

Needles, Haystacks, and Hidden Factors

Hybridizing fundamental and statistical factor models.

GUY MILLER

Most evaluations of financial risk depend crucially on factor models of asset returns. Factor models explain relations among asset returns by positing common factors, influences that unify the returns of related assets and differentiate returns of assets with contrasting characteristics. Returns to the factors describe the strength of these influences over a particular period.

All assets owe a portion of their returns to the common factors, through exposures that quantify the sensitivity of each asset to each factor return. This is expressed through the returns attribution relation:

$$r = Xf + s \quad (1)$$

where if there are N assets and M common factors, r is an $N \times 1$ vector of asset returns; X is an $N \times M$ matrix of asset exposures to the common factors; f is an $M \times 1$ vector of common factor returns; and s is a vector of specific returns, so called because each is supposed to be specific to an asset and independent of other assets and the common factors.

The attribution relation divides any asset return into a factor component and a diversifiable specific component. In risk management applications, the factor component must be hedged.

The attribution relation (1) is a simple linear model that (if it works) accomplishes an astonishing simplification—empirically determining a very large number of cross-asset covariances ($N[N-1]/2$, or over 100,000 for

GUY MILLER is an executive director and head of Equity Modeling Research at MSCI Barra in Berkeley, CA.
guy.miller@mscibarra.com

500 stocks) by means of assigning a much smaller number of factor exposures and finding the factor return covariances. From the almost incomprehensibly large number of relations that one might obtain in the historical record of returns, the simplification serves to select the few that have some structural or causal basis and so might recur in the future. Factor models aim to filter coincidental relations from the data that will be applied to a risk forecast.

How do factor models select credible relations between assets? The attribution relation itself suggests a classification of methods. The basic ingredients are asset returns on one side and factor exposures and factor returns on the other. If the model prescribes exposures, factor returns can be extracted from a cross-sectional linear regression of the asset returns. Factors associated with prescribed exposures are called *fundamental* factors.

If the factor returns themselves are prescribed, the exposures can be estimated by regressing the time series of asset returns against the history of factor returns. Such factors are sometimes called *macroeconomic*; examples include factors that reflect the returns on a stock index or on commodities prices.

Both these methods identify a source of commonality—a shared characteristic or a defined driver of return—from the outset. The identification generally involves financial insight and judgment.

There is a third approach. Statistical methods can uncover relations from the patterns in records of historical returns, and order them by strength. The strongest relations are assumed not to be coincidental; for these relations factor exposures and factor returns are extracted simultaneously from the historical data. Weaker relations are dismissed as accidental, and are excluded from risk forecasts.

Factors selected this way are called *statistical* factors. The usual method for finding them is an eigenvalue analysis of the historical asset return covariance matrix, also referred to as a principal components analysis or PCA. Zangari [2003] provides a good introduction to equity factor models and PCA methods.

In the community of financial practitioners, there is no shortage of opinions on the relative virtues and shortcomings of the different approaches. They surface everywhere, from books and journal articles to the more troubled variety of cocktail party. The qualitative considerations are straightforward and uncontroversial.

Fundamental factors demand the most extensive input data and enforce the strongest prior assumptions

about sources of returns commonality. Fundamental methods therefore apply data to the factors in the most heavily filtered and concentrated way. This means they are least likely to incorporate coincidences into the factor structure by mistake. If it omits no important factors and if imperfections in its factor exposures are sufficiently small enough, a fundamental model will always surpass macroeconomic and statistical models.

Yet fundamental modelers might neglect or incorrectly specify some factor exposures. Although fundamental modeling minimizes *statistical* errors in making inferences from the data, it can fall prey to inaccuracies in the modeling framework itself and produce *structural* or *systematic* errors.

Statistical factor models represent the other extreme. They make no explicit assumptions about the factor structure. They are therefore viewed as forgoing some precision for the sake of impartiality.

To establish what each method offers and whether different methods might be combined successfully, quantitative comparisons are obviously needed. Connor [1995] provides a classic empirical study based on monthly returns data. He finds that well-constructed fundamental factor models are better at capturing sources of commonality than statistical ones.

