

Measurement of Market Integration and Arbitrage

Zhiwu Chen

Peter J. Knez

University of Wisconsin-Madison

We develop a measurement theory of market integration, based on two notions of “integrated markets.” First, two markets cannot be perfectly integrated in any sense if one can construct two portfolios, one from each market, that have identical payoffs but different prices. In that case, the law of one price is violated across the markets. Second they cannot be integrated in a stronger sense if there are cross-market arbitrage opportunities. Two measures of market integration are developed respectively reflecting these notions. The smaller the measures, the more closely integrated (in the respective senses) the markets. Among other things, they are interpreted as measuring pricing discrepancy between markets.

In financial economics, it is common to assume that markets are integrated, either across market types or across national borders. For instance, to derive their well-known options pricing formula, Black and Scholes (1973) assume that the stock, bond, and options markets are all perfectly integrated and that no cross-market arbitrage opportunities are possible. In inter-

We would like to thank Franklin Allen, Ravi Bansal, Phil Dybvig, Cheng-Gao Feng, William Goetzmann, Lars Hansen, Chi-Fu Huang, Ravi Jagannathan, Cyrus Ramezani, Jose Scheinkman, Ken West, and an anonymous referee for comments on the article and John Cochrane, George Constantinides, Mark Fedenia, Wayne Ferson, Gyu-Taeg Oh, and Steve Ross for conversations with them. We are grateful for the research assistance by Dean Johnson and Walter Prahl. In addition, we thank Walter Rudin and Bill Sethares for their help on the algorithms. This article was presented at the 1994 Western Finance Association meetings in Santa Fe, New Mexico. Any remaining errors are ours alone. Address all correspondence to Zhiwu Chen, Graduate School of Business, University of Wisconsin-Madison, 975 University Avenue, Madison, WI 53706.

national asset pricing, models are derived by assuming either perfectly integrated or completely segmented markets.¹ While these assumptions lead to significant insights, the empirical relevance of such models depends crucially on how closely integrated actual markets are. Thus, it is important to develop a formal measurement theory of market integration that one can apply to study cross-market pricing relations.

The achievement of that goal is, however, hampered by the current lack of a formal definition of "integrated markets." As a survey by Adler and Dumas (1983) shows, there are only some vague and improper notions of market integration in the existing literature. For example, one notion is that integrated markets should fluctuate together, and, hence, that low correlations between markets indicate segmentation. As pointed out by Adler and Dumas, this reasoning is flawed, since even two stocks traded on the same exchange do not necessarily move together. According to another notion, integrated national markets should imply highly correlated consumption rates. Consequently, low correlations between national consumption rates are indicative of segmentation. This sequence of inference is correct except that it relies on the empirical validity of the consumption-based asset pricing theory as well as the existence of complete markets. While we can list other notions of market integration,² a common feature among them is that they rely on some parametric asset pricing model. Therefore, when the underlying pricing model is called into question empirically, so is the corresponding notion of market integration.

To develop a formal measurement theory of market integration that is independent of asset pricing models, we start with the idea that two markets cannot be integrated in any sense if it is possible to construct two portfolios, one from each market, that have identical payoffs but different prices. That is, *closely integrated* markets should assign to similar payoffs prices that are *close*. This idea makes sense because, even though markets may have different price-forming microstructures, the degree of segmentation is ultimately reflected by how differently they assign prices to payoffs. Two markets are said to be *perfectly integrated* if the law of one price holds across them.

This intuitive notion of integrated markets allows one to formulate the measurement problem of market integration as a study of

¹ For example, see Solnik (1983), Stulz (1981), and Wheatley (1988). See Adler and Dumas (1983) for a comprehensive review.

² See, among others, Adler and Dumas (1983), Cho, Eun, and Senbet (1986), Gultekin, Gultekin, and Penati (1989), Jorion and Schwartz (1986), Korajczyk and Viallet (1990), Stehle (1977), and Wheatley, (1988) for studies on international asset pricing and market integration.

the relationship between the pricing structures of two markets. Note that the pricing structure of a market is completely summarized by its implied pricing functionals. Under certain conditions each pricing functional can also be uniquely represented by a stochastic discount factor.³ Then, an equivalent way to define perfect market integration is to require that the two markets have *at least one stochastic discount factor in common*. Consequently, when two markets violate perfect integration, we can use the minimum distance between the respective sets of stochastic discount factors for the two markets to gauge the extent of violation. Indeed, we call this mean-square distance the *weak integration measure*. This measure is zero if and only if the two markets are perfectly integrated. A nonzero value for this measure indicates the degree of segmentation: the lower the measure, the more closely integrated the two markets.⁴ It can also be interpreted as the maximum pricing difference between the two markets and over their common payoffs. Thus, this measure reflects precisely the magnitude of pricing discrepancy between the two markets. To some extent, the weak integration measure also indicates the minimum amount of cross-market frictions (e.g., cross-border tariffs) that is necessary to prevent investors from taking advantage of the pricing discrepancy.⁴

The above notion of closely integrated markets is, however, somewhat weak since, even if two markets are perfectly integrated in this sense, arbitrage opportunities may still be possible between them. That is, there may be portfolios, constructed from both markets, that offer positive payoffs in the future but have negative prices. Intuitively, a more reasonable notion of market integration should require the absence of arbitrage across the two markets and, as a result, the extent that this is violated should indicate the extent that the two markets are not integrated. We refer to this notion as the *strong-form market integration* and the other as the *weak-form market integration*. To measure the extent of segmentation in the strong sense, we define the minimum mean-square distance between the respective sets of nonnegative stochastic discount factors for the two markets as a *strong integration measure* or simply *strong measure*. This measure is zero if and only if arbitrage profits are not possible between

³ A stochastic discount factor is sometimes referred to as a pricing operator, a pricing kernel, or the intertemporal marginal rate of substitution. See, for example, Harrison and Kreps (1979), Hansen and Richard (1987), Hansen and Jagannathan (1991, 1992), and Ross (1976, 1978).

⁴ We thank Phil Dybvig for pointing out this cross-market friction connection. See Black (1974) and Stulz (1981) for discussions on cross-market barriers. For example, in the early use of the law of one price in international trade, it was argued that the possible existence of cross-market price differences in traded physical commodities such as metals is due to the transportation costs needed to move the commodities around. It is evident that the minimum amount of transportation costs must be greater than or equal to the maximum amount of price difference per unit of transaction. Otherwise, arbitrage profits would exist.

the two markets. This measure can be interpreted as the mini-max bound on squared pricing differences in using the respective pricing rules of the two markets to value both traded and nontraded payoffs. In other words, it reflects not only how differently two markets will value their common payoffs but also how differently they will price other potential payoffs such as derivative securities.

Clearly, our measurement theory can be applied to estimate empirically the degree of market integration (both domestic and international) without relying on asset pricing models. There is at least one more area in which this theory can be applied. As noted earlier, most models in finance are based on the integrated markets assumption. Consequently, in empirical tests of such models, the price data used by an econometrician can make a crucial difference. When the price data set used in a test is collected from a group of assets that violate the law of one price or the no-arbitrage condition, any pricing theory, regardless of its merits, is destined to fail the test. It is shown in this article that the weak integration measure provides a pricing model-independent lower bound on such false rejections of pricing models. In other words, to avoid such incorrect rejections, one can first use the integration measures to detect any pricing inconsistency in the data and then subject the model to the data set once it is checked.

In addition to developing and characterizing the market integration measures, we also present estimation procedures. For illustrative purposes, we choose the New York Stock Exchange and the NASDAQ market to check whether they are integrated over the period 1973-1991. Each day, these two exchanges open and close at the same times. Hence, one should expect them to be highly, if not perfectly, integrated. In Section 7, we show the sample values of the weak and the strong integration measures between the two exchanges.

Related work includes the papers by Braun (1990), Knez (1991), Snow (1991), and Chen and Knez (1994). They test whether the Hansen and Jagannathan (1991) mean-variance frontiers for stochastic discount factors intersect between two groups of assets. It will be clear that their tests are not tests on market integration or cross-market arbitrage. For example, two markets may be perfectly integrated, and yet their Hansen-Jagannathan frontiers may still not intersect, since the stochastic discount factors included in each frontier constitute only a subset of the set of all admissible stochastic discount factors for the corresponding asset group. As will be clear, some of our discussion is built on the seminal work by Hansen and Jagannathan (1991, 1992). For instance, we have benefitted from their results on the computation of certain stochastic discount factors.

This article is organized as follows. Section 1 characterizes the set of admissible stochastic discount factors implied by either the law of

one price or the no-arbitrage condition. Section 2 introduces the two market integration measures. In Section 3, we study the various properties of the two measures. Section 4 extends the market integration measure to a multimarket context. Section 5 presents algorithms for estimating the two measures and shows the convergence of the algorithms. Section 6 is devoted to an illustration of the market integration measurement method. We conclude the article in Section 7. The proof of each result is given in the Appendix.

1. Preliminaries: The Absence of Arbitrage

Consider a frictionless economy endowed with a probability space $\{\Omega, \mathcal{F}, Pr\}$, on which the linear space of square-integrable random variables, denoted by $\mathcal{L}^2$, is defined. When equipped with the mean-square inner product $\langle x | y \rangle \equiv E(xy)$ for any $x, y \in \mathcal{L}^2$, $\mathcal{L}^2$ is known to be a Hilbert space. Its associated norm is defined by $\|x\| \equiv \sqrt{E(x^2)}$.

In this economy, the market has N traded assets $\{x_i\}$, indexed by the set N , where each $x_i \in \mathcal{L}^2$ is the total return per dollar invested in the i th asset. For a portfolio $\alpha \in \mathfrak{R}^N$, $x = \sum_{i=1}^N \alpha_i \cdot x_i$ is its total payoff, and $\sum_{i=1}^N \alpha_i$ is its current price. Given this set-up, the set of all obtainable asset payoffs in the economy is

$$\mathcal{M} \equiv \left\{ x \in \mathcal{L}^2 : \exists \alpha \in \mathfrak{R}^N \text{ s.t. } \sum_{i=1}^N \alpha_i \cdot x_i = x \right\}, \quad (1)$$

which is a (mean-square) closed subspace of $\mathcal{L}^2$. Define a pricing correspondence $\pi(\cdot)$ on $\mathcal{M}$ by

$$\pi(x) \equiv \left\{ \sum_{i=1}^N \alpha_i : \alpha_i \in \mathfrak{R} \text{ and } \sum_{i=1}^N \alpha_i \cdot x_i = x \right\} \quad \forall x \in \mathcal{M}.$$

Without further economic restrictions, $\pi(x)$ can be a set for each $x \in \mathcal{M}$, since there can be many portfolios that give the same payoff x .

Definition 1. The law of one price (LOP) holds on the market $(\mathcal{M}, \pi)$ if $\pi(x)$ is a singleton for every marketed payoff x .

Another way to define the LOP is to require that any two portfolios generating the same future payoff must have the same price. From now on, let $X \equiv \{x_i\}$ be the N -dimensional basis of $\mathcal{M}$. The following result is adopted from Chamberlain and Rothschild (1983).

Lemma 1. The LOP holds on $(\mathcal{M}, \pi)$ if and only if there exists a unique marketed payoff $d^* \in \mathcal{M}$ such that

$$E(x_i \cdot d^*) = \pi(x_i) = 1 \quad \forall i \in N. \quad (2)$$

The proof of Lemma 1 relies on the fact that when the LOP holds, $\pi(\cdot)$ is a continuous linear pricing functional. Then, according to the Riesz representation theorem,⁵ there exists a *unique* random payoff $d^* \in \mathcal{M}$ representing $\pi(\cdot)$ on $\mathcal{M}$.

We refer to each $d \in \mathcal{L}^2$ satisfying Equation (2) as an *admissible stochastic discount factor* for the market $(\mathcal{M}, \pi)$, and let D be the set of all admissible stochastic discount factors for $(\mathcal{M}, \pi)$. Since d^* is the unique one in D that is also in $\mathcal{M}$, it is easy to see that

$$D = d^* + \mathcal{M}^\perp$$

where $\mathcal{M}^\perp$ is the orthogonal complement of $\mathcal{M}$ in $\mathcal{L}^2$. In other words, $d \in D$ if and only if there exists some $m \in \mathcal{M}^\perp$ such that $d = d^* + m$. The set D possesses two important properties. First, when the LOP holds, D is closed, convex, and nonempty (due to the continuity of the $\mathcal{L}^2$ norm and the linearity of the inner product). Second, D may become smaller as the number of securities N increases; that is, $D(N) \supseteq D(N+1)$. This is true because, when an asset is added to the market with a given price, each admissible stochastic discount factor will have to satisfy not only the equation system in (2) but an additional equation corresponding to the new asset.

