$ implies no shrinkage at all, while at the other extreme $\gamma_i = 0$ implies that the sample returns are completely ignored. The derivation of eq.(7) assumes that the return distributions are normal, and, more importantly, that they are stable over time. In practice, these assumptions may not hold, and a better estimation of the stock's future expected return may be obtained with a higher (or lower) shrinkage intensity (i.e. with a lower or higher value of γ_i). This idea is employed in Levi and Welch (2017) in the context of the estimation of betas, and in Levy and Roll (2023, 2024) in the context of the estimation of Sharpe ratios. In order to allow for “extra shrinkage” relative to equations (8-9), we introduce a market-wide shrinkage parameter γ_μ defined on the interval $[0,1]$, and we generalize eq.(9) to:

$$\gamma_i = \frac{1}{1 + \frac{\sigma_i^2}{T\sigma_\mu^2}\left(\frac{1}{\gamma_\mu}-1\right)}. \quad (10)$$

The standard eq.(9) is obtained as a special case when $\gamma_\mu = \frac{1}{2}$. Values of γ_μ lower than $\frac{1}{2}$ imply more shrinkage relative to eq.(9) (i.e. a lower γ_i), and values higher than $\frac{1}{2}$ imply less shrinkage. In the extremes, $\gamma_\mu = 0$ implies $\gamma_i = 0$, i.e. this is the case of maximal shrinkage where the sample average return is completely ignored; $\gamma_\mu = 1$ implies $\gamma_i = 1$, i.e. the case of no shrinkage at all. In the empirical analysis described below we investigate the value of γ_μ that yields the best out-of-sample portfolio GM. We take the prior expected return μ and the cross-sectional standard deviation of expected returns σ_μ as the average monthly return of all stocks included in our

¹⁶ This is an improvement over the simplistic eq.(4), which applies the same shrinkage intensity to all stocks, regardless of their volatilities.

analysis, and the cross-sectional standard deviation of the average returns, which are 1.48% and 1.74%, respectively.

There are different alternatives for shrinking the covariance matrix. One alternative is to take Σ^{target} in eq.(5) as the identity matrix times a constant (Ledoit and Wolf 2004a). Another is to employ a factor model to construct Σ^{target} (Ledoit and Wolf 2003).¹⁷ Here we follow Ledoit (1995), Ledoit and Wolf (2004b), and Wolf and Wunderli (2012), and employ eq.(5) with Σ^{target} taken as a matrix with the mean sample variance on the diagonal, and the mean sample covariance on the off-diagonal.

When returns are i.i.d., it is possible to analytically determine the optimal shrinkage intensity γ_σ . However, when the return parameters change over time by an unknown process, this is not possible (similar to the case of the expected returns). In the empirical analysis below we will look for the combination of shrinkage intensities γ_μ and γ_σ that yield the highest out-of-sample GM.

4.3 Empirical Analysis

To find which $(\gamma_\mu, \gamma_\sigma)$ combination yields the highest out-of-sample GM, we form the approximate MGM portfolio based on the shrunk sample estimates (using equations 5 and 8-10), with the $(\gamma_\mu, \gamma_\sigma)$ shrinkage intensities. We do this for all $\gamma_\mu, \gamma_\sigma$ values from 0 to 1, with 0.001 intervals (i.e. we examine 1,000,000 $(\gamma_\mu, \gamma_\sigma)$ combinations), recording the out-of-sample performance of each combination.

¹⁷ Another approach is to separately shrink the correlation matrix and the standard deviations, see, for example, Pedersen, Babu and Levine (2021). The shrinkage in eq.(5) is linear. See Ledoit and Wolf (2017) for an extension to non-linear shrinkage.