Macroeconomic factor models are the least successful of the three; like statistical models, they must estimate a very large number of factor exposures from limited data, and the calculated exposures consequently are often subject to considerable measurement error. They are also likely to omit sources of commonality that lack an obvious returns proxy.

I thus do not discuss them, but rather concentrate on establishing how well statistical models work. My approach complements Connor's direct historical assault on the problem of factor model performance, and extends the examination to higher-frequency data (weekly and daily). It differs from Connor's work in that it emphasizes expectations of success in out-of-sample applications—i.e., in capturing recurring financial relations—instead of success in summarizing historical relations.

I review the nature of the trade-off in choosing between fundamental and statistical models, and examine the possibilities offered by hybridizing the two techniques. Under some conditions, a statistical model may be applied to improve on an imperfectly specified fundamental model. Implementation of a hybrid model is demonstrated in a practical application.

WHEN DO STATISTICAL MODELS WORK?

What conditions are necessary for statistical modeling to find a factor? To answer this question, consider a search for one single factor hidden among the specific returns of one large number of assets.

Analytics

Suppose there are N assets. Of these, n are actually exposed to the factor, and the remaining $N - n$ assets are unexposed. By restricting the non-zero exposures to a sub-population of assets, it is possible to learn when searches might fail for a factor that operates locally within an industry or sector. Can statistical methods succeed when they seek needles in haystacks?

One important determinant of success must be the strength of the factor return relative to the level of specific return noise in which it is embedded. The asset level return is given by the attribution relation (1). Each asset is assumed to have the same specific risk, so that the variance of the specific return s is the same for all assets—call it σ_s^2 . As only the product of the factor exposure and the factor return matters, without loss of generality the mean squared exposure among the n exposed assets may be set equal to 1:

$$\frac{1}{n} \sum_{j=1}^n X_j^2 = 1 \quad (2)$$

Denote the variance of the factor return by σ_f^2 . With the normalization (2), the variance ratio σ_f^2 / σ_s^2 conveniently quantifies the relative strength of the factor and specific returns at the asset level. The higher the variance ratio, the more clearly the factor relationship should appear.

Asset returns are observed over T time periods. During this time the statistical model assumes that the exposures and variance ratio do not change appreciably. To explore potential vulnerability of the method to variation in the factor, suppose that the factor operates only for T periods. For the remaining $T - \tau$ periods of the study history, the factor is “switched off.” Thus the factor persists—intermittently or in an uninterrupted sequence—for a fraction τ / T of the statistical search history.

The issue here is to see whether the factor can be detected and its exposures extracted from historical data. Of course, to be of practical use it is essential to find and

characterize a factor before the end of its lifetime (the time over which its volatility and exposures persist).

Statistical analyses can retrieve useful exposures and factor returns from the data if:

$$\min(T, N) \frac{\tau}{T} \frac{n}{N} \frac{\sigma_f^2}{\sigma_s^2} > 1 \quad (3)$$

Here, *useful* means that the estimation error in the factor returns is comparable to or lower than the average factor return over the study history, and that the exposure error squared and summed over all N assets is lower than n , the sum of the squared exposures. It means that the statistically derived factor returns and exposures resemble the true underlying returns and exposures; they are not dominated by coincidence. To be found and exploited by a statistical factor model, the factor must be relatively strong, long-lived, and distributed across a significant fraction of the asset population.

To understand how the success criterion (3) arises, imagine that the factor returns are known perfectly, so that half of the work of a statistical factor model is already done. A regression of an asset's return history against the factor returns provides an estimate $\hat{X}$ of the asset's exposure to the factor. The estimate is subject to some estimation error. It is related to the true exposure X via:

$$\hat{X} \approx X \pm \frac{1}{\sqrt{\tau}} \frac{\sigma_s}{\sigma_f} \quad (4)$$

Demanding that the sum over assets of the squared errors be less than the sum of squared exposures results in the condition $\tau(n/N)(\sigma_f^2/\sigma_s^2) > 1$. Conversely, if the exposures are perfectly known, the demand that the sum of squared errors in the factor returns be less than the sum of squared factor returns yields the condition $n(\tau/N)(\sigma_f^2/\sigma_s^2) > 1$.

