



Asset Allocation Optimization Methodology

Morningstar Methodology Paper
December 12, 2011

©2011 Morningstar, Inc. All rights reserved. The information in this document is the property of Morningstar, Inc.

Reproduction or transcription by any means, in whole or in part, without the prior written consent of Morningstar, Inc., is prohibited.

Contents

Introduction	3
Asset Allocation	3
Limitations of Mean-Variance Optimization	4
Scalability and Frequency Conversion	7
Asset Class Assumptions Modeling	9
Comparison of Models	9
Model #1: Lognormal	10
Model #2: Johnson	14
Model #3: Bootstrapping	16
Optimization and Forecasting	18
Reward—Expected Arithmetic Mean Return	19
Reward—Expected Geometric Return	19
Risk—Standard Deviation	20
Risk—SMDD Standard Deviation	20
Risk—Conditional Value at Risk	21
Risk—Downside Deviation Below a Target Return	22
Risk—Downside Deviation Below the Arithmetic Mean	22
Risk—First Lower Partial Moment Below a Target Return	23
Risk—First Lower Partial Moment Below the Arithmetic Mean	23
Optimization—Mean-Variance Optimization	24
Optimization—MVO with Resampling	25
Optimization—Scenario-Based	26
Forecasting	26
Appendix A: Johnson Distributions	27
Appendix B: Gaussian Copulas	30
Appendix C: Scenario-Based Optimization	31
Smoothed Multivariate Discrete Distributions (SMDD)	31
Expected Utility, Certainty Equivalents, and the Geometric Mean	34
Density Functions	35
Downside Risk, Value at Risk, and Conditional Value at Risk	37
References	40

Introduction

Asset Allocation

Asset allocation is the process of dividing investments among different kinds of asset categories, such as stocks, bonds, real estate, and cash, to achieve a feasible combination of risk and reward that is consistent with an investor's specific situation and goals. When the process involves portfolio optimization, it consists of three general steps. In the first step, the investor specifies asset classes and models forward-looking assumptions for each asset class' return and risk as well as co-movements among the asset classes. When a scenario-based approach is used, returns are simulated based on these forward-looking assumptions. In the second step, an optimization algorithm arrives at percentage allocations to different asset classes, and these allocations are known as the asset mix. In the third step, asset mix return and wealth forecasts are projected over various investment horizons and probabilities to illustrate potential outcomes. For example, an investor can view an estimate of what the portfolio value would be three years from now if its returns were in the bottom 5% of the projected range during this period.

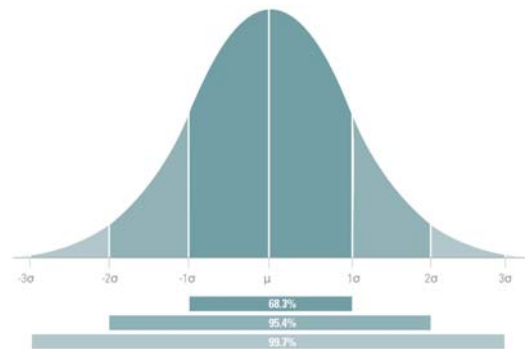
Mean-variance optimization (MVO) has been the standard for creating efficient asset allocation strategies for more than half a century. But MVO is not without its limitations. This document describes these limitations and discusses asset-class modeling and portfolio optimization techniques that overcome some of these shortcomings.

Limitations of Mean-Variance Optimization

First, traditional MVO cannot take into account "fat-tailed" asset class return distributions, which better match real-world historical asset class returns. For example, consider the monthly total returns of the S&P 500 Index back to 1926¹. There are 1,025 months between January 1926 and May 2011. The monthly arithmetic mean and standard deviation of the S&P 500 Index over this time period are 0.943% and 5.528%, respectively. Based on a normal distribution, the return that is three standard deviations away from the mean is -15.64%, calculated as $(0.943\% - 3 \times 5.528\%)$.

¹ Based on Ibbotson Associates' method for extending the S&P 500 TR returns back to January 1926.

Introduction (continued)

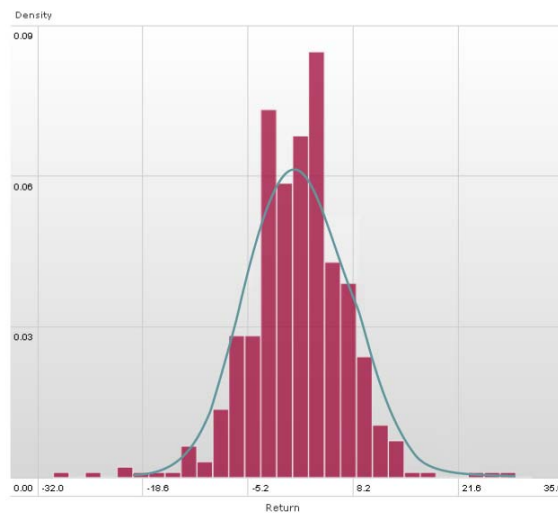


In a normal distribution, 68.27% of the data values are within one standard deviation from the mean, 95.45% within two standard deviations, and 99.73% within three standard deviations. This implies that there is a 0.13% probability that returns would be three standard deviations below the mean, where 0.13% is calculated as $(100\% - 99.73\%)/2$. In other words, the normal distribution estimates that there is a 0.13% probability of returning less than -15.64%, which means that only 1.3 months out of those 1,025 months between January 1926 and May 2011 ought to have returns below -15.64%, where 1.3 months is arrived at by multiplying the 0.13% probability by 1,025 months of return data. However, when examining historical data during the period, there are 10 months where this occurs, which is almost eight times more than the model would predict. Following are the 10 months in question.

Month	Return (%)
Jun 1930	-16.25
Oct 2008	-16.79
Feb 1933	-17.72
Oct 1929	-19.73
Apr 1932	-19.97
Oct 1987	-21.54
May 1932	-21.96
May 1940	-22.89
Mar 1938	-24.87
Sep 1931	-29.73

Introduction (continued)

The normal distribution model also assumes a symmetric bell-shaped curve, and this implies that the model is not well-suited for asset classes with asymmetric return distributions. The histogram below shows the number of historical returns that occurred in the return range of each bar. The curve layered over the histogram graph shows the probability the normal distribution predicts. In the graph below, the left tail of the histogram is longer, and there are actual historical returns that the normal distribution does not predict.



As demonstrated in an article² by Xiong and Idzorek in 2011, skewness (asymmetry) and excess kurtosis (larger than normal tails) in a return distribution can have a significant impact on the optimal allocations in a portfolio selection model where a downside risk measure, such as conditional value at risk (CVaR, defined later in this document), is used as the risk parameter. Intuitively, aside from lower standard deviation, investors should prefer assets with positive skewness and low kurtosis. By ignoring skewness and kurtosis, investors who rely on MVO alone may be creating portfolios that are riskier than they realize.

² Xiong, James X. and Thomas M. Idzorek, "The Impact of Skewness and Fat Tails on the Asset Allocation Decision," Financial Analysts Journal, March/April 2011, pp. 23-35.

Introduction (continued)

Second, the traditional MVO assumes that co-variation of the returns on different asset classes is linear. This means that the relationship between the asset classes is consistent across the entire range of returns. However, the degree of co-variation among equity markets tends to go up during global financial crises. Furthermore, a linear model is clearly an inadequate representation of co-variation when the relationship between two asset classes is based at least in part on optionality such as the relationship between stocks and convertible bonds. Fortunately, nonlinear co-variation can be modeled using a scenario- or simulation-based approach.

Third, the traditional MVO framework is limited by its ability to only optimize asset mixes for one risk metric, standard deviation. As discussed above, using standard deviation as the risk measure ignores skewness and kurtosis in return distributions. Alternative optimization models that incorporate downside risk measures can have a significant impact on optimal asset allocations.

Fourth, traditional MVO is a single-period optimization model that uses the arithmetic expected mean return as the measure of reward. An alternative is to use expected geometric mean return. If returns were constant, geometric mean would equal arithmetic mean. When returns vary, geometric mean is lower than arithmetic mean. Most importantly, while the expected arithmetic mean is the forecasted result for the next one period, the expected geometric mean forecasts the long-term rate of return. Therefore for an investor who is going to regularly rebalance a portfolio to a given asset mix over a long period of time, the expected geometric mean is the relevant measure of reward when selecting the asset mix. (See Poundstone [2005])

Despite its many limitations, the normal distribution has many attractive properties. It is very easy to work with in a mathematical framework, as its formulas are very simple. The normal distribution is very intuitive, as 68.27% of the data values are within one standard deviation on either side of the mean, 95.45% within two standard deviations, and 99.73% within three standard deviations.