We employ the January 1975 – December 2024 sample period (600 months). At any month t we estimate the sample parameters based on the previous 60 months (i.e. months $t-1, t-2, \dots, t-60$), and record the approximate MGM portfolio's out-of-sample performance in the following 12 months (periods $t, t+1, \dots, t+11$). Every month, we take all available U.S. common stocks (CRSP share codes 10 and 11) with complete return records in months $t-60$ to $t+11$. From this population we draw 10 random sets of 100 stocks. For each set we form the approximate MGM portfolio based on the shrunk parameters with $(\gamma_\mu, \gamma_\sigma)$, and record the portfolio's 12 out-of-sample returns. Thus, for every $(\gamma_\mu, \gamma_\sigma)$ combination and every month we record the out-of-sample returns of 10 random portfolios (for consistency, the same random sets of 100 stocks are used for all $(\gamma_\mu, \gamma_\sigma)$ combinations). The GM of every $(\gamma_\mu, \gamma_\sigma)$ combination is calculated over all of the strategy's returns (in all months and all random sets). The results are shown in Figure 5. GMs are annualized and reported in decimal form. Light colors represent higher out-of-sample GMs, as indicated by the vertical bar on the right.

The highest out-of-sample GM, 14%, is obtained with the combination $(\gamma_\mu = 0.004, \gamma_\sigma = 0.060)$. These values imply that very little weight is assigned to the sample parameters (recall that $\gamma_\mu = 0$ and $\gamma_\sigma = 0$ implies ignoring the sample values altogether). This is consistent with previous studies showing that the sample parameters are very noisy estimates of the out-of-sample parameters (see, for example, Carhart 1997, Jagannathan and Ma 2003, and Levy and Roll 2024). Still, the sample parameters do contain useful information: the optimum is obtained when they are given a positive weight. The out-of-sample GM of the zero-information equally-weighted portfolio $(\gamma_\mu = 0, \gamma_\sigma = 0)$ is 13.8%.

It may certainly be possible that more sophisticated shrinkage methods can improve the out-of-sample GM. The point of this exercise is to demonstrate that the analytical solution for the approximate MGM, which relies only on the expected returns and covariance matrix, allows for a straightforward application of standard shrinkage methods, regardless of the specific shrinkage method employed.