Although the LOP presents the most general requirement for a well-functioning capital market, it does not rule out the violation of some other minimum conditions for any market equilibrium. For instance, even when the LOP holds, some positive payoffs may have negative or zero prices.

Definition 2. No arbitrage opportunity exists on market $(\mathcal{M}, \pi)$ if for all $x \in \mathcal{M}$ such that $x \geq 0$ a.s. and $\|x\| > 0$, $\pi(x) > 0$.

This definition requires that all nonnegative payoffs that are positive with positive probability have positive prices. Ross (1978) and Harrison and Kreps (1979) show that the absence of arbitrage is equivalent to the existence of some continuous, positive, and linear extension of π to all of $\mathcal{L}^2$.

⁵ The Riesz representation theorem states that, if $\pi(\cdot)$ is a continuous linear functional on a Hilbert space $\mathcal{M}$, then there is a *unique* $d \in \mathcal{M}$ such that $\pi(x) = E(xd)$ for every $x \in \mathcal{M}$. Furthermore every $d \in \mathcal{M}$ uniquely defines a continuous linear functional in this way. See Luenberger (1969).

Lemma 2 For any market $(\mathcal{M}, \pi)$, the absence of arbitrage holds if and only if there is some $d \in \mathcal{D}$ such that $d > 0$ a.s.

By Lemmas 1 and 2, the absence of arbitrage implies the LOP, but not vice versa. In other words, the set $\mathcal{D}$ can be nonempty, but it may not contain any admissible stochastic discount factor that is positive a.s.

2. Measurement of Market Integration

This section formally defines the market integration measures. For this purpose, take as given two markets: $(\mathcal{M}_A, \pi_A)$ and $(\mathcal{M}_B, \pi_B)$, where, for instance, $\mathcal{M}_A$ and π_A are respectively the payoff span and the pricing functional of market A. For future reference, let $\mathcal{M}_A^\perp$ and $\mathcal{M}_B^\perp$ respectively denote the orthogonal complements of $\mathcal{M}_A$ and $\mathcal{M}_B$ in $\mathcal{L}^2$.

Assumption 1. The LOP holds separately on markets A and B. Furthermore, there exist $x \in \mathcal{M}_A$ and $y \in \mathcal{M}_B$ such that $\pi_A(x) \neq 0$ and $\pi_B(y) \neq 0$.

The above assumption may not hold for two arbitrary markets. If so, one can further divide the market violating the LOP into smaller groups so that the assets in each group jointly satisfy Assumption 1. In empirical applications, one can view any two groups of assets as belonging to two markets.

Let D_A and D_B be the sets of admissible stochastic discount factors, respectively, for markets A and B. Under Assumption 1, both sets D_A and D_B are nonempty. In particular, by Lemma 1, there is a unique stochastic discount factor for market A, denoted by d_A^* , that is in $\mathcal{M}_A$. Similarly, there is a unique $d_B^* \in D_B$ that is in $\mathcal{M}_B$. However, Assumption 1 does not say that pricing is necessarily consistent between the two markets. Intuitively, if two markets are perfectly integrated, at least they should assign the same price to the same payoff. This observation leads to the following definition.

Definition 3. Markets A and B are said to be *perfectly integrated in the weak sense* if $D_A \cap D_B \neq \emptyset$, where $\emptyset$ denotes the empty set.

By the Riesz representation theorem, each stochastic discount factor in D_A and D_B implies a unique continuous linear pricing functional on $\mathcal{L}^2$. Thus, the markets are perfectly integrated in the weak sense; that is, $D_A \cap D_B \neq \emptyset$, if and only if they share at least one pricing rule. The duality nature of the Riesz theorem allows us to work with either the pricing functionals themselves or, equivalently, the stochastic discount factors. However, the latter offers us implementational convenience.

For instance, in $\mathcal{L}^2$, we can use its norm as a distance metric.

We claim that the definition of weak-form market integration is necessary and sufficient to capture the common notion of integrated markets. To see this point, first suppose that $D_A \cap D_B = \emptyset$. Then, when the two markets are viewed as parts of one combined market, there will not exist any random variable $d \in \mathcal{L}^2$ simultaneously satisfying

$$\begin{aligned} E(x \cdot d) &= \pi_A(x) \quad \forall x \in \mathcal{M}_A \\ E(y \cdot d) &= \pi_B(y) \quad \forall y \in \mathcal{M}_B. \end{aligned}$$

In that case, the LOP will be violated across the two markets, which is inconsistent with any acceptable notion of market integration. Hence, Definition 3 is necessary. On the other hand, it is not necessary to define perfect market integration by requiring $D_A = D_B$. For instance, suppose $\mathcal{M}_A \subset \mathcal{M}_B$, which means $\mathcal{M}_A^\perp \supset \mathcal{M}_B^\perp$. Then, since $D_A = d_A^* + \mathcal{M}_A^\perp$ and $D_B = d_B^* + \mathcal{M}_B^\perp$, we have $D_A \neq D_B$ even if the two markets 'price asset payoffs consistently. Thus, Definition 3 is also sufficient.

Another way to interpret Definition 3 is that two markets are perfectly integrated if and only if the LOP holds across them. To see this, note that, when the two markets are treated as parts of one combined market, the combined market will have its set of admissible stochastic discount factors given by $D = D_A \cap D_B$ and its payoff span by $\mathcal{M} = \mathcal{M}_A + \mathcal{M}_B$. Then, $D \neq \emptyset$ and only if $D_A \cap D_B \neq \emptyset$.

In practice, markets may not be perfectly integrated. It is thus necessary to develop a measure to reflect the degree of integration between two markets. Our first measure is based on the idea that two "closely" integrated markets should assign to a given payoff prices that are "close."

Definition 4. For any two markets A and B satisfying Assumption 1, the *weak integration measure* is a real-valued function $g(\cdot, \cdot)$:

$$g(A, B) \equiv \min_{d_A \in D_A, d_B \in D_B} \|d_A - d_B\|^2. \quad (3)$$

Geometrically, $g(A, B)$ measures the minimum squared distance between the two sets D_A and D_B . Note that the set $D_A - D_B$ is convex, closed, and nonempty, since both D_A and D_B are so by Assumption 1. It is well known that every nonempty, convex, and closed set in a Hilbert space has a unique minimum norm element (Luenberger 1969, p. 69). Hence, so does $D_A - D_B$. This implies that a solution to the problem in Equation (3) must exist.

Theorem 1. Any two markets A and B satisfying Assumption 1 are perfectly integrated in the weak sense if and only if $g(A, B) = 0$.

Perfect weak-form integration is thus made testable by estimating $g(A, B)$ directly. When implemented, such a test is also independent of any parametric asset pricing models. From a theoretical perspective, $g(A, B) = 0$ represents a benchmark cross-market relation.

It is worth noting that in general, $g(A, B) \leq \|d_A^* - d_B^*\|^2$, where d_A^* and d_B^* are the unique minimum-norm stochastic discount factors respectively in D_A and D_B . As the following example shows, even if $d_A^* \neq d_B^*$, one may still have $g(A, B) = 0$ and perfect weak-form integration. This observation is important because it means one cannot draw a definite conclusion regarding integration by comparing d_A^* and d_B^* alone.

Example 1. Consider the case where there are two possible states of nature, s_1 and s_2 . Suppose market A has one asset paying \$1 in s_1 and \$0 in s_2 with a current price of \$0.5, while market B has one asset paying \$0 in s_1 and \$2 in s_2 with a current price of \$0.8. It is easy to verify, that $D_A = \{d_A = (0.5, \lambda) : \lambda \in \mathbb{R}\}$ and $D_B = \{d_B = (\gamma, 0.4) : \gamma \in \mathbb{R}\}$. The only marketed state price vector for market A, d_A^* , is clearly (0.5, 0). Similarly, $d_B^* = (0, 0.4)$. Thus, $\|d_A^* - d_B^*\|^2 = 0.5^2 + 0.4^2 = 0.41$. On the other hand, the state price vector (0.5, 0.4) is in both D_A and D_B , which means $g(A, B) = 0$. Therefore, $g(A, B) < \|d_A^* - d_B^*\|^2$ in this example.

Proposition 1. For two markets A and B satisfying Assumption 1, $g(A, B)$ is nondecreasing in both N_A and N_B , where N_A and N_B are the numbers of basis assets respectively in $\mathcal{M}_A$ and $\mathcal{M}_B$.

Proposition 1 basically says that it is harder for two large markets to be perfectly integrated than for two smaller ones. This makes intuitive sense because, with more assets traded on large markets, it is more demanding to maintain pricing consistency across the many assets. For application purposes, this means that, if perfect weak-form integration is rejected with a given set of assets for each market, it will definitely be rejected when additional assets from either market are added to the test.

The notion of weak-form integration is indeed not as strong as one would like since arbitrage opportunities may still be present across the markets A and B even when $g(A, B) = 0$. Therefore, we need a stronger notion of integration.

Assumption 2. For the two markets A and B, (i) there is no arbitrage opportunity on either market and (ii) there exist payoffs $x \in \mathcal{M}_A$ and $y \in \mathcal{M}_B$ such that $\pi_A(x) > 0$ and $\pi_B(y) > 0$.

Let $D_A^{++} \equiv \{d \in D_A : d > 0 \text{ a.s.}\}$ and $D_B^{++} \equiv \{d \in D_B : d > 0 \text{ a.s.}\}$. Under Assumption 2, both sets D_A^{++} and D_B^{++} are nonempty and convex, but not closed. For technical reasons, we define the absence

of cross-market arbitrage or strong-form integration by using the L^2 -closures of D_A^{++} and D_B^{++} , which are respectively denoted by D_A^+ and D_B^+ .

Definition 5. Two markets A and B satisfying Assumption 2 are *perfectly integrated in the strong sense* if $D_A^+ \cap D_B^+ \neq \emptyset$.

Clearly, there is a positive admissible stochastic discount factor in D_A^{++} that is arbitrarily close to some element in D_B^{++} if and only if $D_A^+ \cap D_B^+ \neq \emptyset$. Intuitively, perfect strong-form integration is equivalent to the absence of arbitrage across the two markets (see Lemma 2). The motivation given for the definition of perfect weak-form integration also applies here for the above definition.

Definition 6. For two markets A and B satisfying Assumption 2, the *strong integration measure* $a(A, B)$ is given by

$$a(A, B) \equiv \min_{d_A \in D_A^+, d_B \in D_B^+} \|d_A - d_B\|^2. \quad (4)$$

Theorem 2. Under Assumption 2, markets A and B are perfectly integrated in the strong sense if and only if $a(A, B) = 0$.

A nonzero value for $a(A, B)$ should therefore reflect the extent that cross-market arbitrage is possible between A and B . Since $D_A^+ \subseteq D_A$ and $D_B^+ \subseteq D_B$, it holds by definition that

$$g(A, B) \leq a(A, B),$$

implying that strong-form integration is indeed more demanding than weak-form integration. Like $g(A, B)$, the measure $a(A, B)$ is nondecreasing in N_A and N_B .

Note that the relation defined by either weak-form or strong-form market integration is not transitive. In other words, suppose that markets A and B are perfectly integrated in the weak sense and so are markets B and C . This does not necessarily imply that markets A and C are perfectly integrated in the weak sense. Thus, in applications, one should be cautious in drawing inferences about whether a sequence of markets is perfectly integrated. In Section 5, we give an example to demonstrate why such a relation is not transitive.

3. Properties of the Market Integration Measures

This section shows that the market integration measures reflect how differently two markets assign prices to asset payoffs. Such an exercise is important because it allows us to see that our measurement theory is not ad hoc and it instead has an intuitive economic foundation.