Introduction (continued)

According to the Central Limit Theorem, sums of independently and identically distributed random variables with finite variance tend toward a normal distribution, regardless of the underlying distribution. This has an important implication for forecasting long-term returns. A long-term return relative (one plus the return in decimal form; e.g., 1.05 for 5%) is the product of short-term return relatives. Therefore the logarithm of a long-term return relative is the sum of the logarithms of short-term return relatives. If the short-term log-return relatives are independent and identically distributed random variables with finite variance, then the long-term log-return relative can be approximated well by a normal distribution. This means that the long-term return approximately follows a lognormal distribution. In other words, even though log-return relatives measured over small time intervals (e.g. daily) may exhibit skewness and kurtosis, when they are aggregated over a long measurement period (e.g. five years), the aggregated returns are approximately lognormally³ distributed. The exception to this rule is when the log-return relatives follow extremely fat-tailed distributions in which standard deviation is infinite and thus not defined.

In summary, the asset allocation capabilities in Morningstar Direct allow users to choose from a number of return distribution assumptions to model asset class behavior, including traditional bell-curve shaped return distributions (lognormal) as well as fat-tailed and skewed distributions. Users can then use a conventional MVO, a resampled MVO, or a scenario-based optimization to create optimal asset allocation strategies. They can elect to create strategies that will produce the highest single- or multi-period expected return either for a given level of volatility or for a given level of downside risk.

Scalability and Frequency Conversion

Scalability is another property that is important to return distribution modeling. It means that when the frequency of the underlying return data changes—for example, from monthly to yearly or vice versa—only the parameters of the distribution function change, but not its shape. When a distribution is scalable, one can simulate returns in a frequency different from that of the parameters, for example, using mean return and standard deviation derived from monthly data to simulate annual returns. One of the models offered by Morningstar Direct for asset class assumption modeling is not scalable (Johnson) but this is not an issue at present since the tool simulates returns at the same frequency as that of the underlying assumptions, currently set at monthly.

³ A random variable X has a lognormal distribution if its logarithm has a normal distribution. However, when we say that a return has a lognormal distributions, it is understood that we are actually speaking of the return relative.

Introduction (continued)

The tool allows users to display returns and standard deviations at a frequency different from that of the underlying return data, as it is common practice to display these statistics in annual terms. For example, when the display frequency is changed to annual, return simulations can first be performed in monthly terms. When a portfolio is constructed either via the optimization process or entered by the user, the portfolio weights are applied to monthly simulated returns first to form a stream of portfolio returns. Then the monthly portfolio returns from the first 12-consecutive simulations are accumulated to form the equivalent of a calendar-year return. Monthly returns from the second dozen simulations are accumulated to form the second-annual period, and so on. Portfolio statistics such as return and risk measures are derived from this newly formed annual return stream.

Asset Class Assumptions Modeling

Comparison of Models

In the first step of an asset allocation optimization process, the investor specifies asset classes and models forward-looking assumptions for each asset class' return and risk as well as co-movements among asset classes. The general practice is to use an index or a blended index as a proxy to represent each asset class, although it is also possible to incorporate an investment such as a fund as the proxy or use no proxy at all. When a historical data stream such as an index or an investment is used as the proxy for an asset class, it can serve as a starting point for the estimation of forward-looking assumptions.

An important part of the assumption formulation process is for the investor to determine what he or she believes the return patterns of asset classes and the joint behavior among asset classes will be. The more common practice is based on the belief that these return behaviors can be modeled by a parametric return distribution function, in other words, mathematical formulas with a small number of parameters that define the return distribution. The alternative is to directly use historical data without specifying a return distribution model, also known as Bootstrapping.

Following are the return distribution functions available in Morningstar Direct Asset Allocation:

- ▶ Lognormal
- ▶ Johnson

Note that the prefix "log" in the name of a distribution function means that the natural logarithmic form of the return relative, i.e. $\ln(1 + R)$, is the random variable that is assumed to follow a given distribution function. For example, lognormal means that $\ln(1 + R)$ is normally distributed. If return relatives are lognormally distributed, returns cannot fall below negative 100%, making it a more realistic characterization of the behavior of the market returns than does the normal distribution. On the following pages we will describe the bootstrapping method and each of the distribution function options in detail. Following is a table summarizing and comparing the characteristics of each method.

	Lognormal	Johnson	Bootstrapping
Parametric model	Yes	Yes	No
Asymmetry modeling	No	Yes	Yes
Fat tail modeling	No	Yes	Yes
Co-variation modeling	Linear	Linear	Non-linear

Asset Class Assumptions Modeling (continued)

Model #1: Lognormal

The normal distribution, also known as the Gaussian distribution, refers to the bell-shaped (symmetrical) distribution curve commonly identified with mean-variance optimization (MVO). It is characterized by two parameters—mean and standard deviation. Mean is the probability-weighted arithmetic average of all possible returns and is the measure of reward in MVO. Variance is the probability-weighted average of the square of difference between all possible returns and the mean. Standard deviation is the square root of variance and is the measure of risk in MVO.

The prefix "log" means that the natural logarithmic form of the return relative, $\ln(1 + R)$, is normally distributed. The lognormal distribution is asymmetrical, skewing to the right. Because the logarithm of 0 is $-\infty$, the lowest return possible is -100%, which reflects the fact that an unleveraged investment cannot lose more than 100%.

Despite the many limitations discussed in the Introduction section of this document, the lognormal distribution has many attractive properties. It is very easy to work with in a mathematical framework. It is scalable; therefore, mean and standard deviation can be derived from a frequency different from that of the return simulation. Limitations of the model include its inability to model the skewness and kurtosis empirically observed in historical returns. In other words, the lognormal distribution assumes that the skewness and excess kurtosis of $\ln(1 + R)$ are both zero. (These two measures are discussed in detail below, in the Johnson distribution portion of this section of the document.)

Morningstar Direct allows users to model standard deviation of return by using user-defined numbers or by calculating the statistic based on historical data. The tool offers the following methodologies for modeling the expected arithmetic mean return:

- ▶ Historical
- ▶ Building Blocks
- ▶ CAPM
- ▶ Black-Litterman
- ▶ User-Defined

Asset Class Assumptions Modeling (continued)

The Building Blocks and CAPM methods calculate forward-looking arithmetic returns by adding historical risk premium(s) and a current market premium. Morningstar believes that historical risk premiums are preferred over historical average returns since risk premiums have been found to be more consistent and stable over time, lending more confidence in predicting future returns. The main difference between Building Blocks and CAPM methods is in the risk premium calculation. The Building Blocks method calculates risk premium(s) by taking the arithmetic difference between two historical data series, while the CAPM method uses a regression approach. Both methodologies are described in detail in the Ibbotson Stocks, Bonds, Bills and Inflation (SBBI) Yearbook.

The Black-Litterman method incorporates a user's individual views of the market in the expected returns. A set of expected returns are first calculated from a given risk-free rate, equity risk premium, and set of market capitalization weights for each asset class before applying the views. This methodology is also described in detail in the Ibbotson Stocks, Bonds, Bills and Inflation (SBBI) Yearbook.