For a statistical factor model to work properly, both conditions must be met. The success criterion combines the conditions to select the more stringent of the two.*

Simulations

Suppose a factor is known to exist. Unless the assets exposed to the factor are known—and, more particularly, unless the factor exposures are known with tolerable accuracy—this knowledge cannot be applied effectively to improve the investment process.

To quantify the success of a statistical factor in capturing a true underlying factor, the exposures are therefore a vital concern. The relationship between the statistical exposures $\hat{X}$ and the underlying exposures X can be described by the quality:

$$Q_X^2 \equiv \frac{(\hat{X}'X)^2}{(\hat{X}'\hat{X})(X'X)} \quad (5)$$

The exposure quality is related to the R^2 of a regression that attempts to fit the statistical exposures with the true exposures. If the two exposure vectors have the same normalization, the statistical exposure vector $\hat{X}$ may be expressed as:

$$\hat{X} = Q_X X + \hat{X}_\perp \quad (6)$$

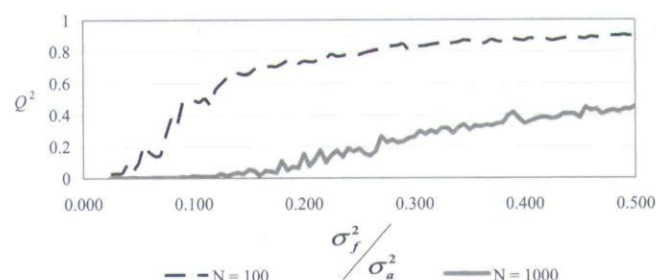
where $\hat{X}_\perp$ is a vector orthogonal (unrelated) to the true exposures X . The square Q_X^2 is preferred to Q_X itself because statistical methods suffer from a sign ambiguity; if the sign of the factor return is reversed along with the sign of the exposures, the factor contribution to asset returns remains the same. Squaring Q_X immunizes it against the sign ambiguity.

Another interpretation of Q_X^2 relates to the success of the statistical exposures in explaining asset returns out-of-sample. Suppose that in a regression of asset returns against the true exposures the factor explains a fraction r^2 of the asset-level returns variance. If the true exposures are replaced with statistical exposures and the regression repeated, the explained fraction will fall approximately to $Q_X^2 r^2$. The quality Q_X^2 describes the fraction of underlying explanatory power that survives in the statistical factor.

The success criterion (3) establishes general guidelines for when statistical factors will reveal and acceptably reproduce the behavior of a hidden factor. Simulations offer a more detailed view of how much is captured under various conditions. A computer runs the returns process in Equation (1) over many observation periods, producing a simulated data history. A statistical factor search is performed on this history, and the success of the statistical factor evaluated.

The benchmark parameters for the simulations are $T = 100$, $n = 20$, and $\sigma_f^2/\sigma_s^2 = 0.1$. The quantity σ_f^2/σ_s^2 is typically a few percent for industry factors (residual to

EXHIBIT 1 Dependence on Strength of Hidden Factor



the broad market component shared by all industry factors) and moderate style factors. Thus the benchmark case assumes a rather strong factor.

In the simulations it is always assumed that the factor and its exposures persist with relatively little change over the observation period T . If each period is a day, the benchmark history extends for just under half a year. If each period is a week, the benchmark history is two years long.

To explore the needle in a haystack problem, the universe of assets for the factor search can be of size $N = 100$ (less dilution) or $N = 1,000$ (substantial dilution). One might think of the number N/n as the number of distinct industries in the market. Then the case $N = 100$ would correspond to a search for a missing sector factor, while $N = 1,000$ corresponds to a search for missing commonality in one of 50 granular industries.