3.1 The weak integration measure as the mini-max bound on squared pricing errors between markets

To establish a pricing error interpretation for $g(A, B)$, take any $d_A \in D_A$ and define for any $x \in \mathcal{L}^2$

$$\hat{\pi}_{d_A}(x) \equiv E(x \cdot d_A). \quad (5)$$

By the Riesz representation theorem, $\hat{\pi}_{d_A}$ is a continuous linear pricing functional that extends π_A to all of $\mathcal{L}^2$. That is, $\hat{\pi}_{d_A}$ agrees with π_A on $\mathcal{M}_A$. Suppose one uses $\hat{\pi}_{d_A}$ to price an arbitrary payoff marketed on B , $x \in \mathcal{M}_B$. Then, the pricing error induced must satisfy the Cauchy-Schwarz inequality:

$$|\hat{\pi}_{d_A}(x) - \pi_B(x)| = |E(x \cdot d_A) - E(x \cdot d_B)| \leq \|d_A - d_B\| \cdot \|x\| \quad (6)$$

where $d_B \in D_B$. Since the inequality (6) holds for any $d_B \in D_B$ and any $x \in \mathcal{M}_B$, we must have

$$\sup_{x \in \mathcal{M}_B, \|x\|=1} |\hat{\pi}_{d_A}(x) - \pi_B(x)|^2 \leq \min_{d_B \in D_B} \|d_A - d_B\|^2. \quad (7)$$

Thus, when one-uses the stochastic discount factor d_A to price any unit-norm payoff marketed on market B, the maximum squared pricing error is bounded above by the minimum squared distance between d_A and D_B .

Similarly, for any $d_B \in D_B$, define

$$\hat{\pi}_{d_B}(x) \equiv E(x \cdot d_B) \quad \forall x \in \mathcal{L}^2. \quad (8)$$

When $\hat{\pi}_{d_B}$ is used to price any unit-norm payoff in $\mathcal{M}_A$, the maximum squared pricing error is given by the minimum squared distance between d_B and the set D_A :

$$\sup_{x \in \mathcal{M}_A, \|x\|=1} |\hat{\pi}_{d_B}(x) - \pi_A(x)|^2 \leq \min_{d_A \in D_A} \|d_A - d_B\|^2. \quad (9)$$

Proposition 2. Suppose markets A and B satisfy Assumption 1. Then,

$$\begin{aligned} g(A, B) &= \inf_{d_A \in D_A} \sup_{x \in \mathcal{M}_B, \|x\|=1} |\hat{\pi}_{d_A}(x) - \pi_B(x)|^2 \\ &= \inf_{d_B \in D_B} \sup_{x \in \mathcal{M}_A, \|x\|=1} |\hat{\pi}_{d_B}(x) - \pi_A(x)|^2 \end{aligned} \quad (10)$$

that is, the weak integration measure is the mini-max bound on the squared pricing errors in using the pricing rules implied by one market to price any unit-norm payoffs marketed on the other market.

This pricing error interpretation is related to a result in Hansen and Jagannathan (1992). In their study, they propose to test an asset

pricing model by measuring the minimum distance between the candidate stochastic discount factor specified by the model and the set of admissible stochastic discount factors. They find that this minimum distance is actually the maximum pricing error induced by the pricing model on the unit-norm sphere of the marketed payoff span.

To grasp the importance of the cross-market pricing error interpretation, look at, for instance, the options pricing literature. A central assumption behind the Black-Scholes pricing formula or the binomial options pricing models is that one can use the pricing rules implied by the stock and bond markets to price such derivative securities as options contracts. Given our preceding discussion, a necessary condition for this assumption to hold is that the weak integration measure between the primitive and the derivative securities markets is zero, which is an empirically testable hypothesis.

3.2 Other representations of the weak integration measure

In the previous subsection, the weak integration measure was shown to be a mini-max bound on squared pricing errors in using the pricing rules of one market to value every (unit-norm) payoff of another. With the aid of functional analytical tools, we can also express this measure in terms of how differently two markets will price their common payoffs. The main trick is to augment each market with the payoffs, and their associated prices, from the other market so that the augmented payoff span for either market is the same.

Given markets A and B, let $\mathcal{F}$ be the linear span of their common payoffs, that is, $\mathcal{F} \equiv \mathcal{M}_A \cap \mathcal{M}_B$. Without loss of generality, assume that $\mathcal{F}$ has a k -dimensional orthogonal basis f . By the projection theorem of Hilbert space, we can decompose both $\mathcal{M}_A$ and $\mathcal{M}_B$ as follows: $\mathcal{M}_A = \mathcal{F} \oplus \mathcal{F}_A^\perp$ and $\mathcal{M}_B = \mathcal{F} \oplus \mathcal{F}_B^\perp$, where $\oplus$ stands for the direct sum operator and $\mathcal{F}_A^\perp$ and $\mathcal{F}_B^\perp$ are the orthogonal complements of $\mathcal{F}$ in, respectively, $\mathcal{M}_A$ and $\mathcal{M}_B$. Since the dimension of $\mathcal{F}_A^\perp$ is $N_A - k$, let $e_A \equiv \{e_{A,1}, \dots, e_{A,N_A-k}\}$ be the orthogonal basis of $\mathcal{F}_A^\perp$. This means that the random payoffs in the bases f and e_A together form an N_A -dimensional basis of $\mathcal{M}_A$. Similarly, there exists an $(N_B - k)$ -dimensional orthogonal basis, e_B , for $\mathcal{F}_B^\perp$ such that the payoffs in f and e_B form a basis for $\mathcal{M}_B$. As before, we use M to denote the joint payoff span of $\mathcal{M}_A$ and $\mathcal{M}_B$.

Consider adding to market A the $(N_B - k)$ payoffs in e_B with their corresponding market B prices. The augmented payoff span for market A is $\overline{\mathcal{M}}_A \equiv \text{span}[\mathcal{M}_A, e_B]$. Since the payoffs in e_B are only marketed on B, we have $\overline{\mathcal{M}}_A = \mathcal{M}$. Extend the pricing functional π_A to $\overline{\mathcal{M}}_A$ as follows:

$$\bar{\pi}_A(x) = \pi_A(x_1) + \pi_B(x_2) \quad (11)$$

for each $x \in \overline{\mathcal{M}}_A$ such that $x = x_1 + x_2$, where $x_1 \in \mathcal{M}_A$ and $x_2 \in \mathcal{F}_B^\perp$. In other words, $\bar{\pi}_A$ is obtained by assigning market A prices to the market-A portion and market B prices to the market-B portion of each payoff. It is easy to check that, under Assumption 1, $\bar{\pi}_A$ is still continuous and linear. Then, according to the Riesz representation theorem, there exists a unique $\bar{d}_A \in \overline{\mathcal{M}}_A$ that represents $\bar{\pi}_A$ on $\overline{\mathcal{M}}_A$. By construction, the projection of $\bar{d}_A$ onto $\mathcal{M}_A$ must be d_A^* . Define

$$\bar{D}_A \equiv \{\bar{d}_A + \epsilon : \epsilon \in \overline{\mathcal{M}}_A^\perp\} \quad (12)$$

where $\overline{\mathcal{M}}_A^\perp$ is the orthogonal complement of $\overline{\mathcal{M}}_A$ in $\mathcal{L}^2$. By the discussion in Section 1, every random variable in $\bar{D}_A$ is an admissible stochastic discount factor for the augmented market A or $(\overline{\mathcal{M}}_A, \bar{\pi}_A)$. Note that $\bar{D}_A \subseteq D_A$, since $\bar{d}_A \in D_A$ and $\overline{\mathcal{M}}_A^\perp \subseteq \mathcal{M}_A^\perp$.

Similarly, augment market B with the payoffs in e_A and extend its pricing functional to all of the resulting payoff span:

$$\overline{\mathcal{M}}_B \equiv \text{span}[\mathcal{M}_B, e_A].$$

Then, the extended pricing functional is given by

$$\bar{\pi}_B(x) = \pi_B(x_1) + \pi_A(x_2) \quad (13)$$

for any $x \in \overline{\mathcal{M}}_B$ such that $x = x_1 + x_2$, where $x_1 \in \mathcal{M}_B$ and $x_2 \in \mathcal{F}_A^\perp$. The set of admissible stochastic discount factors for the augmented market B is $\bar{D}_B = \bar{d}_B + \overline{\mathcal{M}}_B^\perp$, where $\bar{d}_B$ is in $\overline{\mathcal{M}}_B$ and represents $\bar{\pi}_B$ and $\overline{\mathcal{M}}_B^\perp$ is the orthogonal complement of $\overline{\mathcal{M}}_B$ in $\mathcal{L}^2$.

Let f_A^* and f_B^* be, respectively, the projections of d_A^* and d_B^* onto subspace 3. Then, f_A^* for instance, should price the payoffs in 3 consistently with the pricing functionals π_A and $\bar{\pi}_A$.

Lemma 3. Suppose markets A and B satisfy Assumption 1, and define

$$\bar{g}(A, B) \equiv \min_{d_A \in \bar{D}_A, d_B \in \bar{D}_B} \|d_A - d_B\|^2. \quad (14)$$

Then, $\bar{g}(A, B) = \|f_A^* - f_B^*\|^2$.

Lemma 4. For the markets A and B satisfying Assumption 1, it is true that

- (i) $D_A \cap D_B = \bar{D}_A \cap \bar{D}_B$;
- (ii) $g(A, B) = 0$ if and only if $\bar{D}_A = \bar{D}_B$.

Theorem 3. For any two markets A and B satisfying Assumption 1;

$$g(A, B) = \bar{g}(A, B) = \|f_A^* - f_B^*\|^2. \quad (15)$$

To study market integration, one can thus equivalently focus attention on the common payoff span $\mathcal{F}$. In particular, the weak integration measure is given by the squared distance between the projected stochastic discount factors of the two markets onto $\mathcal{F}$.

Proposition 3. *For markets A and B satisfying Assumption 1,*

$$g(A, B) = \sup_{x \in \mathcal{F}, \|x\|=1} |\pi_A(x) - \pi_B(x)|^2; \quad (16)$$

that is, the weak integration measure equals the maximum squared difference between the respective prices assigned by markets A and B to any unit-norm common payoff in $\mathcal{F}$.

The above proposition essentially completes our formalization of the notion that closely integrated markets should assign to their common payoffs prices that are close: when the markets are closely integrated, $g(A, B)$ will be close to zero; conversely, when $g(A, B)$ is close to zero, the markets are then closely integrated. Thus, two markets are perfectly integrated in the weak sense if and only if, for each common payoff in $\mathcal{F}$, they give the same price. This proposition also suggests that the weak integration measure gives the minimum amount of cross-market frictions (e.g., cross-border security transaction taxes) per unit norm of common payoff that is necessary to prevent investors from taking advantage of the pricing inconsistencies (see footnote 4).

Corollary 1. *Suppose markets A and B satisfy Assumption 1. Then, if there is no payoff simultaneously marketed on markets A and B (i.e., $\mathcal{M}_A \cap \mathcal{M}_B = \{0\}$), perfect weak-form integration holds trivially.*

By this corollary, when two markets do not share any marketed payoffs, there is no benchmark payoff with respect to which pricing consistency can be examined across markets. Consequently, perfect weak-form integration holds trivially. However, this is not to say that perfect strong-form integration also holds. Furthermore, it is important to remember that, even if no identical security is cross-listed and traded on both markets, it does not necessarily hold that $\mathcal{M}_A \cap \mathcal{M}_B = \{0\}$, since it may still be possible to mimic some payoffs of one market by forming portfolios of the other market's assets. This is particularly true when there are only finitely many probable states of nature. For instance, suppose that there are 200 states of nature with nonzero probability and that markets A and B each have 200 linearly independent securities traded. Then, even if none of the 200 securities on A is the same as any of the 200 securities on B, every marketed payoff on A must also be marketed on B, that is, $\mathcal{M}_A = \mathcal{M}_B$.

Corollary 2. For the markets A and B under Assumption 1,

$$g(A, B) = \{\pi_A(f) - \pi_B(f)\}' E(ff')^{-1} \{\pi_A(f) - \pi_B(f)\} \quad (17)$$

where the basis of $\mathcal{F}$, f , is treated as a column vector; $E(ff')$ is the second moment matrix of f ; and $\pi_A(f)$ and $\pi_B(f)$ are the vectors of prices assigned to the k basis payoffs in f , respectively, by markets A and B .