The arithmetic mean and standard deviation are often estimated from historical returns as follow:

$$[1] \quad M = \frac{\sum_{i=1}^n R_i}{n}$$

$$[2] \quad S = \sqrt{\frac{\sum_{i=1}^n (R_i - M)^2}{n - 1}}$$

Where:

M	=	Arithmetic mean return
R_i	=	Return in period i
n	=	Total number of periods
S	=	Standard deviation of return

Asset Class Assumptions Modeling (continued)

The mean and standard deviation of the natural log of return relative—the two parameters underlying the lognormal distribution—can be derived from the arithmetic mean return (M) and standard deviation of return (S) as follows :

$$[3] \quad \sigma = \sqrt{\ln \left[1 + \left(\frac{S}{1+M} \right)^2 \right]}$$

$$[4] \quad \mu = \ln(1+M) - 0.5 \cdot \sigma^2$$

Where:

σ	=	Standard deviation of log-relative return, $\ln(1+R)$
μ	=	Mean of log-relative return, $\ln(1+R)$

Correlation coefficients are another set of parameters required when modeling asset class returns. The correlation coefficients for the lognormal model can either be calculated over a historical time period or entered by the user. When using historical return data, the correlations of returns can be estimated using the following formulas:

$$[5] \quad C^{A,B} = \frac{\sum_{i=1}^n (R_i^A - M^A) \cdot (R_i^B - M^B)}{\sqrt{\sum_{i=1}^n (R_i^A - M^A)^2} \cdot \sqrt{\sum_{i=1}^n (R_i^B - M^B)^2}}$$

Once values have been set for the expected arithmetic means, standard deviations, and correlations of return, correlations among log-return relatives can be calculated as follows:

$$[6] \quad \rho^{A,B} = \frac{1}{\sigma^A} \cdot \frac{1}{\sigma^B} \cdot \ln \left[1 + \frac{S^A \cdot S^B \cdot C}{(1+M^A) \cdot (1+M^B)} \right]$$

Asset Class Assumptions Modeling (continued)

Where:

$C^{A,B}$	=	Correlation of returns between asset classes A and B
R_i^A	=	Return for asset class A in period i
M^A	=	Expected arithmetic mean return for asset class A
$\rho^{A,B}$	=	Correlation of log-return relatives between asset classes A and B
σ^A	=	Standard deviation of log-relative return for asset class A
S^A	=	Standard deviation of return for asset class A

Return simulation is a step that happens between the steps of setting return assumptions and optimization in the asset allocation process when a scenario-based to optimization is used. Morningstar Direct offers users many options for asset class assumption modeling and optimization. The only combination of user choices that does not require return simulation is when the lognormal distribution is used in assumption modeling in conjunction with conventional MVO (single-period optimization using arithmetic mean as the reward and standard deviation as the risk measures). To simulate asset class returns, log-return relative means, standard deviations, and correlation coefficients are determined using formulas [3], [4], and [6]. Then simulated log-return relatives are converted back into returns:

$$[7] \quad \tilde{R} = \exp(\tilde{r}) - 1$$

Where:

$\tilde{R}$	=	Simulated return
$\tilde{r}$	=	Simulated log-return relative

Asset Class Assumptions Modeling (continued)

Model #2: Johnson

The Johnson distributions are a four-parameter parametric family of return distribution functions that can be used in modeling skewness and kurtosis. Skewness and kurtosis are two important distribution properties that are zero in the normal distribution and take on limited values in the lognormal model (as implied by the mean and standard deviation).

The Johnson distribution's four parameters are mean, standard deviation, skewness, and excess kurtosis. Mean and standard deviation are described in detail in the previous Lognormal section of this document. Skewness and excess kurtosis are measures of asymmetry and peakedness. The normal distribution is a special case of Johnson with skewness and excess kurtosis of zero. The lognormal distribution is also a special case that is generated by setting the skewness and excess kurtosis parameters to the appropriate values. Positive skewness means that the return distribution has a longer tail on the right side than the left side, and negative skewness is the opposite. Excess kurtosis is zero for a normal distribution. A distribution with positive excess kurtosis is called leptokurtic and has fatter tails than a normal distribution. A distribution with negative excess kurtosis is called platykurtic and has thinner tails than a normal distribution. Aside from lower standard deviation, investors should prefer assets with positive skewness and lower excess kurtosis.

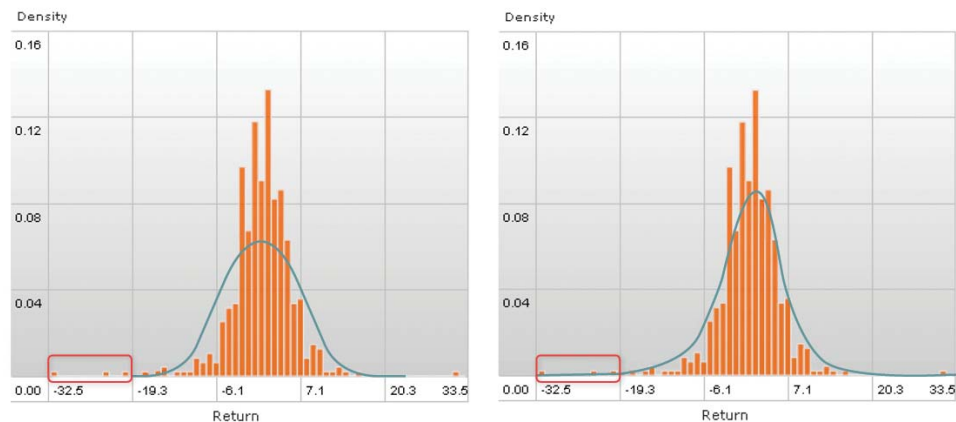
Skewness and excess kurtosis are often estimated from historical return data using the following formulas:

$$[8] \quad \text{Skewness} = \frac{n}{(n-1) \cdot (n-2)} \cdot \sum_{i=1}^n \left(\frac{R_i - M}{S} \right)^3$$

$$[9] \quad EK = \left[\frac{n \cdot (n+1)}{(n-1) \cdot (n-2) \cdot (n-3)} \cdot \sum_{i=1}^n \left(\frac{R_i - M}{S} \right)^4 \right] - \frac{3 \cdot (n-1)^2}{(n-2) \cdot (n-3)}$$

Asset Class Assumptions Modeling (continued)

The difference between modeling with lognormal and Johnson distributions is demonstrated in the following two illustrations. One can clearly observe that the lognormal curve in the first graph does not cover those bars below -20% in the histogram representing actual historical return. In other words, the lognormal model assigns negligible or even no probability to these extreme negative returns which had indeed occurred in the history. The second graph contains the same histogram of historical returns. Although drawn on a slightly different scale from the first graph, one can see that the Johnson curve covers both the left side of the histogram and the peak in the center much better, assuming skewness of -0.8 and excess kurtosis of 3.



Morningstar Direct provides several options for modeling mean and standard deviation, and they are described in the previous section, the Lognormal section of this document. Skewness and excess kurtosis jointly determine the shape of the Johnson distribution curve, and users may specify any combination of these two measures as long as

$$ExcessKurtosis > Skewness^2 - 2.$$

While Johnson distributions are fairly easy to use, they have two limitations. First, they are not scalable. Second, correlation coefficients between pairs of asset classes cannot be modeled directly. Morningstar Direct's solution to this limitation is Gaussian Copulas as presented in Appendix B of this document. Appendix A contains details of Johnson distributions.

Asset Class Assumptions Modeling (continued)

Model #3: Bootstrapping

Instead of modeling asset class behavior with a parametric distribution function, users can choose to use historical data (or any scenario data regardless of source) directly using the bootstrapping technique. It is a sampling process with replacement. The process can be easily understood conceptually with an example. Visualize a very large sheet of paper containing a table with historical returns where columns are asset classes and rows are months in history. Cut this sheet into strips horizontally so that each month is a strip, and place them in a pot and mix them up. Draw a strip from the pot, write down the returns for each asset class, put that strip back in the pot, and repeat the process by drawing another strip. After a large number of repetitions there is a new table of simulated return streams.

The bootstrapping method has several advantages. One is that there are behaviors in some asset classes that cannot be properly modeled by any distribution function. For instance, there may be a reason why a return that is plausible under a parametric distribution in a certain range is in reality unachievable. Or, multiple peaks may exist that cannot be captured by a distribution function model, such as the case illustrated below.

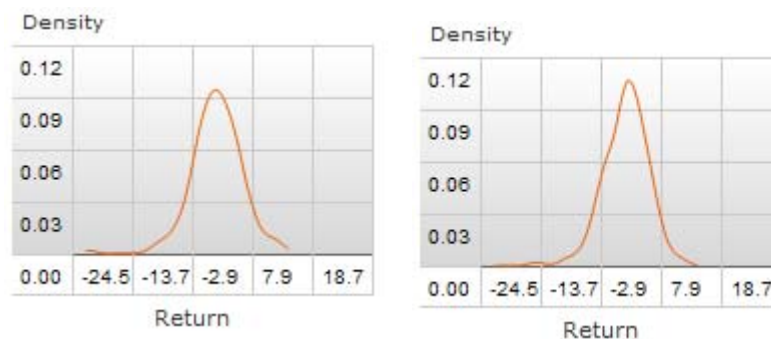


Asset Class Assumptions Modeling (continued)

Another advantage is that any relationship that exists between asset class returns will be preserved, in contrast with parametric distribution function models presented in the rest of this document where co-variation among various transformations of asset class returns are assumed to be linear. The ability to model nonlinear co-variation among asset class returns is particularly interesting if one believes that these relationships changes under certain market conditions, e.g. the tendency for equity markets to move together during a financial crisis. It is also helpful when there is not a clear linear pattern in co-variation of returns.