Exhibit 1 shows how the success of the statistical factor depends on the strength of the hidden factor. For purposes of illustration, the scale of strengths extends from $\sigma_f^2/\sigma_s^2 = 0$ to 0.5, but in real cases a value of 0.1 would be very high. At this strength and with $N = 100$, the success criterion (3) is satisfied—the expression on the left-hand side of the inequality is equal to 2. The exhibit shows that at $\sigma_f^2/\sigma_s^2 = 0.1$ the exposure quality is approximately equal to 0.5, confirming that a strong, broad sector distinction is weakly captured. Less spectacular sector distinctions and narrow industry commonalities of any plausible strength (represented by the $N = 1,000$ curve) are invisible to the statistical method.

Exhibit 2 shows how the success of the statistical factor search depends as well on the number of exposed assets. Again, the results of the simulation agree with the success criterion (3), which states that for $N = 100$ and $\sigma_f^2/\sigma_s^2 = 0.1$ there have to be at least 10 exposed assets for the statistical factor to be useful.

EXHIBIT 2

Dependence on Number of Exposed Assets in Universe

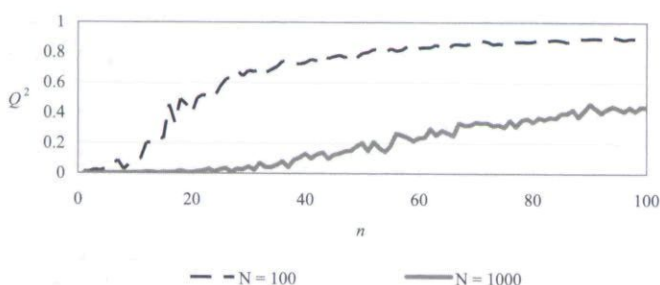


Exhibit 2 shows that the exposure quality reaches approximately 0.5 on the $N = 100$ curve near $n = 20$; there the statistically derived exposures are half signal and half noise. The statistical factor works quite well above $n = 30$. When a strong hidden factor operates across a large fraction of the search universe, it is successfully captured by the statistical factor.

If a factor operates to affect the same assets for a long time, a long history of observations may help in finding it. Exhibit 3 shows how history length influences the quality of the statistically estimated factor exposures. The length is given in terms of the number of observation periods in the history, and extends from 20 to 100. The extreme right-hand side of the plot thus corresponds to a little less than a half year for daily data, two years for weekly data, and eight years for monthly data—a longer history than would probably be used in practice.

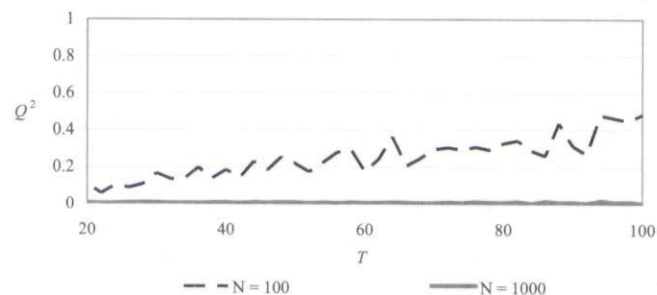
Exhibit 3 shows that even very strong sector distinctions (represented by the $N = 100$ curve) are difficult to capture with several years of monthly data, or with a few years of weekly data. Typical style factors, which extend across the entire universe but are several times weaker, ($\sigma_f^2 / \sigma_s^2 \sim 0.02$), would also be largely invisible. Daily data can reveal broad factors of ordinary strength, but still fail to expose strong factors that last for less than a year.

Implications

The implications of the analytic condition (3) and the simulations are clear. Well-constructed fundamental factor models can find factors that elude statistical models. These include industrial factors that operate in a relatively narrow group of assets within the broader search universe, and technical factors (e.g., momentum) with exposures that change over a few hundred periods or less—a year

EXHIBIT 3

Dependence on History Length



of daily data. Statistical factors can capture broad, strong distinctions that persist for several years, especially if they are based on daily data. Weekly returns data offer much less sensitivity, and monthly data are unlikely to reveal much that a good fundamental model would omit.

IMPLEMENTATION FOR ASSET MANAGEMENT

Why and how do we append statistical factors based on daily returns data to a fundamental factor model?