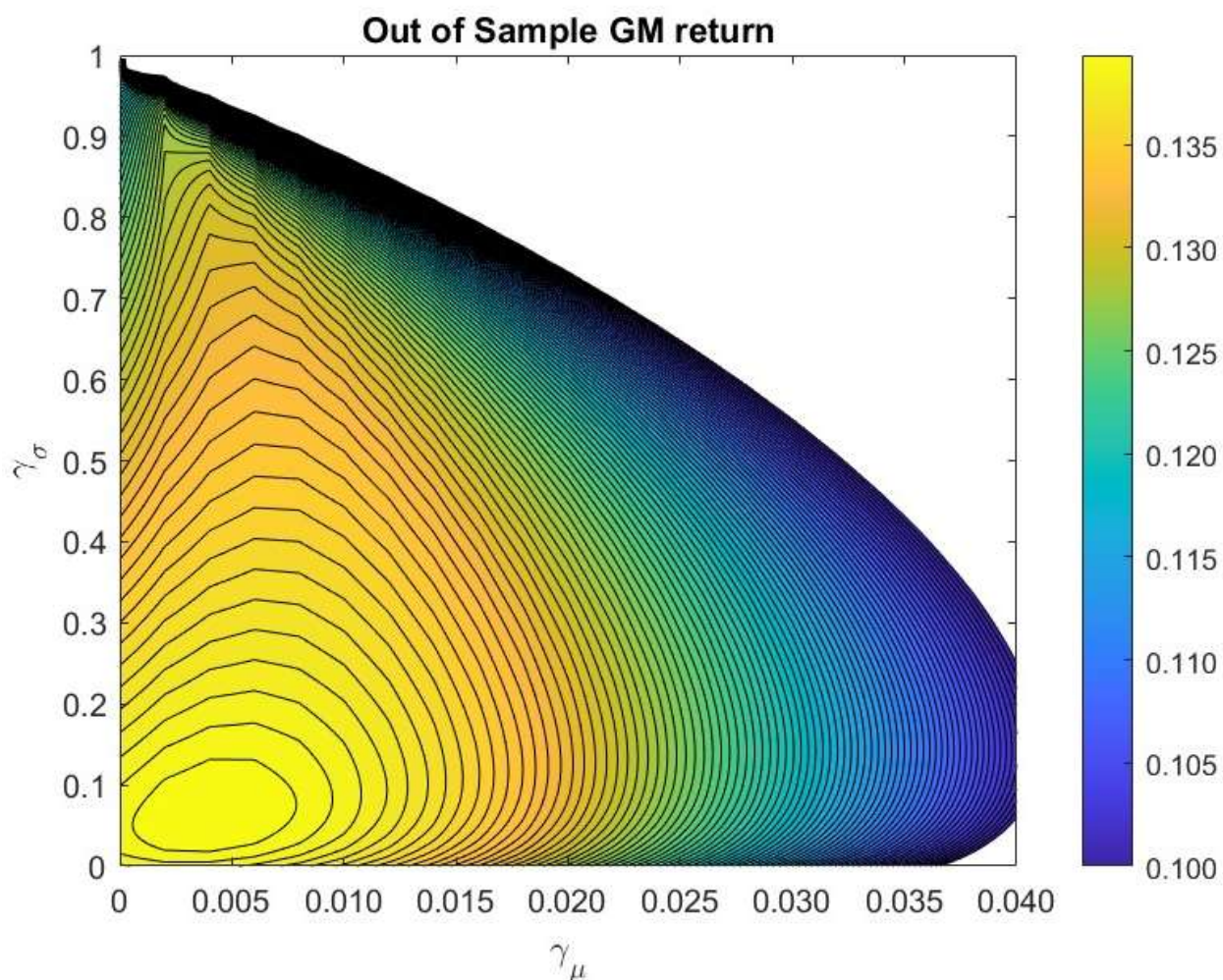


Figure 5. The annualized out-of-sample GM as a function of the shrinkage intensities (γ_μ , γ_σ). The highest GM is obtained at the combination ($\gamma_\mu = 0.004$, $\gamma_\sigma = 0.060$). The in-sample parameters are based on the preceding 60 months, and the out-of-sample returns are calculated over the following 12 months.

5. Conclusions

The portfolio with the Maximal Geometric Mean (MGM) has the following well-known property: as the investment horizon increases, this portfolio yields a higher terminal wealth than any other portfolio, with a probability that approaches 1. This property makes the MGM portfolio an attractive choice for long-run investors. The MGM portfolio is also advantageous for investors with short or intermediate horizons: when borrowing is realistically limited, this portfolio yields a higher expected utility than the maximal Sharpe portfolio, for all but the most extremely risk-averse investors. These virtues make GM maximization an important goal for investors and fund managers. How can the MGM portfolio be constructed in practice? This paper addresses this question.

Given a set of risky assets, one could, in principle, seek the combination that yields the MGM numerically. However, this is a high-dimensional problem with potentially many local maxima. In addition, it requires an assumption about the shape of the return distributions, which are unknown. Moreover, in practical applications where the ex-ante parameters are estimated from ex-post returns, it is unclear how to apply statistical shrinkage to the sample returns. Here we suggest a different approach. We derive an analytical solution for the approximate MGM portfolio, based on the mean-variance approximation of the GM. This approach is very simple: it is equivalent to deriving the tangency portfolio, with the difference that a quantity r^* , endogenously determined by the stocks expected returns and covariances, replaces the risk-free rate. We show that the approximate MGM approach yields a portfolio very close to the exact MGM, and this holds true for different assumptions regarding the return distributions. The advantages of the analytic approximation is that it does not require knowledge of the return distribution, and it allows for the

application of standard methods of statistical shrinkage of the covariance matrix and the expected returns.

Harry Markowitz invented the mean-variance framework, which has become the foundation of modern portfolio theory. He was also a strong advocate of Geometric Mean maximization as an investment objective. The present paper shows that these two themes are harmonically intertwined – mean-variance optimization offers an efficient and elegant way to achieve the goal of maximizing the Geometric Mean.