In Morningstar Direct, users can use historical return data to create scenarios. Furthermore, the user can define several sub-periods and weight their likelihood. In the strips of paper example provided above, it is like having multiples of the same strip of paper so that the strip has more chance of getting drawn in the sampling process. By default, each period is weighted equally. If a user believes, however, that the asset classes will behave like they have historically except that certain periods of the historical data should receive greater weight, then users can overweight specific periods. For example, consider the S&P 500 data from 1980-2009. In this case there are three decades of data. If one thinks that recent behavior is more informative of how an asset class will behave in the future, one can identify that period and overweight it while still including the older data to some degree.

In the example illustrated below, the chart on the left shows the distribution of the S&P 500 data where the entire period of 1980-2009 is given equal weights. In the graph on the right, the first two first two decades are assigned 10% weights while the last decade is emphasized with 80% weight.



Optimization and Forecasting

In the context of asset allocation, optimization is the process of identifying asset mixes that have the highest level of reward for any given level of risk. Morningstar Direct offers two options for measuring reward and six for measuring risk. The choice of reward measure conveys whether an investor desires a single or multi-period view point. The choice of risk depends on whether the investor is more concerned about all return variability or the downside risk of the portfolio and the most appropriate measure to convey the latter. These two types of risk measures can produce very different optimization results, especially for asset class distributions that exhibit skewness. In addition, Morningstar Direct provides users with three types of optimization techniques, and which options are available depends on the combination of risk-return parameters chosen.

Choices of return measures are as follow:

- ▶ Expected arithmetic mean return: single-period viewpoint
- ▶ Expected geometric mean return: multi-period viewpoint

The following risk measures are offered, with the standard deviation being a measure of all variability and the others of downside risk:

- ▶ Standard deviation and/or SMDD standard deviation
- ▶ Conditional value at risk (CVaR)
- ▶ Downside deviation below the arithmetic mean return
- ▶ Downside deviation below a target return
- ▶ First lower partial moment below the arithmetic mean return
- ▶ First lower partial moment below a target return

Optimization techniques are:

- ▶ Mean-variance (MVO)
- ▶ MVO with resampling
- ▶ Scenario-based

This section of the document describes each of these reward and risk measures and optimization techniques in details. Specifically, this section will provide two sets of formulas for each measure. One set is the theoretical, statistical formula expressed as "expected value," $E[.]$, or probability weighted average. However, many people relate better to a discrete version of the formulas where there are a finite number of possible returns that are all equally likely to occur. Because understanding these formulas is the first step in discriminating and selecting among different measurements, it is beneficial for this document to supply the discrete version.

Optimization and Forecasting (continued)

Reward—Expected Arithmetic Mean Return

Morningstar Direct provides users with the choice of expected arithmetic mean return or expected geometric mean return for measuring reward in the risk-reward optimization. The expected arithmetic mean return is the measure of reward in conventional mean-variance optimization (MVO). (Unfortunately, many users of MVO confuse it with the expected geometric mean return and thus misspecify the inputs to MVO and misinterpret its outputs.) The conventional MVO is a single-period model in which the expected arithmetic mean return is a forecast of return over the next period of investment. For a discrete distribution, the expected arithmetic mean return is the simple average given by the formula stated in equation [1] of this document. The theoretical version expresses it as the probability-weighted average of all possible returns, as shown in formula [10].

$$[10] \quad M = E[\tilde{R}]$$

Reward—Expected Geometric Return

The geometric mean return measures how fast wealth accumulates. It is a more familiar statistic than arithmetic mean return because it is a more standard measure of performance. By selecting this reward measure, the user is taking a multi-period viewpoint, in contrast to the conventional MVO which is a single-period model concerned with maximizing expected return just for the next period. In other words, the geometric mean is a more relevant measure of performance for an investor who is investing for a long time and will be rebalancing their portfolio back to the same asset allocation every period. Optimizing on expected geometric mean return rather than expected arithmetic mean return can lead to meaningful differences in the efficient asset mixes, especially at the riskier end of the efficient frontier. The geometric mean is the same as the arithmetic mean when returns are constant. When returns vary it is always below the arithmetic mean.

Equation [11] below shows the theoretical version of the geometric mean formula, while equation [12] is the discrete form:

$$[11] \quad G = \exp\{E[\ln(1 + \tilde{R})]\} - 1$$

$$[12] \quad G = \left[\prod_{i=1}^n (1 + R_i) \right]^{1/n} - 1$$

Optimization and Forecasting (continued)

Risk—Standard Deviation

The standard deviation is a measure of dispersion from the mean and represents the volatility of an investment's return, encompassing both the upside and the downside. The square of the standard deviation is the variance, and standard deviation is the most commonly known risk measure as it is what has been used in conventional mean-variance optimization (MVO) for many decades. One reason is that standard deviation is a very intuitive number. In a normal distribution, 68.27% of the data values lie within one standard deviation away from the mean, 95.45% within two standard deviations, and 99.73% within three standard deviations, etc.

Equation [13] below shows the theoretical version of standard deviation formula, while the discrete version is in equation [2] of this document:⁴

$$[13] \quad S = \sqrt{E\left[\left(\tilde{R} - M\right)^2\right]}$$

Risk—SMDD Standard Deviation

The prefix "SMDD" is the abbreviation for smoothed multivariate discrete distribution. This version of standard deviation differs from the conventional version in that the former is associated with scenario-based optimization while the latter with the MVO. Therefore, by providing these two choices of standard deviation, Morningstar Direct gives users the choice of the optimization technique. For technical details on SMDD, please refer to Appendix C of this document.

⁴ Equation [2] shows the formula for estimating standard deviation from a sample on n observations. For a discrete distribution where there are only n possible outcomes, the denominator needs to be changed from $n-1$ to n . [italicize n in all 3 places.]

Optimization and Forecasting (continued)

Risk—Conditional Value at Risk

Conditional Value at Risk (CVaR) is a measure of downside risk which is a more appropriate measure for investors who are more concerned about downside risk. This is especially important when investing in assets with asymmetrical or fat-tailed return distributions. It is also known as expected shortfall, mean shortfall, expected tail loss, and tail VaR. Conditional Value at Risk (CVaR) is closely related to the more well-known risk measure Value at Risk (VaR). VaR is the maximum loss not exceeded with a given probability (i.e. percentile) over a given period of time. For example, if the 5 percent VaR of a portfolio is 30% per year, it means that over the next year there is a 5% chance that the portfolio will lose 30% or more of its value. The CVaR is the probability weighted average of the possible losses conditional on the loss being equal to or exceeding the VaR. Two asset classes with the same VaR could have very different CVaR figures, depending on how the return distributions of these asset classes are shaped in the left tail. Xiong and Idzorek (2011) and Kaplan (2012, Chapter 26) chose to use the CVaR as the downside risk measure instead of the VaR because the former is subadditive while the latter is not. Being subadditive means that the risk of a combination of investments is at most as large as the sum of the individual risks. This is intuitive as it is consistent with the benefit of diversification.

The formula for CVaR expressed in terms of VaR is as follow:

$$[14] \quad CVaR(p) = -E[\tilde{R} \mid \tilde{R} \leq -VaR(p)]$$

Where:

p = the given probability of loss

While the CVaR, downside deviation and first lower partial moment are all measures of downside risk, the former specifies the probability of loss (e.g. 5 percent) while the latter two specify a target return for which there is a penalty for falling below (e.g. -15% return).

Optimization and Forecasting (continued)

Risk—Downside Deviation Below a Target Return

Downside deviation, as its name indicates, is a measure of downside risk which is a more appropriate measure for investors who are more concerned about the downside risk. This is especially important when investing in assets with asymmetrical or fat-tailed return distributions. Downside deviation is a well-known statistic as it is the denominator of the Sortino Ratio (see Sortino [2001]). The formula for downside deviation is similar to that of the standard deviation, but for returns above a certain target, the deviations (from the target rather than the mean) are set to zero, thus focusing on those returns below the target. Just like the standard deviation and in contrast with the First Lower Partial Moment, by taking the square of deviations in the numerator, the downside deviation measure penalizes larger deviations by a greater degree than smaller ones. In Morningstar Direct, users can specify a particular rate of return to use as the target.