Performance Attribution and Hybrid Modeling

There are several reasons that models combining statistical and fundamental factors might be desirable. First and most important, they represent complementary approaches to the problem of forecasting risk. Fundamental factors include distinct commonalities that may exist in a relatively narrow cross-section of assets, or over a limited lifetime.

Industry membership factors are examples of the former. It is unlikely that statistical factors would distinguish biotechnology companies from traditional pharmaceutical firms—although investors do. It is also unlikely that statistical factors based on data at daily or lower frequencies would capture a short-term momentum factor that distinguishes recent winners from losers. Yet if a fundamental model should happen to omit a broad and reasonably persistent factor (the most dangerous possibility), statistical factors might reveal it and help quantify its risk.

A second reason is that people who use fundamental factor models use them not only to forecast risk, but also to analyze the performance of their investment

strategies. The factor model decomposes portfolio returns into returns to industry and style factors and to diversifiable asset-level bets. By comparing the return to each category with its contribution to risk, a manager can form an idea of whether different strategy components are performing according to expectations, or whether the strategy performs badly for assets of a certain industry sector or style extreme. As the fundamental factors are particularly convenient for performance analysis, we want to find a way to incorporate one or more statistical factors without modifying the attribution structure.

The solution to this problem is straightforward: Keep the fundamental factor model, and search for statistical factors among its residual returns. First decompose returns according to the attribution Equation (1) using the fundamental exposures and derived factor returns. Take the attributed specific return history and perform a statistical analysis to complete the factor attribution:

$$S_{\text{fundamental}} = \hat{X}\phi + s \quad (7)$$

Here $\hat{X}$ is a matrix of statistically estimated exposures, and ϕ is the vector of returns to the statistical factors. The residual s to the statistical attribution is assumed to include no common factor of importance—although of course the statistical analysis may have overlooked weak factors, transient factors, or factors localized to a particular industry.

The goal of this construction is to introduce the statistical factors without disturbing the structure of the fundamental factor model unnecessarily. Returns attribution is unaffected. The fundamental factor covariances can also be left completely intact (or almost completely). The factor covariance matrix consists of three units: the block of fundamental factor covariances, the block of statistical factor covariances, and the two off-diagonal blocks of covariances between the fundamental and statistical factors.

The statistical factors are often constructed so that their mutual covariances vanish. They nevertheless may covary with the fundamental factors. To create a composite covariance matrix, one takes the statistical factor history that has been extracted with the statistical analysis and calculates correlations between the statistical and fundamental factors. Each off-diagonal correlation is scaled by the forecast standard deviations of the two factors it relates to create the forecast covariance.

Note that this technique does not ensure that the composite factor covariance forecast will be positive-semidefinite. Although our case studies have not yet produced a non-positive example, as a precaution it is prudent to check the composite matrix for negative eigenvalues, and if there are some, to replace them with small but positive eigenvalues.

This straightforward implementation ignores potential difficulties in aggregating the risk deduced from high-frequency statistical factor returns to longer time intervals. For example, if the statistical factor exhibits some mean reversion, its returns on one day would be partially cancelled by returns on the next. The monthly returns variance would be lower than the daily variance times the number of trading days in the month.

In principle, effects of this kind can be accommodated by modeling the serial returns dependencies. In practice, the statistical factors may not be of high enough quality to permit such corrections. Fortunately, if the data frequency is not too high, the serial dependencies should not be strong enough to modify the statistical factor portion of the risk forecast severely. If the statistical factors are providing mild corrections to the fundamental model, their serial relations can be ignored as small corrections to the corrections.

Finally, if the statistical factors are not too strong, they will account for only a few percent of any asset-level returns variance. This means that the specific risk forecasting submodel of the fundamental risk model need not be modified to account for the statistical factors. The effect of the statistical factors is to introduce small correlations between the assets and modify portfolio-level risk, not to make substantial changes in asset-specific risk levels.

Optimization and Portfolio Construction

Even a reasonably successful statistical factor can have significant estimation errors in its exposures. A manager trying to match a portfolio's factor exposures with those of a benchmark may be led by the errors into hedging non-existent risk. Given the noisiness of statistical exposures, should hybrid statistical factors be used together with fundamental factors in portfolio optimization?