In applying the measurement theory, one can thus estimate the weak integration measure by using solely the basis payoffs in f and their market prices respectively assigned by the two markets. To find the common payoffs in f , one may have to appeal to certain factor analytic procedures such as the interbattery factor analysis method used by Chu, Eun, and Senbet (1986). However, depending on one's own preferences, such an estimation method may not be the most efficient one. In Section 7, we provide a more direct estimation procedure for $g(A, B)$ that does not require any information concerning $\mathcal{F}$.

3.3 The strong integration measure as the mini-max bound on squared pricing errors over conceivable payoffs

As before, define for any nonnegative stochastic discount factor $d_A^+ \in D_A^+$,

$$\pi_A^+(x) \equiv E(x \cdot d_A^+) \quad (18)$$

By the Riesz representation theorem, π_A^+ is a continuous, linear, and nonnegative extension of π_A to all of $\mathcal{L}^2$. For any $d_B^+ \in D_B^+$, define similarly a nonnegative extension of π_B to all of $\mathcal{L}^2$, π_B^+ .

Proposition 4. For any markets A and B satisfying Assumption 2, we have

$$a(A, B) = \inf_{d_A^+ \in D_A^+, d_B^+ \in D_B^+} \sup_{x \in \mathcal{L}^2, \|x\|=1} |\pi_A^+(x) - \pi_B^+(x)|^2; \quad (19)$$

that is, the strong integration measure is the mini-max bound on the squared pricing differences in respectively using the nonnegative stochastic discount factors of the two markets to price every unit-norm payoff of $\mathcal{L}^2$, marketed or not.

The strong measure, therefore, reflects how differently two markets will price any conceivable unit-norm payoffs. This is in contrast with the result in Proposition 2, where the weak integration measure is shown to reflect only how well the markets will price each other's marketed payoffs. In other words, the weak measure does not capture how differently the implied pricing rules of the two markets will value payoffs that are marketed neither on A nor on B (e.g., payoffs on derivative securities), while the strong measure does. This observation

further proves that the latter measure represents a stronger notion of market integration.

Mathematically, we can relate the two measures as follows. For any $d_A^+ \in D_A^+$, its projection onto $\mathcal{F}$ is f_A^* by construction: $d_A^+ = f_A^* + (d_A^+ - f_A^*)$, where $(d_A^+ - f_A^*)$ is orthogonal to $\mathcal{F}$ by the projection theorem. Similarly, for any $d_B^+ \in D_B^+$, its projection onto $\mathcal{F}$ is f_B^* and $(d_B^+ - f_B^*)$ is orthogonal to $\mathcal{F}$. Then, using Theorem 3, we have

$$\begin{aligned} a(A, B) &= \inf_{d_A^+ \in D_A^+, d_B^+ \in D_B^+} \|f_A^* - f_B^* + (d_A^+ - f_A^*) + (d_B^+ - f_B^*)\|^2 \\ &= g(A, B) + \inf_{d_A^+ \in D_A^+, d_B^+ \in D_B^+} \|(d_A^+ - f_A^*) + (d_B^+ - f_B^*)\|^2. \end{aligned}$$

4. Extensions to the Multimarket Context

Suppose that a collection of markets are pair-wise perfectly integrated in the weak sense. Then, does this mean that they are also integrated collectively? To answer this question, we need to generalize the pair-wise market integration notion to the case with multiple markets.

Assumption 3. There exist K markets, $(M_1, \pi_1), \dots, (M_K, \pi_K)$, where M_k and π_k are respectively the payoff span and the pricing functional of the k th market, such that (i) the LOP holds on each of the K markets and (ii) for each $k = 1, \dots, K$, there exists some $x \in M_k$ such that $\pi_k(x) \neq 0$.

Let D_k be the set of admissible stochastic discount factors for the k th market. Assumption 3 ensures that D_k is nonempty, convex, and closed for each k .

Definition 7. The K markets under Assumption 3 exhibit *perfect weak-form multimarket integration* if $D_1 \cap D_2 \cap \dots \cap D_K \neq \emptyset$.

The definition of weak-form multimarket integration simply states that there should be a stochastic discount factor that is admissible to all markets in the collection. When this is true, the LOP will, according to Lemma 1, hold across all the markets. Clearly, a necessary condition for $D_1 \cap D_2 \cap \dots \cap D_K \neq \emptyset$ is that $D_k \cap D_l \neq \emptyset$ for any $k, l = 1, \dots, K$. But the latter is not sufficient for the former, as the following example illustrates. Therefore, in applications, even if a collection of markets is pair-wise perfectly integrated, one cannot immediately conclude that they are collectively integrated.

Example 2. Consider three markets in a two-state economy. Figure 1 depicts the geometry of the sets D_1 , D_2 , and D_3 . The payoff spans of the three markets are straight lines passing through the origin, not shown in Figure 1. Each set D_k should be graphically perpendicular

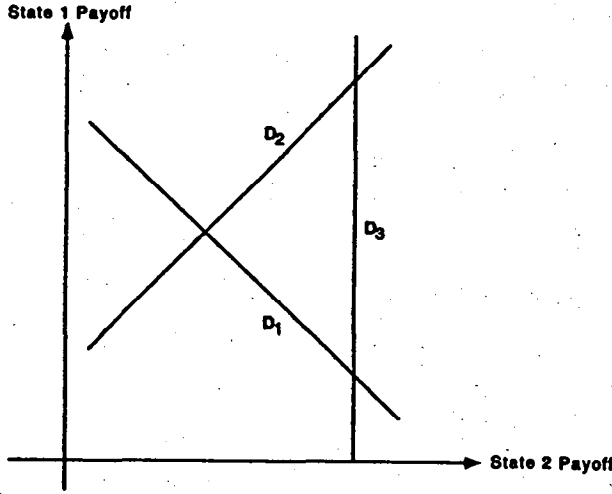


Figure 1

The case with two states and three markets

The line marked D_i , for $i = 1, 2, 3$, corresponds to the set of stochastic discount factors for market i . Here, the three markets are pairwise integrated, but not integrated as a whole; that is, $D_i \cap D_j \neq \emptyset$ for $i, j = 1, 2, 3$, but $D_1 \cap D_2 \cap D_3 = \emptyset$.

to the straight line representing $\mathcal{M}_k$. One can see from the figure that $D_1 \cap D_2 \neq \emptyset$, $D_2 \cap D_3 \neq \emptyset$, and $D_1 \cap D_3 \neq \emptyset$, but $D_1 \cap D_2 \cap D_3 = \emptyset$. In other words, the three markets are pair-wise perfectly integrated in the weak sense but not collectively so.

We generalize the pair-wise weak integration measure below:

$$g(\mathcal{M}_1, \dots, \mathcal{M}_K) \equiv \min_{d_1 \in D_1, \dots, d_K \in D_K} \sum_{k=2}^K \|d_1 - d_k\|^2. \quad (20)$$

As for the two-market case, a solution to Equation (20) is ensured under Assumption 3. Furthermore, weak-form multimarket integration obtains on the collection of K markets if and only if $g(\mathcal{M}_1, \dots, \mathcal{M}_K) = 0$. The proof of this statement is a straightforward generalization of that for Theorem 1. The following is the mini-max squared pricing error interpretation of the multimarket integration measure:

Proposition 5. For the K markets satisfying Assumption 3,

$$g(\mathcal{M}_1, \dots, \mathcal{M}_K) = \inf_{d_1 \in D_1} \left\{ \sum_{k=2}^K \sup_{x \in \mathcal{M}_k, \|x\|=1} |E(x \cdot d_1) - \pi_k(x)|^2 \right\}. \quad (21)$$

In words, when one uses the pricing rules of market 1 to price the payoffs of each of the other $(K - 1)$ markets, there may be pricing errors. Then, the lowest bound on the sum of the maximum such (squared) pricing errors is equal to the multimarket integration measure.

We can similarly extend Proposition 6 and other results to this multimarket context. However, we choose to omit the details here.

5. Applications

The measurement theory developed in the preceding sections can be applied in at least two areas. The first area concerns empirical tests of cross-market pricing relations. In particular, it allows one to measure market integration without relying on any parametric asset pricing models. The other area is related to tests of asset pricing models. Our framework provides a tool with which one can check the internal pricing consistency of an asset data set before subjecting any asset pricing model to it. In this section, we first discuss this application and then present algorithms for the estimation of the market integration measures.

5.1 Measurement of internal pricing inconsistency and tests of asset pricing models

In asset pricing models, it is common to assume, implicitly or explicitly, that all assets are traded in an integrated market where the LOP holds. Therefore, each test of any asset pricing model should be conducted on a set of assets that satisfy the LOP because, if the set of assets violates the LOP, any asset pricing model, regardless of its merits, will be rejected. We show in this subsection that, when an asset pricing model is rejected on a set of assets; the weak integration measure will reflect the extent that the rejection is due to the pricing inconsistency inherent in the data set.

Consider an asset pricing model whose candidate stochastic discount factor is $y \in \mathcal{L}^2$. As suggested by Hansen and Richard (1987), an asset pricing model can be indexed by the candidate stochastic discount factor it specifies. Thus, whether the pricing model is supported by a set of assets depends on whether y is an admissible stochastic discount factor. Suppose an econometrician collects a set of assets, on which the model is to be tested, from two (sub-) markets A and B (say, the NYSE and the NASDAQ): $(\mathcal{M}_A, \pi_A)$ and $(\mathcal{M}_B, \pi_B)$. Again, let M , p , and D be, respectively, the payoff span, pricing functional, and set of admissible stochastic discount factors of the combined 'market.

Define the linear pricing functional on M implied by y , as

$$\pi_y(x) \equiv E(x \cdot y) \quad \forall x \in \mathcal{M}.$$

When y is used to price the unit-norm payoffs in M , its maximum squared pricing error is

$$\delta(y) \equiv \sup_{x \in \mathcal{M}, \|x\|=1} |\pi(x) - \pi_y(x)|^2. \quad (22)$$

Hansen and Jagannathan (1992) call $\delta(y)$ a measure of *model specification* error. Under the assumption that markets A and B are perfectly integrated, they show

$$\delta(y) = \min_{d \in D} \|y - d\|^2. \quad (23)$$

That is, $\delta(y)$ the minimum squared distance between y and the set D . If the pricing model is empirically supported by the combined market, then $\delta(y) = 0$. Otherwise, the larger $\delta(y)$, the more sample information against the model. For more discussion on using $\delta(y)$ to measure the relative performance of a pricing model, the reader is referred to Hansen and Jagannathan (1992).

However, if markets A and B are not perfectly integrated in the weak sense, Equation (23) will not hold and $S(y)$ will not reflect the true pricing model specification error. To see this point, let's still adopt Assumption 1 but assume that A and B are not integrated. That is, D_A and D_B are nonempty but $D = D_A \cap D_B = \emptyset$. In this case, the right side of Equation (23) is not well defined. Let $\delta_A(y)$ be the maximum squared pricing error induced by y over market A; that is,

$$\delta_A(y) \equiv \sup_{x \in \mathcal{M}_A, \|x\|=1} |\pi_A(x) - \pi_y(x)|^2.$$

By Equation (23), $\delta_A(y) = \min_{d_A \in D_A} \|y - d_A\|^2$. Similarly, if we let $\delta_B(y)$ denote the maximum squared pricing error induced by y over market B, it holds that $\delta_B(y) = \min_{d_B \in D_B} \|y - d_B\|^2$. Together, in using y , the maximum squared pricing error over all unit-norm payoffs in $\mathcal{M}$ has to satisfy

$$\delta(y) \geq \max\{\delta_A(y), \delta_B(y)\}. \quad (24)$$

Since $D_A \cap D_B = \emptyset$ by assumption, y can be in at most one of the two sets, D_A and D_B . Suppose, for instance, $y \in D_A$; that is, the pricing model is empirically valid for market A. Then, $\delta_A(y) = 0$ but $\delta_B(y) > 0$. By Equation (24), we have $\delta(y) > 0$. In other words, even when the model performs well on one market, it will still be rejected on the combined market. Evidently, such a rejection is purely due

to the pricing inconsistency across the two markets, not due to the model itself.