Equation [15] below shows the theoretical version of the formula, while equation [16] is the discrete version:

$$[15] \quad DD(\tau) = \sqrt{E[\max(\tau - \tilde{R}, 0)^2]}$$

$$[16] \quad DD(\tau) = \sqrt{\frac{\sum_{i=1}^n \max(\tau - R_i, 0)^2}{n}}$$

Where:

τ = the target rate of return

Risk—Downside Deviation Below the Arithmetic Mean

An alternative to a user-specified target is to use the return distribution's own arithmetic mean in measuring downside deviation. In this case, the target return differs for each possible asset mix.

Optimization and Forecasting (continued)

Risk—First Lower Partial Moment Below a Target Return

The first lower partial moment (FLPM) is similar to the downside deviation measure in that it measures the deviation below a target rate of return, but the latter penalizes larger deviations to a greater degree. The FLPM is also associated with a well-known measure, the Omega measure as defined by Shadwick and Keating (2002). As demonstrated by Kaplan and Knowles (2004), the Omega measure can be expressed with the FLPM as its denominator⁵.

Equation [17] below shows the theoretical version of the formula, while equation [18] is the discrete version:

$$\begin{aligned} [17] \quad FLPM(\tau) &= E[\max(\tau - \tilde{R}, 0)] \\ [18] \quad FLPM(\tau) &= \frac{\sum_{i=1}^n \max(\tau - R_i, 0)}{n} \end{aligned}$$

Risk—First Lower Partial Moment Below the Arithmetic Mean

An alternative to a user-specified target is to use the return distribution's own arithmetic mean in measuring first lower partial moment. In this case, the target return can differ for each possible asset mix.

⁵ Kaplan and Knowles (2004) shows that the Omega measure in Shadwick and Keating (2002) can be expressed in terms of FLPM as: $\Omega(\tau) = 1 + (M - \tau)/FLPM(\tau)$.

Optimization and Forecasting (continued)

Optimization—Mean-Variance Optimization

Harry Markowitz's Mean-variance optimization (MVO) has been the standard for creating efficient asset allocation strategies for more than half a century. It identifies asset mixes that are efficient in terms of expected arithmetic mean return as the measure of reward and standard deviation as the measure of risk. It is also a single-period optimization as it is based the arithmetic mean as opposed to the geometric mean. In Morningstar Direct, the system defaults to using the MVO technique when the user selects expected arithmetic mean and standard deviation as the risk and reward measures, regardless of which distribution model was used in developing asset class assumptions. If the user changes either of these conditions, a scenario-based optimization is required and automatically run as the MVO becomes an inappropriate technique in incorporating those models. Some of the limitations of the MVO are discussed in the Introduction section of this document. In addition, see the next section on Resampling for another criticism and a potential solution to a limitation of MVO.

When running the MVO optimizer, if asset class assumptions are generated using any of the parametric models, the assumptions for means, standard deviations, and correlations specified are plugged directly into MVO. The formulas below give the relationships between any given asset mix and its expected arithmetic mean return and standard deviation of return.

$$[19] \quad R^P = \sum_{j=1}^N w^j \bullet M^j$$

$$[20] \quad S^P = \sqrt{\sum_{j=1}^N \sum_{k=1}^N w^j \bullet w^k \bullet S^j \bullet S^k \bullet C^{j,k}}$$

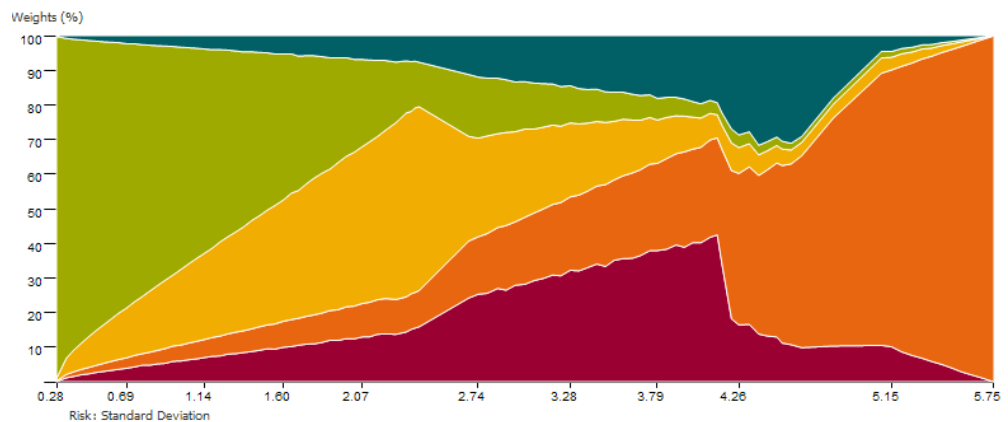
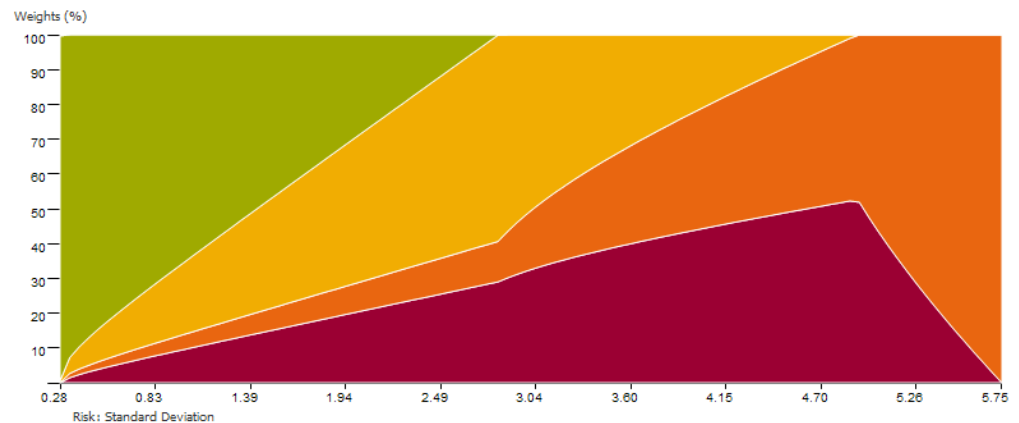
Where:

R^P	=	Expected arithmetic mean return of the asset mix
w^j	=	Asset mix weight in asset class j
M^j	=	Expected arithmetic mean return of asset class j
N	=	Total number of asset classes
S^P	=	Standard deviation of the asset mix
w^k	=	Asset mix weight in asset class k
S^j	=	Standard deviation of asset class j
S^k	=	Standard deviation of asset class k
$C^{j,k}$	=	Correlation coefficient between asset classes j and k

Optimization and Forecasting (continued)

Optimization—MVO with Resampling

MVO with resampling grew out of the work of a number of authors, but is most closely associated with the work of Richard Michaud (2008). Michaud's main criticism of the conventional MVO is that the asset mix produced by the MVO can be very sensitive to small changes in the assumptions for the expected arithmetic mean return, the standard deviation, and the correlation coefficient matrix, as if they were known with 100% certainty—unrealistic in the real world. Such sensitivity can lead to highly concentrated allocation in a small number of asset classes. Michaud's proposed solution is the resampling technique, and the illustrations below show an example comparing conventional MVO on top and resampling at the bottom. The latter has more diversified asset mixes, with each color representing an asset class.



Optimization and Forecasting (continued)

The resampling technique is a combination of Monte Carlo simulation and conventional MVO. In resampling, several thousand hypothetical lognormally distributed return series are generated through Monte Carlo simulation. From each return stream, the arithmetic mean, standard deviation, and correlation coefficients are calculated to produce a hypothetical set of assumptions. Each set of assumptions is used as an input to MVO to form a set of asset mixes. These asset mixes are then grouped into a set of bins by their standard deviation under the original lognormal inputs. All of the asset mixes in a bin are then averaged together to produce one asset mix per standard deviation, forming the resampled frontier. Because this technique is still based on the conventional MVO, it is only available if expected arithmetic mean return and standard deviation are the chosen measures of reward and risk.

Optimization—Scenario-Based

The conventional MVO optimization and resampling techniques are only available when expected arithmetic mean return and standard deviation are the chosen measures of reward and risk. Otherwise, a simulation-based approach is necessary and is automatically applied.