The answer seems to be that they should. The reason is that because the errors are random, they are unlikely to produce large phantom exposures unless there are very few assets N_p in the portfolio. The errors tend to be diversified away. Furthermore, although the statistical

factor captures only a fraction of the hidden factor, it alerts the optimization to at least a portion of the hidden risk, which is better than omitting it altogether.

CASE STUDY IN THE JAPANESE EQUITY MARKET

I can demonstrate practical application of statistical factors in a specific market, comparing the performance of MSCI Barra's JPE3 model for risk in the Japanese equity market and an experimental JPE3-based hybrid model.

We investigated the JPE3 factors when we found that JPE3 risk forecasts had failed for several optimized portfolios in the period 1999–2004. The risk forecasts were significantly lower than the realized portfolio risks. The dissonance between the model forecasts and realized risks was traced to the portfolios' specific returns, which failed to diversify as expected. This indicated that a common factor missing from the model had been at work.

The portfolios were long-only portfolios selected from contemporary TOPIX constituents (an index of about 1,400 stocks) and constrained to hold about 100 names. The troublesome portfolios included both tracking portfolios, which attempted only to minimize the departure of their returns from the TOPIX, and tilt portfolios, which tried to outdo the returns of the TOPIX benchmark by overloading on high-value assets with the aim of achieving higher return with reasonable risk.

Although no model is perfect, the forecast biases were unexpected and unwelcome. JPE3 had been released as a commercial product just a few years before, and its performance in a wide array of tests had been exemplary. Where had the hidden factors come from?

Among the problematic portfolios, risk forecast inaccuracies had been weaker or non-existent in earlier periods, which indicated that any factors responsible were transient. Furthermore, the forecast biases disappeared when the long-only constraint was removed.

As the long-only constraint is most influential among low-capitalization stocks where benchmark holdings are small, it is likely that the optimization exposed a hidden factor by trying to neutralize large-cap style exposures with the exposures of a few small-cap positions. The problem would be unlikely to appear at full strength in actual portfolios, because it is unusual for a single portfolio to combine positions of such diverse capitalization and liquidity. As these points became clear, the problem became less alarming, but it remained vexing. (Grinold and Kahn [2000] describe the consequences of long-only policies.)

EXHIBIT 4 Out-of-Sample Statistical Factor Significance

Period	Factor 1	Factor 2	Factor 3
1996/01–1999/12	3.73	3.26	1.92
2000/01–2004/06	2.75	3.42	2.71
1996/01–2004/06	3.25	3.36	2.33

Despite numerous attempts and some partial successes, we could not account fully for the hidden factors. It seemed appropriate to try a remedy with statistical factors. When statistical factors based on monthly returns failed, we tried techniques based on daily returns.

The standard hybrid model includes the three strongest statistical factors, selected from a year of daily returns patterns in the JPE3 estimation universe of roughly 1,400 stocks. Variations using shorter histories of 120 or even 60 trading days failed to improve on that standard model. Other variations that attempted to confine the statistical factors to the financial or technological sectors did not perform as well as the standard hybrid model. Experiments applying separate sets of statistical hybrid factors in separate capitalization ranges also produced no improvement.

The standard hybrid model succeeded in finding real factors, verified in out-of-sample testing. Statistical factor exposures obtained with returns data through a given month were used in regressions of the fundamental model residual returns $s_{\text{fundamental}}$ for the next month. Typical t-statistics for the statistical factors (the ratios of the estimated factor return to the contribution expected from estimation error) were high and statistically very significant.

Exhibit 4 shows root mean-squared t-statistics for all three statistical factors over several historical periods.

The bias statistic quantifies forecast performance by averaging the squared ratio of portfolio return (in this case the active return or difference in return between the managed portfolio and the TOPIX benchmark) to the forecast risk, and taking the square root:

$$\text{bias} = \sqrt{\frac{1}{T} \sum_{t=2}^T \left(\frac{r}{\sigma_r} \right)^2} \quad (8)$$

where r is the portfolio return in question, and σ_r is its standard deviation, as forecast by the risk model.