Proposition 6. Suppose that y is a candidate stochastic discount factor specified by an asset pricing model and that markets A and B satisfy Assumption 1. Then, the maximum squared pricing error for the combined market is bounded below:

$$\delta(y) \geq g(A, B) + \max\{\delta_A(y) + \mu_A(y), \delta_B(y) + \mu_B(y)\} \quad (25)$$

where

$$\mu_A(y) \equiv \min_{d_A \in D_A, d_B \in D_B} 2E\{(y - d_A)(d_A - d_B)\} \quad (26)$$

$$\mu_B(y) \equiv \min_{d_A \in D_A, d_B \in D_B} 2E\{(y - d_B)(d_A - d_B)\}. \quad (27)$$

Note that $\delta_A(y)$, $\delta_B(y)$, $\mu_A(y)$, and $\mu_B(y)$ are all asset pricing model-dependent. The weak integration measure $g(A, B)$ is the only term in Equation (25) that provides a pricing model-independent lower bound on the model specification error measure $\delta(y)$. Thus, it reflects the extent to which *any* asset pricing model can be falsely rejected.

In testing asset pricing models, we should thus first verify the pricing consistency of the data set used. If the price data pass this step, we can then confidently accept or reject the pricing model under consideration, based on its performance in explaining the included asset prices. Otherwise, if it fails, any inferences regarding the empirical validity of the model on this set of assets will be biased. To conduct such a verification, we can, for instance, divide the collected assets into subgroups and, via the procedures discussed next, estimate the weak integration measure between the subgroups.

5.2 Estimation Algorithms

To estimate the market integration measures, we need to solve a minimum distance problem between two convex and closed sets. Although there are readily available methods for solving the classic minimum distance problem between a point and a closed and convex set, solution to the minimum distance problem between two such sets is not that simple. The difficulty arises from the fact that, for two general closed and convex sets, there may, as demonstrated by the following example, exist more than one solution, even though the minimum distance between them is unique.

Example 3. Consider a two-state economy where the payoff span for both markets A and B is the vertical axis, that is, $\mathcal{M}_A = \mathcal{M}_B = \{(0, x) : x \in \mathfrak{R}\}$. Thus, no state 1 payoff is traded on either market. See Figure 2 for the geometric representation of this economy. The two sets

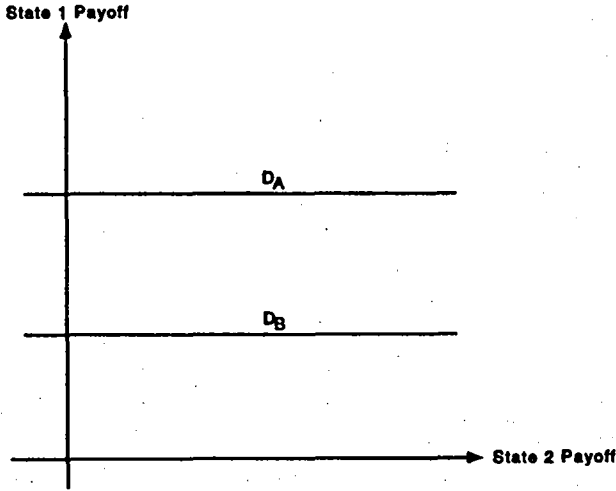


Figure 2

The case with two states and two markets

For both markets, the payoff span is the vertical axis. The line marked D_i , for $i = A, B$, corresponds to the set of stochastic discount factors for market i .

of admissible state price vectors, D_A and D_B , are two lines parallel to the horizontal axis. In this case, the distance between any state price vector $d_A \in D_A$ and a corresponding point $d_B \in D_B$ gives the minimum distance between D_A and D_B . Thus, the solution to the minimization problem in Equation (3) or (4) is not unique. There does not exist a general closed-form solution to these minimization problems.

We instead propose two algorithms to solve the minimum distance problems. First, look at $g(A, B) = \min_{d_A \in D_A, d_B \in D_B} \|d_A - d_B\|^2$. Our algorithm for solving this problem involves two iterative steps. In the first step, we find the minimum squared distance between a point in D_A , say $\hat{d}_A$, and the set D_B . Let $\hat{d}_B$ be the stochastic discount factor in D_B whose distance from $\hat{d}_A$ gives the minimum distance between $\hat{d}_A$ and D_B . In the second step, we project the point $\hat{d}_B$ back onto D_A to find the minimum squared distance between $\hat{d}_B$ and D_A . Restarting from the projected point of $\hat{d}_B$ onto D_A , we repeat these two steps back and forth until we reach a fixed point. The basic logic behind this algorithm is depicted in Figure 3.

Formally, given any $\tilde{d}_A \in D_A$, we intend to find its minimum squared distance from D_B , denoted conveniently by $g(\tilde{d}_A, D_B)$. Define

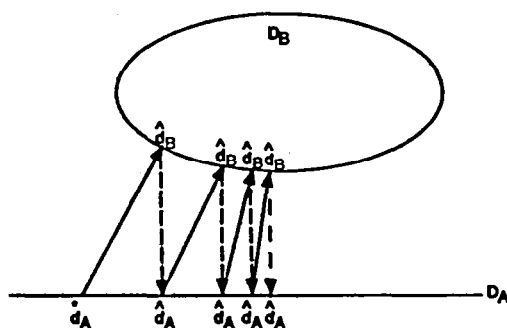


Figure 3

Illustration of Algorithms 1 and 2

D_A and D_B are, respectively, the set of stochastic discount factors for markets A and B. For instance, Algorithm starts at the stochastic discount factor d_A^1 and then projects back and forth between D_A and D_B until it converges to the minimum distance points in D_A and D_B .

a nearest point map $P_B(\cdot) : \mathcal{L}^2 \rightarrow D_B$ by

$$\|\hat{d}_A - P_B(\hat{d}_A)\|^2 \equiv \min_{d_B \in D_B} \|\hat{d}_A - d_B\|^2 = g(\hat{d}_A, D_B). \quad (28)$$

That is, for each $d_A \in D_A$, the distance between d_A and $P_B(d_A)$ is always the minimum distance between d_A and the set D_B . The nearest point map has the following properties: (i) for any $d_B \in D_B$, $P_B(d_B) = d_B$; (ii) for any $x \in \mathcal{L}^2$, $P_B^2(x) = P_B(P_B(x)) = P_B(x)$; and (iii) $P_B(\cdot)$ is nonlinear.

As shown in Hansen and Jagannathan (1992) and, for completeness, in the Appendix of this article, the solution to Equation (28) is as follows:

$$g(\hat{d}_A, D_B) = [E(X_B \hat{d}_A) - \pi_B(X_B)]' [E(X_B X_B')]^{-1} \\ \times [E(X_B \hat{d}_A) - \pi_B(X_B)], \quad (29)$$

which means that $g(\hat{d}_A, D_B)$ reflects the extent that there will be errors when one uses $\hat{d}_A$ to price market B payoffs, and the nearest point map is

$$P_B(\hat{d}_A) = \hat{d}_A - X_B' [E(X_B X_B')]^{-1} [E(X_B \hat{d}_A) - \pi_B(X_B)]. \quad (30)$$

In addition, the market B payoff that leads to the maximum pricing error, $g(\hat{d}_A, D_B)$, is $x_{B, \hat{d}_A} \equiv X'_B \alpha_{B, \hat{d}_A}$, where

$$\alpha_{B, \hat{d}_A} = \frac{[E(X_B X'_B)]^{-1} [\pi_B(X_B) - E(X_B \hat{d}_A)]}{\sqrt{g(\hat{d}_A, D_B)}} \quad (31)$$

(see, again, Hansen and Jagannathan (1992) or the Appendix).

Replacing $\hat{d}_A$, D_B , X_B , and π_B , respectively, with $\hat{d}_B$, D_A , X_A , and π_A , we can similarly define a nearest point map $P_A(\cdot) : \mathcal{L}^2 \rightarrow D_A$ as in Equation (28) and derive the corresponding expressions for $g(D_A, \hat{d}_B)$, $P_A(\hat{d}_B)$, and $x_{B, \hat{d}_A}$ as in Equations (29), (30), and (31).

Algorithm 1. Estimating the weak integration measure

Step 0. Let I be the number of iterations. For the first iteration, set $I = 0$ and $\hat{d}_A^I = d_A^* = X'_A [E(X_A X'_A)]^{-1} \pi_A(X_A)$, where $\hat{d}_A^I$ is the stochastic discount factor in D_A used for the I th iteration [see Equation (48) in the Appendix for the derivation of d_A^*].

Step 1. Compute $g(\hat{d}_A^I, D_B)$, $\hat{d}_B^I = P_B(\hat{d}_A^I)$, and $x_{B, \hat{d}_A^I} = X'_B \alpha_{B, \hat{d}_A^I}$, as given in Equations (29), (30), and (31), where $\hat{d}_B^I$ is the stochastic discount factor in D_B , whose distance from $\hat{d}_A^I$ is the shortest distance between $\hat{d}_A^I$ and D_B for the I th iteration.

Step 2. Compute $g(D_A, \hat{d}_B^I)$, $\hat{d}_A^{I+1} = P_A(\hat{d}_B^I)$ and $x_{A, \hat{d}_B^I} = X'_A \alpha_{A, \hat{d}_B^I}$, as given in Equations (29), (30), and (31) with the appropriate changes mentioned above.

Step 3. Let $I = I + 1$. Repeat steps 1 and 2 for a preset number of times (or, apply some other stopping rule), and then stop.

The algorithm for computing the strong integration measure works similarly except that, in steps 1 and 2, the following problem is repeatedly solved for any $d_A^+ \in D_A^+$:

$$a(d_A^+, D_B^+) = \min_{d_B^+ \in D_B^+} \|d_A^+ - d_B^+\|^2 \equiv \|d_A^+ - P_B^+(d_A^+)\|^2 \quad (32)$$

where $P_B^+(\cdot) : \mathcal{L}^2 \rightarrow D_B^+$ is defined to be the nearest point map onto D_B^+ .

Let $(x)^{\text{pos}} \equiv \max\{x, 0\}$; that is, it is the payoff x truncated at the origin 0. Hansen and Jagannathan (1992) find that a solution to the nearest point problem in Equation (32) is given by

$$P_B^+(d_A^+) = (d_A^+ + \beta_B \cdot X_B)^{\text{pos}}$$

where $\beta_B \in \mathfrak{R}^N$ is chosen so that $P_B^+(d_A^+)$ is in D_B^+ and prices each payoff consistently with π_B . See Hansen and Jagannathan (1992) for a specific procedure that searches for a solution to Equation (33). We can similarly define $P_A^+(\cdot) : \mathcal{L}^2 \rightarrow D_A^+$ as the nearest point map to a point in D_A^+ .

Algorithm 2. Estimating the strong integration measure

Step 0. Let I be the number of iterations. First, set $I = 0$ and $d_A^{I+} = d_A^* = X_A'[E(X_A X_A')]^{-1} \pi_A(X_A)$, where d_A^{I+} is the stochastic discount factor in D_A used for the I th iteration.

Step 1. Compute $a(d_A^{I+}, D_B^+)$ and $d_B^{I+} = P_B^+(d_A^{I+})$ is given in Equations (32) and (33), where d_B^{I+} is the image of the nearest point map, $P_B^+(\cdot)$, to D_B^+ .

Step 2. Compute $a(D_A^+, d_B^{I+})$ and $d_A^{I+1+} = P_A^+(d_B^{I+})$ as given in Equations (32) and (33) with the appropriate changes mentioned above.

Step 3. Let $I = I + 1$. Repeat steps 1 and 2 for a preset number of times and then stop.

Note that in step 0, the algorithm starts at d_A^* . Even though d_A^* may not be nonnegative, the algorithm will correct for that after the first iteration. The central idea behind this algorithm is still to take alternating projections between the two sets D_A^+ and D_B^+ , as seen in Figure 3.

Our next task is to prove that Algorithms 1 and 2 will converge to some fixed points. Given that the two algorithms work similarly, we choose to focus on Algorithm 1. Clearly, steps 1 and 2 in Algorithm 1 together perform a *composite mapping* $T(\cdot) : D_A \rightarrow D_A$,

$$T \equiv P_A P_B, \quad (34)$$

such that, after step 2 for the I th iteration, $\hat{d}_A^{I+1} = T(\hat{d}_A^I) = P_A(P_B(\hat{d}_A^I))$.