As a step between asset class assumptions modeling and optimization, Morningstar Direct uses a multivariate Monte Carlo simulation of the asset classes to create hypothetical/simulated returns data for the parametric models of return distributions. (This step is described in the Asset Class Assumptions Modeling section of this document. If the Bootstrapping method is used, no further Monte Carlo simulation is needed.) Next, a non-linear optimization algorithm is then used to find the asset mixes at every level of reward that has the lowest level of risk. Details on this optimization technique are described in Appendix C of this document.

Forecasting

This refers to the third step in the asset allocation process where asset mix return and wealth forecasts are projected over various investment horizons and probabilities to illustrate potential outcomes. For example, it estimates what the portfolio value would be three years from now if returns were in the bottom 5% of projected range during this period. This process uses the same multivariate Monte Carlo simulation technique mentioned in the Scenario-Based section above. When forecasting, the model takes into account inflation, cash flows, and rebalancing.

Appendix A: Johnson Distributions

A common method of describing deviations from normality is with the coefficient of skewness (s) and the coefficient of kurtosis⁶ (κ). A normal distribution has $s = 0$ and $\kappa = 3$. Deviations from these values indicate non-normality.

The Johnson distributions as described by Hill, Hill and Holder⁷ (1976, hereafter HHH) provide a convenient way to create random variables with distributions that have given values of s and κ from normally distributed random variables. Thus the return on an asset class can be modeled as:

$$\tilde{R} = g^{-1}(\tilde{z}; \varphi)$$

where $\tilde{z}$ is a random variable with a standard normal distribution, φ is a vector of parameters that describe a Johnson distribution, and $g(\cdot; \varphi)$ is the function that maps values of a random variable with the Johnson distribution described by the parameter vector φ into the standard normal random variable from which it is derived.

Let φ denote the five-element parameter vector ($\gamma, \delta, \xi, \lambda, \text{type}$) where type denotes the type of Johnson distribution which can be three-parameter lognormal (3PLN), unbounded, bounded, or normal.

A random variable $\tilde{X}$ has a distribution function belonging to the Johnson family as defined by HHH if $\tilde{X}$ can be transformed into a standard normal random variable $\tilde{z}$ as follows:

$$\tilde{z} = g(\tilde{X}; \varphi)$$

⁶ Excess kurtosis is often used. Excess kurtosis is simply $\kappa-3$ so that a normal distribution has an excess kurtosis of 0.

⁷ Hill, I.D., R. Hill, and R. L. Holder, "Fitting Johnson Curves by Moments," Applied Statistics, 25:2, 1976.

Appendix A: Johnson Distributions (continued)

where g is defined as follows:

$$g(x; \varphi) = \begin{cases} \lambda \left[\gamma + \delta \ln(\lambda(x - \xi)) \right], & \lambda x > \lambda \xi \\ \gamma + \delta \sinh^{-1} \left(\frac{x - \xi}{\lambda} \right) & (\lambda = \pm 1) \quad (\text{type=3PLN}) \\ \gamma + \delta \ln \left(\frac{x - \xi}{\xi + \lambda - x} \right), & \xi < x < \xi + \lambda \quad (\text{type=unbounded}) \\ \gamma + \delta x & \quad (\text{type=bounded}) \\ & \quad (\text{type=normal}) \end{cases}$$

The choice of type depends on the values of κ and s ⁸. If $s \approx 0$ and $\kappa = 3$, we set type=normal and set

$$\delta = \frac{1}{S}$$

$$\gamma = -\frac{M}{S}$$

where M and S are the mean and standard deviation of $\tilde{X}$ respectively. Otherwise we use the procedure described below.

Let w be the solution to the equation

$$(w - 1)(w + 2)^2 = s^2$$

Let

$$\kappa^* = w^4 + 2w^3 + 3w^2 - 3$$

⁸ Note that for κ and s to be a valid combination of coefficients of skewness and kurtosis, we must have $\kappa > s^2 + 1$.

Appendix A: Johnson Distributions (continued)

If X follows the three-parameter lognormal distribution, $\kappa = \kappa^*$. So if, $\kappa \approx \kappa^*$, we set $\text{type} = \text{3PLN}$. The remaining parameters are set as follows:

$$\delta = \frac{1}{\sqrt{\ln(w)}}$$

$$\gamma = \frac{\delta}{2} \ln \left[\frac{w(w-1)}{S^2} \right]$$

$$\lambda = \text{sign}(s)$$

$$\xi = M - \lambda \exp \left(\frac{\frac{1}{2\delta} - \gamma}{\delta} \right)$$

If $w = \exp(\sigma^2)$ where σ is as defined in equation [3], we have the familiar two-parameter lognormal distribution with $\lambda = 1$, $\delta = 1/\sigma$, $\gamma = -\mu/\sigma$, and $\xi = -1$, μ being as defined in equation [4].

If κ is significantly less than κ^* , $\text{type} = \text{bounded}$. If κ is significantly greater than κ^* , $\text{type} = \text{unbounded}$. In these cases, the parameters are found using the iterative algorithms in HHH.

Appendix B: Gaussian Copulas

Morningstar Direct uses a Gaussian Copulas to model the covariation of asset class returns when using the Johnson distribution. This means that for each trial, we first generate a vector of correlated standard random variables (one for each asset class) and then transform each standard random variable into an asset class return using a given function so we can write

$$\tilde{R}_i = h_i(\tilde{z}_i)$$

where $\tilde{R}_i$ is the return on asset class i , $\tilde{z}_i$ is the standard normal random variable used to generate $\tilde{R}_i$, and $h_i(\cdot)$ the transformation function for asset class i .

When simulating returns with Johnson distributions, Morningstar Direct uses the correlations that the user specifies for the asset classes as the correlations among the $\tilde{z}_i$'s. For $h_i(\cdot)$, we use $g^{-1}(\cdot; \varphi_i)$ where

$$g^{-1}(z; \varphi) = \begin{cases} \lambda \exp\left(\frac{\lambda z - \gamma}{\delta}\right) + \xi & (\text{type}=3\text{PLN}) \\ \lambda \sinh\left(\frac{z - \gamma}{\delta}\right) + \xi & (\text{type}=\text{unbounded}) \\ \frac{\lambda \exp\left(\frac{z - \gamma}{\delta}\right)}{1 + \exp\left(\frac{z - \gamma}{\delta}\right)} + \xi & (\text{type}=\text{bounded}) \\ \frac{z - \gamma}{\delta} & (\text{type}=\text{normal}) \end{cases}$$

Appendix C: Scenario-Based Optimization⁹

Smoothed Multivariate Discrete Distributions (SMDD)

The foundation of the scenario-based optimization framework is the Smoothed Multivariate Discrete Distributions, or SMDD. A multivariate discrete distribution is a flexible way of representing the joint distribution of a random vector. It places no restrictions on either the shapes of the marginal distributions of the separate random variables or the relationships between them.

The disadvantage of discrete distributions is that they lack continuous smooth functional representations. Hence, many basic calculations involve “brute force” methods such as sorting and counting. (For example, the probability distribution function must be depicted using a histogram, and to create a histogram, one must select an arbitrary number of “buckets” and count the number of points that fall within bucket.) This can lead to difficulties in using them in nonlinear optimization problems such as the portfolio selection problem that we discuss below.

In contrast, multivariate normal distributions can be represented by simple smooth continuous functions that have well known properties. The disadvantage is that they misrepresent important empirical characteristics of capital market returns; namely, the regular recurrence of extreme events and the breakdown of linear co-movements between asset classes during those extreme events.

Historical return series or the output of Monte Carlo simulations of “fat-tailed” distributions can be used as a basis to form a more realistic representation of joint return distributions. The usual approach is to simply use a matrix of return data to form a discrete multivariate distribution. Indeed, this is the starting point for our approach.

In order to gain the advantages of having a continuous smooth function to represent the joint distribution of returns, we simply add a multivariate normally distributed disturbance term vector to the discretely distributed vector. The result is a probability model of the vector of returns that has a smooth continuous probability density function yet retains all of the essential properties of the discrete model. This model is the SMDD.

⁹ This appendix is an adaptation of Appendix 26A in Kaplan (2012).