References

- Arrow, K.J., 1971. *Essays in the Theory of Risk-Bearing*. Amsterdam: North-Holland.
- Bali, T.G., Demirtas, K.O., Levy, H. and Wolf, A., 2009. Bonds versus stocks: Investors' age and risk taking. *Journal of Monetary Economics*, 56(6), pp.817-830.
- Barro, R.J. and Jin, T., 2011. On the size distribution of macroeconomic disasters. *Econometrica*, 79(5), pp.1567-1589.
- Bessembinder, H., Cooper, M.J. and Zhang, F., 2023. Mutual fund performance at long horizons. *Journal of Financial Economics*, 147(1), pp.132-158.
- Best, M.J. and Grauer, R.R., 1991. On the sensitivity of mean-variance-efficient portfolios to changes in asset means: some analytical and computational results. *The Review of Financial Studies*, 4(2), pp.315-342.
- Black, F. and Litterman, R., 1992. Global portfolio optimization. *Financial Analysts Journal*, 48(5), pp.28-43.
- Carhart, M.M., 1997. On persistence in mutual fund performance. *The Journal of Finance*, 52(1), pp.57-82.
- De Mel, S., McKenzie, D. and Woodruff, C., 2008. Returns to capital in microenterprises: evidence from a field experiment. *The Quarterly Journal of Economics*, 123(4), pp.1329-1372.
- Estrada, J., 2010. Geometric mean maximization: an overlooked portfolio approach? *Journal of Investing*, 19(4), p.134.
- Friend, I. and Blume, M.E., 1975. The demand for risky assets. *The American Economic Review*, 65(5), pp.900-922.
- Fortune, P. 2000. Margin requirements, margin loans, and margin rates: practice and principles. *New England Economic Review*, 19.
- Friend, I. and Blume, M.E., 1975. The demand for risky assets. *The American Economic Review*, 65(5), pp.900-922.
- Hansen, L.P. and Singleton, K.J., 1982. Generalized instrumental variables estimation of nonlinear rational expectations models. *Econometrica*, 50(5), pp.1269-1286.
- Jagannathan, R. and Ma, T., 2003. Risk reduction in large portfolios: Why imposing the wrong constraints helps. *The Journal of Finance*, 58(4), pp.1651-1683.
- James, W. and Stein, C., 1961. Estimation with Quadratic Loss. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability, Volume 1: Contributions to the Theory of Statistics*.

- Jensen, M.C., 1968. The performance of mutual funds in the period 1945-1964. *The Journal of Finance*, 23(2), pp.389-416.
- Jorion, P., 1986. Bayes-Stein estimation for portfolio analysis. *Journal of Financial and Quantitative Analysis*, 21(3), pp.279-292.
- Jorion, P., 1991. Bayesian and CAPM estimators of the means: Implications for portfolio selection. *Journal of Banking & Finance*, 15(3), pp.717-727.
- Kelly, J.L., 1956. A new interpretation of information rate. *The Bell System Technical Journal*, 35(4), pp.917-926.
- Latane, H.A., 1959. Criteria for choice among risky ventures. *Journal of Political Economy*, 67(2), pp.144-155.
- Ledoit, O., 1995. *Essays on risk and return in the stock market* (Doctoral dissertation, Massachusetts Institute of Technology).
- Ledoit, O. and Wolf, M., 2003. Improved estimation of the covariance matrix of stock returns with an application to portfolio selection. *Journal of Empirical Finance*, 10(5), pp.603-621.
- Ledoit, O. and Wolf, M., 2004a. A well-conditioned estimator for large-dimensional covariance matrices. *Journal of Multivariate Analysis*, 88(2), pp.365-411.
- Ledoit, O. and Wolf, M., 2004b. Honey, I shrunk the sample covariance matrix. *The Journal of Portfolio Management*, 30(4), pp.110-119.
- Ledoit, O. and Wolf, M., 2017. Nonlinear shrinkage of the covariance matrix for portfolio selection: Markowitz meets Goldilocks. *The Review of Financial Studies*, 30(12), pp.4349-4388.
- Ledoit, O. and Wolf, M., 2022. The power of (non-) linear shrinking: A review and guide to covariance matrix estimation. *Journal of Financial Econometrics*, 20(1), pp.187-218.
- Lee, P.M., 2012. *Bayesian Statistics*. London: Oxford University Press.
- Leshno, M. and Levy, H., 2002. Preferred by “all” and preferred by “most” decision makers: Almost stochastic dominance. *Management Science*, 48(8), pp.1074-1085.
- Levhari, D. and Levy, H., 1977. The capital asset pricing model and the investment horizon. *The Review of Economics and Statistics*, pp.92-104.
- Levi, Y. and Welch, I., 2017. Best practice for cost-of-capital estimates. *Journal of Financial and Quantitative Analysis*, 52(2), pp.427-463.
- Levy, H., 1972. Portfolio performance and the investment horizon. *Management Science*, 18(12), pp. 645-653.
- Levy, H., 2016. Aging population, retirement, and risk taking. *Management Science*, 62(5), pp.1415-1430.