We first show that $P_A(\cdot)$, $P_B(\cdot)$, and $T(\cdot)$ are small nonexpansive mappings. By definition, a mapping $T(\cdot) : D_A \rightarrow D_A$ is *nonexpansive* if, for every $c, d \in D_A$, the following is true:

$$\|T(c) - T(d)\| \leq \|c - d\|.$$

Lemma 5. *The nearest point maps $P_A(\cdot)$ and $P_B(\cdot)$, and the composite mapping, $T(\cdot)$, are all nonexpansive.*

Having established the nonexpansiveness of $T(\cdot)$, we need only to use the results from Browder (1965) and Goldburg and Marks (1985) to show that the repeated applications of the composite mapping $T(\cdot)$ will converge to a fixed point in D_A and a corresponding point

in D_B . See Goebel and Kirk (1990) for further characterizations of nonexpansive mappings.

Theorem 4. Under Assumption 1, D_A and D_B are closed and convex. Assume that D_A is also bounded. Let $\hat{d}_A$ be any admissible stochastic discount factor in D_A . Then, (i) the nonexpansive mapping $T(\cdot)$ must have a fixed point in D_A , and the set of all such fixed points, denoted by F_T , is a convex subset of D_A and (ii) as the number of iterations in Algorithm 1, I , goes to infinity,

$$T^I(\hat{d}_A) = T^{I-j}(T^j(\hat{d}_A)), \quad \text{for any } j, 0 \leq j \leq I,$$

will weakly converge to a fixed point in F_T and $(\hat{d}_A^I, D_B)$ will converge to $g(A, B)$. The same conclusion holds for Algorithm 2.

Theorem 4 thus verifies the convergence of Algorithms 1 and 2. In practice, when one chooses a large enough number I , one can expect to obtain a good approximation of the market integration measures. Alternatively, as is done in the next section, one can use a stopping rule so that the iteration process stops once, the market integration measures can no longer go down, ‘much further.’ Also note that, as I goes to infinity in Algorithm 1, $x_{B, \hat{d}_A}^I$ will converge to the common payoff of markets A and B that leads to the maximum pricing difference between A and B, $g(A, B)$. This is true because, by Proposition 3, as $g(\hat{d}_A^I, D_B)$ converges to $g(A, B)$, $x_{B, \hat{d}_A}^I$ has to converge to the common payoff that solves the maximization problem on the right side of Equation (16). Otherwise, $g(\hat{d}_A^I, D_B)$ would not converge to $g(A, B)$.

6. An Illustration

To illustrate the proposed measurement method, we choose; for the following reasons, the New York Stock Exchange and the NASDAQ exchange, respectively, as markets A and B. First, the two exchanges are open during the same time interval each day. Second, if any pair of exchanges is expected to be perfectly integrated, this pair should be among them; at least we should expect them to be integrated in the weak sense. Third, if they are indeed integrated in the weak sense, we can use them to check empirically whether the strong integration measure provides any additional restriction on cross-market pricing relations.

6.1 Data description

For the period from January 1973 to December, 1991, 228 monthly returns on the NYSE and the NASDAQ stocks were collected from

the CRSP stock files. According to a stocks SIC code, the NYSE and the NASDAQ stocks were then respectively sorted into 119 industry groups. The SIC codes were partitioned into 119 groups so as to ensure that each group has a sufficient number of stocks. For each of the 228 months and on either exchange, there is at least one stock in each of the 119 industry groups that was traded in that month. Separately for each exchange, the stocks in an industry group were equally weighted to generate a monthly portfolio return series. Together, each exchange has 119 portfolios constructed from the stock universe.

The portfolios, instead of individual stocks, were used in the test out of the standard considerations in the empirical finance literature. For instance, some individual stocks may not be traded for certain periods or have missing observations, and aggregating stock returns to obtain portfolio returns can help reduce the effect of data measurement errors.

6.2 Estimation results

Proposition 1 and Corollary 1 provide useful guidance for the implementation of market integration tests. First, the integration measures are nondecreasing functions of the number of assets included in each market. Second, by including more assets in each market, we should make the common payoff span $\mathcal{F}$ as large as possible so that the two markets are truly subject to a market integration test. To see that these implications hold empirically, we choose different subsets of assets from NYSE and the NASDAQ and estimate the integration measures across each pair of subsets. We refer to the number of assets in a subset as *group size*. For instance, when the group size is 20, we mean that 20 portfolios (out of 119) from the NYSE and 20 portfolios from the NASDAQ are selected to form markets A and B, respectively. We examine group sizes 10 and 20.

For a given group size, say N , it becomes somewhat arbitrary, however, which N portfolios out of 119 will be selected from each exchange. Clearly, depending on the choice of the two subsets of N portfolios, the estimated values of the integration measures can be quite different. To control for such biases, we adopt the following randomization procedure:

- Step 1. Randomly select N portfolios from the 119 NYSE portfolios and N from the 119 NASDAQ portfolios, which gives a pair of N asset submarkets.
- Step 2. Following Algorithms 1 and 2, respectively, estimate $q(A, B)$ and $a(A, B)$ on the two submarkets selected in Step 1.

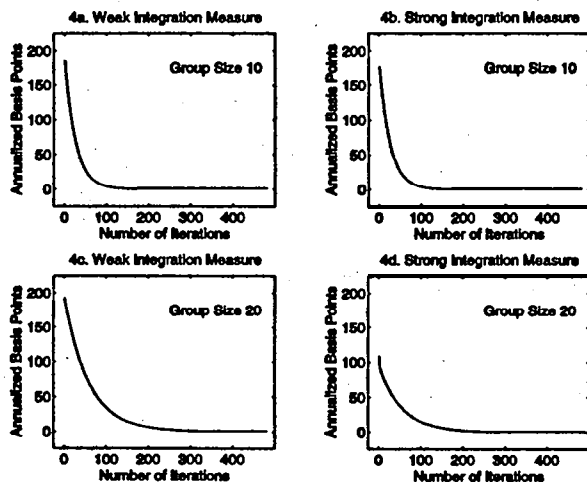


Figure 4

Convergence paths for the market integration measures

For each integration measure and each group size, there is one sample convergence path reported above. In Algorithms 1 and 2, each iteration process was allowed to continue until the sum of the previous five changes in the corresponding integration measure was less than 0.05 in annualized basis points.

- Step 3. Collect the sample values for $g(A, B)$ and $a(A, B)$ obtained in step 2, and then repeat steps 1 and 2 for 500 random draws.

We run the above procedure separately for each group size and each integration measure. The purpose is to test the hypotheses:

- H_0 : Between any pair of submarkets each with N portfolios, there is perfect weak-form integration, that is, $g(A, B) = 0$.
- H'_0 : Between any pair of submarkets each with N portfolios, there is perfect strong-form integration; that is, $a(A, B) = 0$.

The iteration process in Algorithms 1 and 2 stops once *the sum of absolute changes in the estimated integration measure for the last five iterations is less than 0.05 basis points*. This stopping rule is quite satisfactory, as shown by the example convergence paths in Figure 4.

All estimation results are reported in Table 1. Based on 500 random draws, the mean value for the weak integration measure is 0.065 basis points for group size 10 and 0.44 basis points for group size 20. The frequency distributions based on the 500 estimated values of the weak measure are given in Figure 5 for the two group sizes. While the estimated means in both cases are more than two standard deviations away from zero, they are quite small in economic terms. In addition,

Table 1
Estimation of the market integration measures^a

	Group size	Mean	Standard deviation	Min	Max	Average no. of iterations
The weak measure $g(A, B)$	10	0.14	0.06	0.03	0.44	352
	20	0.44	0.08	0.25	0.74	472
The strong measure $\alpha(A, B)$	10	0.25	0.06	0.12	0.45	387
	20	0.47	0.08	0.27	0.78	497

^a This table is generated as follows. Take group size 20 as an example. In step 1, randomly select 20 portfolios from the 119 NYSE portfolios (forming market A) and another 20 from the 119 NASDAQ portfolios (forming market B). In step 2, follow Algorithms 1 and 2 to estimate respectively the weak and strong integration measures, and collect the estimated values. Repeat the random drawing in step 1. and the estimation in step 2, for 500 times. Then, for the given group size, we obtain a sample distribution for each of the two integration measures. The mean, standard deviation, and lowest (min) and highest (max) values of each sample distribution are reported in the following table, and each sample distribution is shown in Figure 5.

the highest values for the measure are also small for both group sizes. Thus, the two markets are closely integrated in the weak sense. This is consistent with our priors; because it would be surprising if the law of one price would be violated between the NYSE and the NASDAQ markets even at the monthly time interval.

The mean value for the strong integration measure is 0.25 basis points for group size 10 and 0.47 basis points for group size 20, and both estimates are again more than two standard deviations away from zero. Economically, these estimates are insignificant, which suggests that the two exchanges are also closely integrated in the strong sense. The frequency distributions of the strong measure are shown in Figure 5, based on 500 random draws.

In Table 1, it is apparent that both integration measures are increasing in group size. For the same group size, the strong integration measure is higher than the weak measure. These are consistent with the predictions of theory.

By Proposition 3, the maximum pricing difference $g(A, B)$ is achieved over some common payoff of A and B. By Theorem 4, this pricing difference-maximizing payoff can be identified by the limit of the portfolio sequence $\alpha_{B, \hat{\alpha}_A}$ in Algorithm 1. Economically, this payoff represents the potentially most profitable portfolio to buy and sell in an arbitrage strategy. Thus, it is interesting to see what positions this limit portfolio may typically involve of each security.⁶ For this pur-

⁶ We thank the referee for pointing out this connection as well as the importance of keeping track of the portfolio weight information in $\alpha_{B, \hat{\alpha}_A}$.

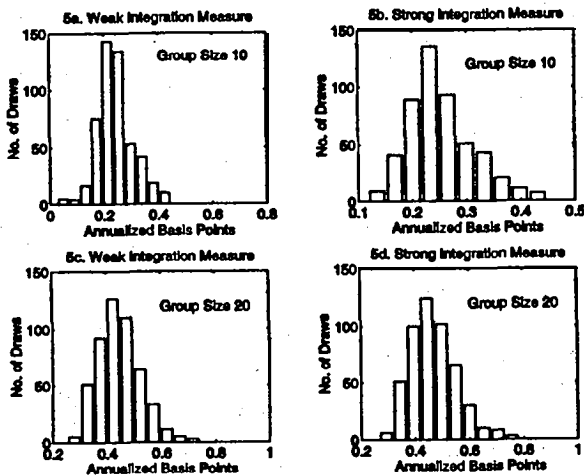


Figure 5

Histograms for the market integration measures

Each histogram was generated as follows. First, a group size was selected (e.g., 50). For this group size the requisite number of portfolios were randomly selected without replacement from the 119 NYSE portfolios, and the same number of portfolios were randomly selected from the 119 NASDAQ portfolios. Then, for the selected portfolios from each exchange, the market integration measures were computed. For each given group size and each integration measure, the above two steps were repeated 500 times to generate the corresponding histogram.

pose, we report in Table 2 the following characteristics of the limit portfolio weights for five pairs of size 20 groups: the lowest weight in any of the 20 securities (min), the highest weight in any security (max), the average weight (mean), the number of negative weights (no. of negatives), and the average size of the negative positions (mean of negatives). First, recall that the pricing difference-maximizing payoff is normalized to have a unit norm. Hence, not surprisingly, the reported portfolio weights in Table 2 are generally quite small. Second, for the five sample pairs, the number of negative portfolio weights ranges from 13 to 16 out of 20. That is, the pricing difference-maximizing common payoff involves short positions in over half of the stock portfolios. Provided that only five sample pairs are reported for group size 20, this is a significant number of short positions.

The above exercise demonstrates that the proposed integration measures depend crucially on the frictionless market assumption. For instance, when short selling is not allowed on either market A or B or both, arbitrage strategies will become harder to implement, as illustrated in Table 2. This leads to the suggestion that future work is necessary to extend the discussion here to economics with trading

Table 2
Summary statistics of sample portfolio weight.^a

Sample	Min	Max	Mean	No. of negatives	Mean of negatives	$g(A, B)$
1	-1.15e-5	6.12e-5	-3.58e-6	13	-6.56e-6	0.40
2	-1.55e-5	4.39e-5	-3.80e-6	15	-5.71e-6	0.49
3	-1.87e-5	2.48e-6	-3.63e-6	16	-4.91e-6	0.42
4	-1.27e-5	4.98e-6	-3.66e-6	16	-5.17e-6	0.45
5	-1.66e-5	5.01e-6	-3.73e-6	15	-5.74e-6	0.43

^a Here five sample pairs of submarkets, each formed out of 20 NYSE and 20 NASDAQ portfolios, are included. For each sample pair, Algorithm 1 is run and the weights of the limit NYSE portfolio that gives the maximum pricing difference between the two submarkets are used to calculate the following summary statistics: the lowest portfolio weight assigned to any of the 20 NYSE portfolios (min), the highest weight (max), the average weight (mean), the number of negative weights (no. of negatives), the average size of the negative weights (mean of negatives), and the corresponding value of $g(A, B)$.

frictions and transaction costs. The analysis in this article serves as a useful benchmark, nonetheless.