Appendix C: Scenario-Based Optimization (continued)

Formally, let

$R =$ an m by n matrix of asset class returns representing the returns on m asset classes in n scenarios in a discrete model

$w =$ an n -element vector of probability weights for the n events of the discrete model ($w \geq 0$; $\sum_{j=1}^n w_j = 1$)

$W =$ the n by n diagonal matrix form of w (that is, $W_{ii} = w_i$; $W_{ij} = 0$ if $i \neq j$)

$\bar{R}_j =$ the j th column of R (an m -element vector)

$\tilde{J} =$ a discrete uniform random variable that can take on values $1, 2, \dots, n$

$\tilde{\varepsilon} =$ a m -element multivariate normally distributed disturbance vector, with mean vector $\bar{0}$ and variance-covariance matrix Ω

$\tilde{\bar{R}} =$ the random vector of returns on the m asset classes

Our model is simply:

$$\tilde{\bar{R}} = \bar{R}_{\tilde{J}} + \tilde{\varepsilon} \quad (C.1)$$

We construct Ω so that the standard deviation of any asset mix is a fixed percentage greater than what is implied by the discrete distribution alone and the correlations among the asset class returns are unaffected. The fixed percentage is specified by a positive parameter, θ . For example, if $\theta = 0.05$ (i.e., 5%) and the standard deviation of its discrete distribution of a given asset mix were 20%, the standard deviation of its smoothed distribution would be 21%.

The variance-covariance matrix of the discrete multivariate distribution is

$$\Sigma[\bar{R}_j] = RWR' - \bar{\bar{R}}W\bar{\bar{R}}' \quad (C.2)$$

where

$$\bar{\bar{R}} = \sum_{j=1}^n w_j \bar{R}_j \quad (C.3)$$

Appendix C: Scenario-Based Optimization (continued)

Let

$$\phi = \sqrt{(1 + \theta)^2 - 1} \quad (\text{C.4})$$

We set

$$\Omega = \phi^2 \Sigma [\tilde{\mathbf{R}}_j] \quad (\text{C.5})$$

Let $\tilde{\mathbf{x}}$ be an m -element vector representing an asset mix (so that $\sum_{i=1}^m x_i = 1$). The return on the asset mix is the random variable

$$\tilde{\mathbf{R}}_p = \tilde{\mathbf{x}}' \tilde{\mathbf{R}} = \mathbf{x}' \tilde{\mathbf{R}}_j + \tilde{\mathbf{x}}' \tilde{\boldsymbol{\varepsilon}} = \mathbf{R}_{pj} + \tilde{\varepsilon}_p \quad (\text{C.6})$$

The standard deviation of the disturbance term is

$$\omega_p = \phi \sigma [\mathbf{R}_{pj}] \quad (\text{C.7})$$

Note that the relationship between the standard deviation of the return on the portfolio under the SMDD model and standard deviation under the discrete model is

$$\sigma [\tilde{\mathbf{R}}_p] = (1 + \theta) \sigma [\mathbf{R}_{pj}] \quad (\text{C.8})$$

With in the increase in volatility, the SMDD model reduces the coefficient of skewness (s) and the coefficient of excess kurtosis (κ) of the discrete model as follows:

$$s [\tilde{\mathbf{R}}_p] = \frac{s [\mathbf{R}_{pj}]}{(1 + \theta)^3} \quad (\text{C.9})$$

$$\kappa [\tilde{\mathbf{R}}_p] = \frac{\kappa [\mathbf{R}_{pj}]}{(1 + \theta)^4} \quad (\text{C.10})$$

Appendix C: Scenario-Based Optimization (continued)

Expected Utility, Certainty Equivalents, and the Geometric Mean

Levy and Markowitz (1979) developed an approximation for expected utility based on a Taylor series. Suppose that $u(\cdot)$ is a twice differentiable von Neumann-Morgenstern utility function. Let $E[\tilde{R}_p]$ denote the expected value of $\tilde{R}_p$.

The second-order Taylor series expansion of $u(1 + \tilde{R}_p)$ around $1 + E[\tilde{R}_p]$ is

$$\begin{aligned} u(1 + \tilde{R}_p) &\approx u(1 + E[\tilde{R}_p]) + u'(1 + E[\tilde{R}_p])(\tilde{R}_p - E[\tilde{R}_p]) \\ &\quad + \frac{1}{2} u''(1 + E[\tilde{R}_p])(\tilde{R}_p - E[\tilde{R}_p])^2 \end{aligned} \quad (C.11)$$

Hence, expected utility can be approximated as follows:

$$E[u(1 + \tilde{R}_p)] \approx u(1 + E[\tilde{R}_p]) + \frac{1}{2} u''(1 + E[\tilde{R}_p]) \sigma^2[\tilde{R}_p] \quad (C.12)$$

The Levy-Markowitz approximation in equation (C.12) can be generalized to a SMDD model of portfolio return. The model expressed in equation (C.6) can be thought of as a two-stage process. In the first stage, one of the n scenarios, $\tilde{J}$, is drawn at random. In the second stage, the portfolio is drawn from a normal distribution with mean $R_{p\tilde{J}}$ and standard deviation of ω_p . Hence,

$$\begin{aligned} E[u(1 + \tilde{R}_p)] &= \sum_{j=1}^n w_j E[u(1 + \tilde{R}_p) | \tilde{J} = j] \\ &\approx \sum_{j=1}^n w_j \left[u(1 + R_{pj}) + \frac{1}{2} u''(1 + R_{pj}) \omega_p^2 \right] \end{aligned} \quad (C.13)$$

The certainty equivalent return, $CE[\tilde{R}_p]$, is the constant rate of return that yields the same utility as the stochastic return. Using the approximation in equation (C.13), we have

$$CE[\tilde{R}_p] \approx u^{-1} \left(\sum_{j=1}^n w_j \left[u(1 + R_{pj}) + \frac{1}{2} u''(1 + R_{pj}) \omega_p^2 \right] \right) - 1 \quad (C.14)$$

Appendix C: Scenario-Based Optimization (continued)

Poundstone (2005) argues that an investor with an infinite time horizon should construct a portfolio each period using the Kelly criterion which means, in effect, that such investors have logarithmic utility. With logarithmic utility, the certainty equivalent return is the geometric mean which we denote as $GM[\tilde{R}_p]$. Applying equation (C.14), we have

$$GM[\tilde{R}_p] \approx \exp \left(\sum_{j=1}^n w_j \left[\ln(1 + R_{pj}) - \frac{1}{2} \left(\frac{\omega_p}{1 + R_{pj}} \right)^2 \right] \right) - 1 \quad (C.15)$$

Density Functions

Treating the SMDD as the two-stage process that we describe above, it follows that the probability density function (pdf) of portfolio return is

$$f_{\tilde{R}_p}(R_p) = \sum_{j=1}^n w_j f_N(R_p; R_{pj}, \omega) \quad (C.16)$$

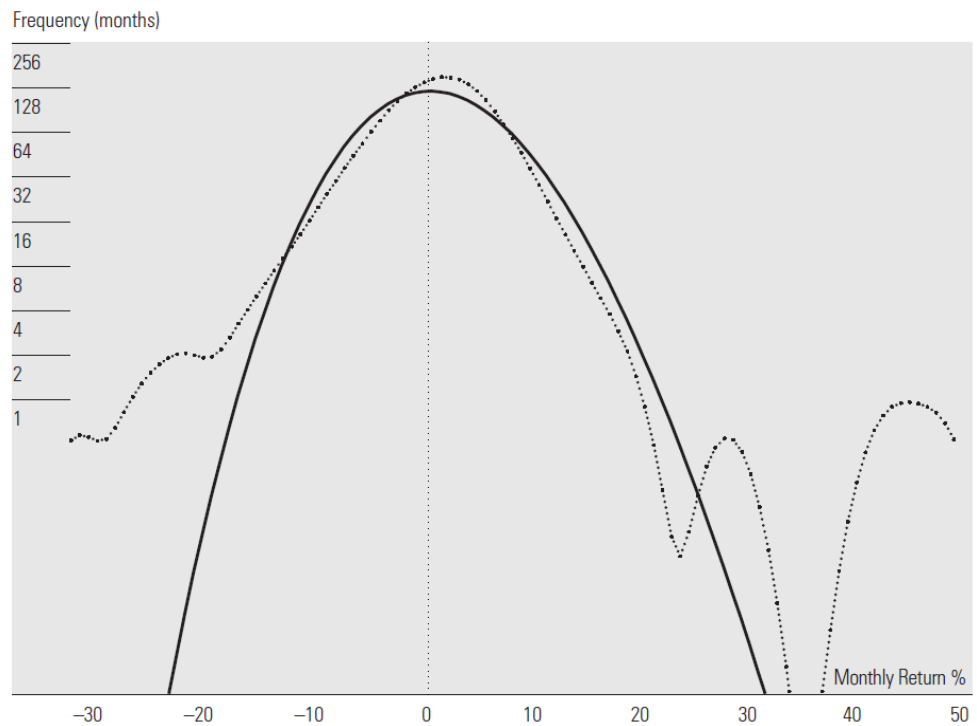
where $f_N(\cdot; \mu, \sigma)$ is the pdf of a normal distribution with mean μ and standard deviation σ :

$$f_N(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left[-\frac{1}{2} \left(\frac{x - \mu}{\sigma} \right)^2 \right] \quad (C.17)$$

Appendix C: Scenario-Based Optimization (continued)

Figure C.1 shows the pdf of the SMDD of the historical real total return on the S&P 500 from January 1926 through August 2009 in contrast to the pdf of a lognormal distribution fitted to the same data. The pdfs are drawn on a logarithmic scale to bring the tails of the distributions into sharp contrast.