The bias statistic may be thought of as the ratio of realized to forecast risk over the study period. The ideal value is 1—values significantly higher than 1 indicate that the model risk forecasts are too low; values much lower than

EXHIBIT 5

Optimization Performance Comparisons: July 1999–June 2004

Portfolio	Model	Tracking Error	Bias
Track 75	JPE3	2.84%	1.36
Track 75	Hybrid	2.59%	1.22
Track 200	JPE3	1.24%	1.50
Track 200	Hybrid	1.18%	1.39

1 indicate forecasts that are too high. If portfolio returns are normally distributed, and the expected return is much lower than the uncertainty in return, an unbiased forecast will produce values between $1 - \sqrt{2/T}$ and $1 + \sqrt{2/T}$ with 95% confidence. Non-normal returns may demand a broader range, but this serves as a conservative discriminant of forecast quality.

It is worth stressing that since the statistical factors captured real returns commonality, they did not simply inflate risk forecasts. For example, over a five-year period starting in April 1999, the JPE3 bias in forecasting the active risk of the MSCI JP portfolio was 1.08, an excellent result. The value produced by the hybrid model was 1.05. Hybridization did not compromise this successful risk forecast.

Exhibit 5 compares the performance of portfolios constructed with JPE3 and the standard JPE3 hybrid. Optimized portfolios were updated in monthly rebalances that used each model to minimize the portfolio's active return risk against the benchmark TOPIX portfolio, subject to the constraint that the portfolio include a certain number of assets (75 and 200 assets for the tracking portfolios). The unbiased range of results for this five-year study of monthly returns is $0.81 < \text{bias} < 1.19$.

The hybrid model improved the JPE3 risk forecasts, but did not completely eliminate the forecast biases. It is remarkable that although the improvements might be modest, the hybrid model risk forecasts were better than the JPE3 forecasts and so too was the performance of portfolios optimized with the hybrid model. The value of the daily returns-based statistical factors was not wholly lost through estimation error in their exposures.

CONCLUSIONS

Although statistical searches for returns commonality are not as sensitive as the directed searches of fundamental factor models, when carefully applied they are useful to supplement fundamental models. As observations of many independent (or nearly independent) return periods are necessary to discern a typical common factor, statistical factors based on daily returns offer the best chances for improving the forecasts of a fundamental factor model; they can be expected to fare much better than statistical factors based on weekly or monthly returns.

Our Japanese case study illustrates these ideas by adding statistical factors based on daily returns to the fundamental factors of a particular equity risk model. The hybridized model displays modest improvement both in risk forecasting and in optimized portfolio construction.

ENDNOTES

The author thanks John Taymuree for allowing report of the results of his simulations; Xiaowei Li for sharing her insights from the JPE3 hidden factor study; Eiji Mori and Bradley Thilges for the many JPE3 Aegis backtests they contributed to that work; and Mark Ferrari and Lisa Goldberg for their editorial suggestions.

*A detailed derivation of the criterion, based on a full analysis of the principal components problem, is available at www.barra.com.

REFERENCES

- Connor, G. "The Three Types of Factor Models: A Comparison of Their Explanatory Power." *Financial Analysts Journal*, May/June 1995, pp. 42-46.
- Grinold, R.C., and R.N. Kahn. *Active Portfolio Management*, 2nd edition. New York: McGraw-Hill, 2000, pp. 419-443.
- Zangari, P. "Equity Risk Factor Models." In R. Litterman, ed., *Modern Investment Management*. Hoboken: John Wiley & Sons, 2003, pp. 334-395.

To order reprints of this article, please contact Dewey Palmieri at dpalmieri@ijjournals.com or 212-224-3675.

©Euromoney Institutional Investor PLC. This material must be used for the customer's internal business use only and a maximum of ten (10) hard copy print-outs may be made. No further copying or transmission of this material is allowed without the express permission of Euromoney Institutional Investor PLC. Copyright of *Journal of Portfolio Management* is the property of Euromoney Publications PLC and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.

The most recent two editions of this title are only ever available at <http://www.euromoneyplc.com>