As a final note, we have not yet worked out a sampling theory for the two measures, $g(A, B)$ and $a(A, B)$. Presumably, we can follow the ideas in Hansen, Heaton, and Luttmer (1993) to accomplish this. In their case, they exploit the concavity of the population criterion function and the convexity of the constraint set to derive the asymptotic distribution of the Hansen-Jagannathan (1992) model specification error measure. Since the criterion functions implied by the market integration measures share a similar structure with theirs, their approach may generalize to the problems specified here. We leave this important task for future research.

7. Conclusions

In this article, we have formalized the idea that closely integrated markets should assign to similar payoffs prices that are close and have defined two notions of market integration. Such a formalization has allowed us to develop two market integration measures. Based on the weak-form integration notion, the weak measure reflects how differently two markets price their common payoffs. Strong-form integration requires the markets not only to assign the same price to the same common payoff but also to price every conceivable payoff in such a way that no arbitrage profits can be generated across the two markets. Correspondingly, the strong integration measure is generally higher than the weak measure. They both indicate the extent that the respective notion of perfect market integration is violated across two markets. Since our development does not rely on any restrictive pricing model assumptions, its applications do not depend on the

empirical validity of any asset pricing model. This is in contrast with the existing literature, in which authors can only conduct joint tests of market integration and a specific asset pricing model.

The proposed measurement framework can be applied in the study of cross-market pricing relations (both domestic and international) as well as in empirical tests of asset pricing models. As an illustration, we chose the NYSE and the NASDAQ as two sample markets. The estimation results suggest that the two exchanges are closely integrated in the weak sense but not in the strong sense. That is, based on the data, they seem to violate the strong-form integration, and cross-market arbitrage may be possible. There is no positive pricing rule that is simultaneously consistent with the NYSE and the NASDAQ.

Appendix

Proof Lemma 1. See Chamberlain and Rothschild (1983). ■

Proof of Lemma 2. See Hansen and Richard (1987). ■

Proof of Theorem 1. Let (d'_A, d'_B) be a solution to Equation (3). By a property of the norm,

$$g(A, B) = \min_{d_A \in D_A, d_B \in D_B} \|d_A - d_B\|^2 = \|d'_A - d'_B\|^2 = 0$$

if and only if $d'_A = d'_B$ a.s., which is true if and only if $D_A \cap D_B \neq \emptyset$. ■

Proof of Proposition 1. As noted earlier, the set D_A may shrink as N_A increases; that is, $D_A(N_A + i) \subseteq D_A(N_A)$ for any integer $i \geq 0$. Similarly, $D_B(N_B + j) \subseteq D_B(N_B)$ for any integer $j \geq 0$. Thus,

$$\begin{aligned} g(A, B) &= \min_{d_A \in D_A(N_A), d_B \in D_B(N_B)} \|d_A - d_B\|^2 \\ &\leq \min_{d_A \in D_A(N_A + i), d_B \in D_B(N_B + j)} \|d_A - d_B\|^2, \end{aligned}$$

for any $i, j \geq 0$. ■

Proof of Theorem 2. It follows from Definitions 5 and 6 and a property of the norm. ■

Proof of Proposition 2. Hansen and Jagannathan (1992) establish that, for any given $d_A \in D_A$, $\min_{d_B \in D_B} \|d_A - d_B\|^2 = \|Q_B(d_A) - d_B^*\|^2$, where $Q_B(\cdot)$ is the projection operator onto $\mathcal{M}_B$. If $Q_B(d_A) = d_B^*$, the desired result holds trivially. Suppose $Q_B(d_A) \neq d_B^*$. Then, note that $x' = \frac{Q_B(d_A) - d_B^*}{\|Q_B(d_A) - d_B^*\|}$ a marketed payoff in $\mathcal{M}_B$, with $\|x'\| = 1$. When x' is

priced by $\hat{\pi}_{d_A}$, its squared pricing error is precisely $x' \cdot (d_A - d_B^*)^2 = \|Q_B(d_A) - d_B^*\|^2$, which, according to Equation (7), is also the upper bound on squared pricing errors for payoffs $x \in \mathcal{M}_B$ with $\|x\| = 1$. Thus,

$$\sup_{x \in \mathcal{M}_B, \|x\|=1} |\hat{\pi}_{d_A}(x) - \pi_B(x)|^2 = \min_{d_B \in D_B} \|d_A - d_B\|^2. \quad (35)$$

Taking the infimum of both sides of the above equation over the set D_A gives the first part of Equation (10). For the other equality, the proof steps are the same. ■

Proof of Lemma 3. We first show that $\bar{g}(A, B) = \|\bar{d}_A - \bar{d}_B\|^2$. To do this, note that, for any $d_A \in \bar{D}_A$ and $d_B \in \bar{D}_B$, there exist ϵ_1 and ϵ_2 such that (i) $d_A = \bar{d}_A + \epsilon_1$ and (ii) $d_B = \bar{d}_B + \epsilon_2$, where both ϵ_1 and ϵ_2 are orthogonal to $\bar{d}_A$ and $\bar{d}_B$. Then,

$$\begin{aligned} \|d_A - d_B\|^2 &= \|\bar{d}_A + \epsilon_1 - (\bar{d}_B + \epsilon_2)\|^2 \\ &= \|\bar{d}_A - \bar{d}_B\|^2 + \|\epsilon_1 - \epsilon_2\|^2 \\ &\geq \|\bar{d}_A - \bar{d}_B\|^2 \end{aligned}$$

which establishes that the minimum squared distance between $\bar{D}_A$ and $\bar{D}_B$ is equal to $\|\bar{d}_A - \bar{d}_B\|^2$.

Next, we show that $\|\bar{d}_A - \bar{d}_B\|^2 = \|f_A^* - f_B^*\|^2$. Since the projections of $\bar{d}_A$ and $\bar{d}_B$ onto $\mathcal{F}$ are, respectively, f_A^* and f_B^* , it holds that $(\bar{d}_A - f_A^*)$ and $(\bar{d}_B - f_B^*)$ are both orthogonal to f_A^* and f_B^* , which implies

$$\|\bar{d}_A - \bar{d}_B\|^2 = \|f_A^* - f_B^*\|^2 + \|(\bar{d}_A - f_A^*) - (\bar{d}_B - f_B^*)\|^2. \quad (36)$$

Let $\mathcal{F}^\perp$ be the orthogonal complement of $\mathcal{F}$ in $\mathcal{M}$, where $\mathcal{M} = \mathcal{M}_A + \mathcal{M}_B$. That is, $\mathcal{F}^\perp = \mathcal{F}_A^\perp + \mathcal{F}_B^\perp$. Clearly, $\mathcal{M} = \bar{\mathcal{M}}_A = \bar{\mathcal{M}}_B$. By construction, $(\bar{d}_A - f_A^*)$ represents in the sense of the Riesz representation theorem the restriction of $\bar{\pi}_A$ to $\mathcal{F}^\perp$. Therefore, it can be used to price consistently the random payoffs in e_A and e_B , which are linearly independent and together form a basis of $\mathcal{F}^\perp$, as follows:

$$E[e_{A,i} \cdot (\bar{d}_A - f_A^*)] = \pi_A(e_{A,i}) \quad \forall e_{A,i} \in e_A \quad (37)$$

$$E[e_{B,i} \cdot (\bar{d}_A - f_A^*)] = \pi_B(e_{B,i}) \quad \forall e_{B,i} \in e_B. \quad (38)$$

(Recall that $e_A \subset \mathcal{M}_A$ and $e_B \subset \mathcal{M}_B$ are, respectively, the bases of $\mathcal{F}_A^\perp$ and $\mathcal{F}_B^\perp$.) For the same reason, the above two equations should hold when $(\bar{d}_A - f_A^*)$ is replaced by $(\bar{d}_B - f_B^*)$. Therefore, the restriction of $\bar{\pi}_A$ to $\mathcal{F}^\perp$ must be identical to the restriction of $\bar{\pi}_B$ to $\mathcal{F}^\perp$. According to the Riesz representation theorem, the random payoff in $\mathcal{F}^\perp$ that represents

the restriction of $\tilde{\pi}_A$ to $\mathcal{F}^\perp$ must be unique, meaning $(\bar{d}_A - f_A^*) = (\bar{d}_B - f_B^*)$ a.s. Hence, the last term in Equation (36) is equal to zero, which completes the proof. ■

Proof of Lemma 4. (i) For every $d \in D_A \cap D_B$, it must price every payoff in $\mathcal{M}$ consistently with both $\tilde{\pi}_A$ and $\tilde{\pi}_B$, implying $d \in \bar{D}_A \cap \bar{D}_B$. On the other hand, $\bar{D}_A \cap \bar{D}_B \subseteq D_A \cap D_B$ due to the fact that $\bar{D}_A \subseteq D_A$ and $\bar{D}_B \subseteq D_B$. This proves the first part.

(ii) By Theorem 1, $g(A, B) = 0$ if and only if $D_A \cap D_B \neq \emptyset$. When $D_A \cap D_B \neq \emptyset$, then $\bar{D}_A \cap \bar{D}_B \neq \emptyset$ by part (i), which means $\bar{g}(A, B) = 0$. From the proof of Lemma 3, this is equivalent to $\bar{d}_A = \bar{d}_B$. Thus, we have $\bar{D}_A = \bar{D}_B = \bar{d}_A + \mathcal{M}^\perp$. Relying on part (i), one can easily prove the reverse implication. ■

Proof of Theorem 3. First, note that for any $d_A \in D_A$ and $d_B \in D_B$, it holds that $(d_A - f_A^*) \perp \mathcal{F}$ and $(d_B - f_B^*) \perp \mathcal{F}$, by the projection theorem for Hilbert space. Thus,

$$\begin{aligned} \|d_A - d_B\|^2 &= \|(d_A - f_A^*) - (d_B - f_B^*) + (f_A^* - f_B^*)\|^2 \\ &= \|(d_A - f_A^*) - (d_B - f_B^*)\|^2 + \|f_A^* - f_B^*\|^2 \\ &\geq \|f_A^* - f_B^*\|^2 \\ &= \bar{g}(A, B) \end{aligned}$$

where the last equality holds by Lemma 3, which implies

$$g(A, B) = \min_{d_A \in D_A, d_B \in D_B} \|d_A - d_B\|^2 \geq \bar{g}(A, B). \quad (39)$$

Next, we have

$$\begin{aligned} \bar{g}(A, B) &= \min_{d_A \in \bar{D}_A, d_B \in \bar{D}_B} \|d_A - d_B\|^2 \\ &\geq \min_{d_A \in D_A, d_B \in D_B} \|d_A - d_B\|^2 \\ &= g(A, B) \end{aligned}$$

due to the fact that $\bar{D}_A \subseteq D_A$ and $\bar{D}_B \subseteq D_B$. Therefore, by Lemma 3, the desired conclusion is true. ■

Proof of Proposition 3. By construction, the projection of any $d_A \in D_A$ onto $\mathcal{F}$ is f_A^* . Thus, for any $d_A \in D_A$, there must be some ϵ_A such that $\epsilon_A \perp \mathcal{F}$ and $d_A = f_A^* + \epsilon_A$. Similarly, for any $d_B \in D_B$, there exists

some ϵ_B such that $\epsilon_B \perp \mathcal{F}$ and $d_B = f_B^* + \epsilon_B$. Hence, for any $x \in \mathcal{F}$,

$$\begin{aligned} |\pi_A(x) - \pi_B(x)|^2 &= |E(x \cdot d_A) - E(x \cdot d_B)|^2 \\ &\quad \text{for any } d_A \in D_A \text{ and } d_B \in D_B \\ &= |E[x \cdot (f_A^* + \epsilon_A - f_B^* - \epsilon_B)]|^2 \\ &= |E[x \cdot (f_A^* - f_B^*)]|^2 \text{ (since } x \perp \epsilon_A \text{ and } x \perp \epsilon_B) \\ &\leq \|f_A^* - f_B^*\|^2 \|x\|^2 \end{aligned}$$

where the last inequality holds by the Cauchy-Schwarz inequality. Taking the supremum of the left-hand side of the above inequality over $\{x \in \mathcal{F} : \|x\| = 1\}$ yields

$$\sup_{x \in \mathcal{F}, \|x\|=1} |\pi_A(x) - \pi_B(x)|^2 \leq \|f_A^* - f_B^*\|^2. \quad (40)$$