Levy, H., 2024a. Mean–Variance Analysis, the Geometric Mean, and Horizon Mismatch. *Journal of Portfolio Management*, 50(8).

Levy, H., 2024b. The maximum geometric mean criterion: revisiting the Markowitz–Samuelson debate: survey and analysis. *Annals of Operations Research*, pp.1-22.

Levy, H., De Giorgi, E.G. and Hens, T., 2012. Two paradigms and Nobel prizes in economics: a contradiction or coexistence? *European Financial Management*, 18(2), pp.163-182.

Levy H., and M. Levy, 2004. Prospect theory and mean-variance analysis, *Review of Financial Studies*, 17(4), pp.1015-1041.

Levy, H. and Markowitz, H.M., 1979. Approximating expected utility by a function of mean and variance. *The American Economic Review*, 69(3), pp.308-317.

Levy, M., 2017. Measuring portfolio performance: Sharpe, alpha, or the geometric mean? *Journal of Investment Management*, 15(3), pp.1-17.

Levy, M. and Ritov, Y.A., 2011. Mean–variance efficient portfolios with many assets: 50% short. *Quantitative Finance*, 11(10), pp.1461-1471.

Levy, M. and Roll, R., 2023. The Shrinkage Adjusted Sharpe Ratio: An Improved Method for Mutual Fund Selection. *The Journal of Investing*, 32(2), pp.7-23.

Levy, M. and Roll, R., 2024. Mutual Fund Selection. *Springer Books*.

Lo, A.W., Orr, H.A. and Zhang, R., 2018. The growth of relative wealth and the Kelly criterion. *Journal of Bioeconomics*, 20, pp.49-67.

Maier, S.F., Peterson, D.W. and Vander Weide, J.H., 1977. A monte carlo investigation of characteristics of optimal geometric mean portfolios. *Journal of Financial and Quantitative Analysis*, 12(2), pp.215-233.

MacLean, L.C., Ziemba, W.T. and Blazenko, G., 1992. Growth versus security in dynamic investment analysis. *Management Science*, 38(11), pp.1562-1585.

Markowitz, H.M., 1976. Investment for the long run: New evidence for an old rule. *The Journal of Finance*, 31(5), pp.1273-1286.

Markowitz, H.M., 2006. *Samuelson and investment for the long run*. In Szenberg M, Ramrattan L, Gottesman AA, eds. *Samuelsonian Economics and the Twenty-First Century* (pp. 252-261). Oxford: Oxford University Press.

Markowitz, H., 2012. Mean-variance approximations to the geometric mean. *Annals of Financial Economics*, 7(01), p.1250001.

Markowitz, H., 2014. Mean–variance approximations to expected utility. *European Journal of Operational Research*, 234(2), pp.346-355.