Figure C.1: Pdfs of the S&P 500 (1/1926–8/2009) on a Logarithmic Scale



Source: Morningstar (2010).

Appendix C: Scenario-Based Optimization (continued)

Similarly, that the cumulative density function (cdf) of portfolio return is

$$F_{\tilde{R}_p}(R_p) = \sum_{j=1}^n w_j F_N(R_p; R_{pj}, \omega) \quad (C.18)$$

where $F_N(\cdot; \mu, \sigma)$ is the cdf of a normal distribution with mean μ and standard deviation σ .

The inverse cdf, $F_{\tilde{R}_p}^{-1}$ has no closed-form solution and must be estimated numerically using a nonlinear-equation-solving routine.

Downside Risk, Value at Risk, and Conditional Value at Risk

For an investor, risk is not merely the volatility of returns, but the possibility of losing money. This observation has led a number of researchers, including Markowitz (1959), to propose “downside” measures of risk as alternatives to standard deviation. Harlow (1991) formalizes a set of “lower partial moment” downside risk measures. Harlow defines the nth lower partial moment for a given target rate of return, τ , as:

$$LPM_n(\tau) = \int_{-\infty}^{\tau} (\tau - R_p)^n f_{\tilde{R}_p}(R_p) dR_p \quad (C.19)$$

Sharpe (1998) shows that if returns follow a normal distribution with mean μ and standard deviation σ , $LPM_1(\tau)$ is

$$LPM_1(\tau) = \sigma^2 f_N(\tau; \mu, \sigma) + (\tau - \mu) F_N(\tau; \mu, \sigma) \quad (C.20)$$

Applying equation (C.20) to a SMDD, we have

$$LPM_1(\tau) = \sum_{j=1}^n w_j \left[\omega^2 f_N(\tau; R_{pj}, \omega) + (\tau - R_{pj}) F_N(\tau; R_{pj}, \omega) \right] \quad (C.21)$$

Appendix C: Scenario-Based Optimization (continued)

Similarly, it can be shown that if returns follow a normal distribution with mean μ and standard deviation σ , $LPM_2(\tau)$ is

$$LPM_2(\tau) = (\tau - \mu)\sigma^2 f_N(\tau; \mu, \sigma) + \left[(\tau - \mu)^2 + \sigma^2 \right] F_N(\tau; \mu, \sigma) \quad (C.22)$$

Applying equation (C.22) to a SMDD, we have

$$LPM_2(\tau) = \sum_{j=1}^n w_j \left[(\tau - R_{pj}) \omega^2 f_N(\tau; R_{pj}, \omega) + \left[(\tau - R_{pj})^2 + \omega^2 \right] F_N(\tau; R_{pj}, \omega) \right] \quad (C.23)$$

Just as variance is often represented by its square root, standard deviation, LPM2 is often represented by its square root, downside deviation which we write as:

$$DD(\tau) = \sqrt{LPM_2(\tau)} \quad (C.24)$$

Another downside risk measure is to use the portfolio's own arithmetic mean as the target. We define:

$$LPM_1^* = LPM_1(AM[\tilde{R}_p]) \quad (C.25)$$

and

$$DD^* = DD(AM[\tilde{R}_p]) \quad (C.26)$$

where

$$AM[\tilde{R}_p] = \sum_{j=1}^n w_j R_{pj} \quad (C.27)$$

Another risk measure that has become popular is value at risk. VaR for a given probability level p is defined as

$$VaR(p) = -F_p^{-1}(p) \quad (C.28)$$

Appendix C: Scenario-Based Optimization (continued)

Related to VaR is conditional value at risk. CVaR is defined as

$$\text{CVaR}(p) = -E[\tilde{R}_p \mid \tilde{R}_p \leq -\text{VaR}(p)] \quad (\text{C.29})$$

We can write

$$\text{CVaR}(p) = \text{VaR}(p) + \frac{1}{p} \text{LPM}_1(-\text{VaR}(p)) \quad (\text{C.30})$$

Multiplying equation (C.30) through by p yields:

$$p\text{CVaR}(p) = p\text{VaR}(p) + \text{LPM}_1(-\text{VaR}(p)) \quad (\text{C.31})$$

Applying integration-by-parts to equation (C.19), we can show that

$$\text{LPM}_1(\tau) = \int_{-\infty}^{\tau} F_{\tilde{R}_p}(r_p) dr_p \quad (\text{C.32})$$

References

Fishburn, Peter C. "Mean-Risk Analysis with Risk Associated with Below-Target Returns." American Economic Review, March 1977.

Harlow, W. Van. "Asset Allocation in a Downside-Risk Framework." Financial Analysts Journal, September/October 1991.

Hill, I.D., R. Hill, and R. L. Holder, "Fitting Johnson Curves by Moments," Applied Statistics, 25:2, 1976.

Kaplan, Paul D. Frontiers of Modern Asset Allocation, Hoboken, NJ: John Wiley & Sons, Inc. 2012.

Kaplan, Paul D. and James A. Knowles. "Kappa: A Generalized Downside Risk-Adjusted Performance Measure." Journal of Performance Measurement, Spring 2004.

Krokhmal, Pavlo, Jonas Palmquist, and Stanislav Uryasev. "Portfolio Optimization with Conditional Value-at-Risk Objective and Constraints." Journal of Risk, 2002.

Levy, Haim and Harry M. Markowitz. "Approximating Expected Utility by a Function of Mean and Variance." American Economic Review, June 1979.

Markowitz, Harry M. Portfolio Selection: Efficient Diversification of Investments. New York: John Wiley & Sons, 1959.

Markowitz, Harry M. "Portfolio Theory: As I See It." Annual Review of Financial Economics, 2010, <http://www.annualreviews.org/doi/pdf/10.1146/annurev-financial-011110-134602>.

Michaud, Richard O. and Robert O. Michaud. Efficient Asset Management: A Practical Guide to Stock Portfolio Optimization and Asset Allocation, second edition. New York: Oxford University Press. 2008.

Ibbotson Stocks, Bonds, Bills, and Inflation 2010 Yearbook, Morningstar, Inc., 2010.

Nolan, John P., "Multivariate Stable Densities and Distribution Functions: General and Elliptical Case," white paper, Deutsche Bundesbank's 2005 Annual Fall Conference

References (continued)

Poundstone, William. *Fortune's Formula*. Hill and Wang, 2005.

Savage, Sam. *The Flaw of Averages*. John Wiley & Sons, 2009.

Shadwick, William F. and Con Keating. "A Universal Performance Measure." *Journal of Performance Measurement*, Spring 2002.

Sharpe, William F. "Morningstar's Risk-Adjusted Ratings." *Financial Analysts Journal*, July/August 1998.

Sortino, Frank A. "From Alpha to Omega." *Managing Downside Risk in Financial Markets*, Frank A. Sortino and Stephen E. Satchell, eds., Reed Educational and Professional Publishing Ltd., 2001.

Taleb, Nassim N., *Fooled by randomness*, New York: Random House, Inc., 2004.

Xiong, James X. and Thomas M. Idzorek, "The Impact of Skewness and Fat Tails on the Asset Allocation Decision," *Financial Analysts Journal*, March/April 2011, pp. 23-35.