- Merton, R.C., 1972. An analytic derivation of the efficient portfolio frontier. *Journal of Financial and Quantitative Analysis*, 7(4), pp.1851-1872.
- Merton, R.C. and Samuelson, P.A., 1974. Fallacy of the log-normal approximation to optimal portfolio decision-making over many periods. *Journal of Financial Economics*, 1(1), pp.67-94.
- Pedersen, L.H., Babu, A. and Levine, A., 2021. Enhanced portfolio optimization. *Financial Analysts Journal*, 77(2), pp.124-151.
- Samuelson, P.A., 1971. The “fallacy” of maximizing the geometric mean in long sequences of investing or gambling. *Proceedings of the National Academy of Sciences*, 68(10), pp.2493-2496.
- Stein, C., 1956, January. Inadmissibility of the usual estimator for the mean of a multivariate normal distribution. In *Proceedings of the Third Berkeley Symposium on Mathematical Statistics and Probability* (Vol. 1, No. 1, pp. 197-206).
- Szpiro, G.G., 1986a. Measuring risk aversion: an alternative approach. *The Review of Economics and Statistics*, pp.156-159.
- Szpiro, G.G., 1986b. Relative risk aversion around the world. *Economics Letters*, 20(1), pp.19-21.
- Vander Weide, J.H., Peterson, D.W. and Maier, S.F., 1977. A strategy which maximizes the geometric mean return on portfolio investments. *Management Science*, 23(10), pp.1117-1123.
- Vasicek, O.A., 1973. A note on using cross-sectional information in Bayesian estimation of security betas. *The Journal of Finance*, 28(5), pp.1233-1239.
- Welch, I. 2022. Simply better market betas. *Critical Finance Review*, 11, pp. 37–64.
- Wolf, M., and D. Wunderli. 2012. Fund-of-funds construction by statistical multiple testing methods. In B. Scherer and K. Winston (eds.), *The Oxford Handbook of Quantitative Asset Management*, pp. 116–135. Oxford: Oxford University Press
- Young, W.E. and Trent, R.H., 1969. Geometric mean approximations of individual security and portfolio performance. *Journal of Financial and Quantitative Analysis*, 4(2), pp.179-199.
- Ziemba, W.T., 2015. A response to Professor Paul A. Samuelson's objections to Kelly capital growth investing. *Journal of Portfolio Management*, 42(1), p.153.

Appendix – Proof of Proposition 1

Merton (1972) shows that the MV frontier is given by:

$$\sigma^2 = \frac{C\mu^2 - 2A\mu + B}{D}, \quad (A1)$$

where A , B , C , and D are scalars defined as:

$$A \equiv \underline{1}' \Sigma^{-1} \underline{\mu}, \quad B \equiv \underline{\mu}' \Sigma^{-1} \underline{\mu}, \quad C \equiv \underline{1}' \Sigma^{-1} \underline{1}, \quad D \equiv BC - A^2 \quad (A2)$$

(see eq.12 in Merton 1972). Thus, for portfolios on the MV frontier we have:

$$\mu - \frac{\sigma^2}{2} = \mu - \frac{C\mu^2 - 2A\mu + B}{2D} = \left(1 + \frac{A}{D}\right) \mu - \frac{C}{2D} \mu^2 - \frac{B}{2D}. \quad (A3)$$

The maximum value of the expression $\mu - \frac{\sigma^2}{2}$ is obtained for the portfolio on the MV efficient frontier with mean return of:

$$\begin{aligned} \frac{C}{D} \mu_{MGM} &= 1 + \frac{A}{D}, \quad \text{or:} \\ \mu_{MGM} &= \frac{A+D}{C}, \end{aligned} \quad (A4)$$

where μ_{MGM} (and σ_{MGM}) denote the expected return (and standard deviation) of the approximate MGM portfolio. The slope of the frontier at any point is given by:

$$\frac{\partial \mu}{\partial \sigma} = \frac{D\sigma}{C\mu - A}, \quad (A5)$$

See Merton (1972), eq.(15). Thus, the slope at the approximate MGM portfolio is:

$$\frac{D\sigma_{MGM}}{C\mu_{MGM} - A} = \frac{D\sigma_{MGM}}{C\left(\frac{A+D}{C}\right) - A} = \sigma_{MGM}. \quad (A6)$$

In other words, for the approximate MGM portfolio the slope is exactly equal to the portfolio's standard deviation. r^* is the point where the tangent at the approximate MGM portfolio crosses the y-axis, see Figure 2. Hence, r^* satisfies:

$$\frac{\mu_{MGM} - r^*}{\sigma_{MGM}} = \sigma_{MGM}, \quad (A7)$$

as these are two expressions for the same slope. This allows us to express r^* as:

$$r^* = \mu_{MGM} - \sigma_{MGM}^2 = \mu_{MGM} - \frac{C\mu_{MGM}^2 - 2A\mu_{MGM} + B}{D}. \quad (A8)$$

Employing (A4), and a little algebraic manipulation yields:

$$r^* = \frac{A-1}{C} = \frac{\underline{1}' \Sigma^{-1} \underline{\mu} - 1}{\underline{1}' \Sigma^{-1} \underline{1}}. \quad \text{QED.}$$