Since both f_A^* and f_B^* are in $\mathcal{F}$, $x' \equiv \frac{f_A^* - f_B^*}{\|f_A^* - f_B^*\|}$ must be in $\mathcal{F}$ with $\|x'\| = 1$. Then, for x' , its squared pricing difference implied by π_A and π_B is

$$|\pi_A(x') - \pi_B(x')|^2 = |E[x' \cdot (f_A^* - f_B^*)]|^2 = \|f_A^* - f_B^*\|^2.$$

This means that the inequality in Equation (40) has to hold with equality. By Theorem 3, $g(A, B) = \|f_A^* - f_B^*\|^2$, which, together with the equality in Equation (40), completes the proof. ■

Proof of Corollary 1. In this case, $\mathcal{F} = \{0\}$. Hence $f_A^* = f_B^* = 0$, resulting in $g(A, B) = 0$ by Theorem 3. ■

Proof of Corollary 2. Since f_A^* can, by construction, be used to price the payoffs in $\mathcal{F}$ consistently with π_A , we have

$$E(f_j \cdot f_A^*) = \pi_A(f_j) \quad \forall f_j \in \mathcal{F}. \quad (41)$$

Because $f_A^* \in \mathcal{F}$, there must exist some $\alpha \in \mathcal{H}^k$ such that $f_A^* = \alpha' \cdot f$. Substituting this into Equation (41) and recombining the expressions yields

$$f_A^* = f' \cdot E(ff')^{-1} \cdot \pi_A(f). \quad (42)$$

Similarly,

$$f_B^* = f' \cdot E(ff')^{-1} \cdot \pi_B(f). \quad (43)$$

Substituting Equations (42) and (43) into Equation (15) in Theorem 3 and arranging the terms produces Equation (17). ■

Proof of Proposition 4. For any $d_A^+ \in D_A^+$ and $d_B^+ \in D_B^+$, we have by the Cauchy-Schwarz inequality,

$$|\pi_A^+(x) - \pi_B^+(x)| = |E(x \cdot d_A^+) - E(x \cdot d_B^+)| \leq \|d_A^+ - d_B^+\| \|x\|,$$

which implies

$$\sup_{x \in \mathcal{L}^2, \|x\|=1} |\pi_A^+(x) - \pi_B^+(x)|^2 \leq \|d_A^+ - d_B^+\|^2. \quad (44)$$

If $d_A^+ = d_B^+$, then both sides of the inequality in Equation (44) are clearly zero, which means that both sides of Equation (19) are zero. Otherwise, if $d_A^+ \neq d_B^+$ for any $d_A^+ \in D_A^+$ and $d_B^+ \in D_B^+$, we can choose payoff $x' = \frac{d_A^+ - d_B^+}{\|d_A^+ - d_B^+\|}$. Then, $\|x'\| = 1$ and $|\pi_A^+(x') - \pi_B^+(x')|^2 = |E[x' \cdot (d_A^+ - d_B^+)]|^2 = \|d_A^+ - d_B^+\|^2$. Therefore, the weak inequality in Equation (44) must hold with strict equality. Replacing the inequality in Equation (44) with equality and taking the minimum of both sides over the sets D_A^+ and D_B^+ yield the desired result in Equation (19). ■

Proof of Proposition 5. For any fixed $d_1 \in D_1$, we have from the proof of Proposition 2 that, for any k ,

$$\sup_{x \in \mathcal{M}_k, \|x\|=1} |E(x \cdot d_1)|^2 = \min_{d_k \in D_k} \|d_1 - d_k\|^2.$$

Summing the above equation over $k = 2, \dots, K$ and taking the infimum of the resulting equation over D_1 produces Equation (21). ■

Proof of Proposition 6. Suppose $\delta_A(y) < \delta_B(y)$. Then,

$$\begin{aligned} \delta(y) &\geq \max\{\delta_A(y), \delta_B(y)\} = \delta_B(y) \\ &= \min_{d_B \in D_B} \|(y - d_A) + (d_A - d_B)\|^2 \\ &= \|y - d_A\|^2 + \min_{d_B \in D_B} \{2E[(y - d_A)(d_A - d_B)] + \|d_A - d_B\|^2\} \\ &\geq \min_{d_A \in D_A} \|y - d_A\|^2 + \min_{d_A \in D_A, d_B \in D_B} \{2E[(y - d_A)(d_A - d_B)]\} \\ &\quad + \min_{d_A \in D_A, d_B \in D_B} \|d_A - d_B\|^2 \\ &= \delta_A(y) + \mu_A(y) + g(A, B), \end{aligned}$$

Where $\mu_A(y)$ is as defined in Equation (26). Similarly, if $\delta_A(y) \geq \delta_B(y)$,

$$\delta(y) \geq \delta_B(y) + \mu_B(y) + g(A, B)$$

where $\mu_B(y)$ is as defined in Equation (27). Combining the above

inequalities yields

$$\delta(y) \geq \max\{\delta_A(y) + \mu_A(y) + g(A, B), \delta_B(y) + \mu_B(y) + g(A, B)\},$$

which gives the desired result in Equation (25). ■

Solution to the nearest point problem in (28). To solve explicitly the classic nearest point problem in Equation (28), project given random variable $\hat{d}_A$ onto $\mathcal{M}_B$, which results in the decomposition $\hat{d}_A = Q_B(\hat{d}_A) + \hat{\epsilon}_A$, where $Q_B(\cdot) : \mathcal{L}^2 \rightarrow \mathcal{M}_B$ is the projection operator from any point in $\mathcal{L}^2$ onto $\mathcal{M}_B$. From the fact that for any $d_B \in D_B$, there exists some ϵ_B such that $\epsilon_B \perp \mathcal{M}_B$ and $d_B = d_B^* + \epsilon_B$, we have

$$\begin{aligned} \|\hat{d}_A - d_B\|^2 &= \|Q_B(\hat{d}_A) + \hat{\epsilon}_A - (d_B^* + \epsilon_B)\|^2 \\ &\geq \|Q_B(\hat{d}_A) - d_B^*\|^2 + \|\hat{\epsilon}_A - \epsilon_B\|^2 \\ &\geq \|Q_B(\hat{d}_A) - d_B^*\|^2. \end{aligned}$$

Since $\hat{\epsilon}_A \perp \mathcal{M}_B$, there must be some $d_B' \in D_B$ such that $d_B' = d_B^* + \hat{\epsilon}_A$ and $\|\hat{d}_A - d_B'\|^2 = \|Q_B(\hat{d}_A) - d_B^*\|^2$. Thus, the minimum distance between $\hat{d}_A$ and D_B is achieved at d_B' ; that is,

$$g(\hat{d}_A, D_B) = \|Q_B(\hat{d}_A) - d_B^*\|^2. \quad (45)$$

To compute $g(\hat{d}_A, D_B)$, note that $Q_B(\hat{d}_A)$ is in $\mathcal{M}_B$. Hence, there exists some $\alpha \in \mathbb{R}^{N_B}$ such that $Q_B(\hat{d}_A) = X_B' \alpha$ and

$$E(X_B X_B' \alpha) = E(X_B \hat{d}_A) \quad (46)$$

where X_B denotes the column vector of the N_B basis payoffs on market B, and Equation (46) states that $Q_B(\hat{d}_A)$ should price every payoff of market B consistently with the stochastic discount factor $\hat{d}_A$. Assuming that $E(X_B X_B')$ is invertible, we have

$$Q_B(\hat{d}_A) = X_B' [E(X_B X_B')]^{-1} E(X_B \hat{d}_A). \quad (47)$$

By construction, d_B^* is in $\mathcal{M}_B$ representing the pricing functional π_B , meaning that there exists some $\alpha^* \in \mathbb{R}^{N_B}$ such that $d_B^* = X_B' \alpha^*$ and that

$$E(X_B X_B' \alpha^*) = \pi_B(X_B)$$

where $\pi_B(X_B)$ stands for the column vector of prices for the N_B payoffs in X_B . Thus,

$$d_B^* = X_B' [E(X_B X_B')]^{-1} \pi_B(X_B). \quad (48)$$

Substituting Equations (47) and (48) into Equation (45) yields Equation (29). Next, according to Equations (28) and (45), the nearest point map is given by $P_B(\hat{d}_A) = \hat{d}_A - Q_B(\hat{d}_A) + d_B^*$, which, together with Equations (47) and (48), implies Equation (30).

By the argument in the proof of Proposition 2,

$$x_{B, \hat{d}_A} = \frac{d_B^* - Q_B(\hat{d}_A)}{\|d_B^* - Q_B(\hat{d}_A)\|} \equiv X_B' \alpha_{B, \hat{d}_A}$$

is a solution to the maximization problem on the left side of Equation (35), where, by Equations (45), (47), and (48), $\alpha_{B, \hat{d}_A}$ should be as given in Equation (31). ■

Proof of Lemma 5. We only need to show that $P_A(\cdot) : \mathcal{L}^2 \rightarrow D_A$ and $T(\cdot) : D_A \rightarrow D_A$ are nonexpansive.

(i) Given the definition of $P_A(\cdot)$ in Equation (28), it is well known that, for any $x \in \mathcal{L}^2$ and any $d_A \in D_A$,

$$\langle x - P_A(x) \mid d_A - P_A(x) \rangle \leq 0 \quad (49)$$

where $\langle \cdot \mid \cdot \rangle$ is the mean-square inner product. See Luenberger (1969) for this result. For any two random variables $x_1, x_2 \in \mathcal{L}^2$, we have

$$\langle x_1 - P_A(x_1) \mid P_A(x_2) - P_A(x_1) \rangle \leq 0 \quad (50)$$

$$\langle x_2 - P_A(x_2) \mid P_A(x_1) - P_A(x_2) \rangle \leq 0. \quad (51)$$

Adding Equations (50) and (51) together and rearranging the terms yields

$$\begin{aligned} 0 &\geq \langle P_A(x_1) - P_A(x_2) \mid P_A(x_1) - P_A(x_2) \rangle \\ &\quad - \langle x_1 - x_2 \mid P_A(x_1) - P_A(x_2) \rangle \\ &= \|P_A(x_1) - P_A(x_2)\|^2 - \langle x_1 - x_2 \mid P_A(x_1) - P_A(x_2) \rangle. \end{aligned}$$

Thus,

$$\begin{aligned} \|P_A(x_1) - P_A(x_2)\|^2 &\leq |\langle x_1 - x_2 \mid P_A(x_1) - P_A(x_2) \rangle| \\ &\leq \|x_1 - x_2\| \|P_A(x_1) - P_A(x_2)\| \end{aligned}$$

where the last inequality holds due to the Cauchy-Schwarz inequality, which implies that $\|P_A(x_1) - P_A(x_2)\| \leq \|x_1 - x_2\|$. Hence, $P_A(\cdot)$ is nonexpansive.

(ii) We have just shown that both $P_A(\cdot)$ and $P_B(\cdot)$ are nonexpansive. Then, for any $x_1, x_2 \in \mathcal{L}^2$, we have

$$\begin{aligned} \|T(x_1) - T(x_2)\| &\leq \|P_A(P_B(x_1)) - P_A(P_B(x_2))\| \\ &\leq \|P_B(x_1) - P_B(x_2)\| \\ &\leq \|x_1 - x_2\| \end{aligned}$$

which completes the proof. ■

Proof of Theorem 4. For the proof of part (i), see Browder (1965). The convergence of $T^I(\hat{d}_A)$ to a point in F_T is proven in Goldberg and Marks (1985). Here, we show only that $g(\hat{d}_A^I, D_B)$ will converge to $g(A, B)$ as I goes to ∞ . For this purpose, let $\hat{d}_B^I = P_B(\hat{d}_A^I)$, where $P_B(\cdot)$ is the nearest point map to D_B during the I th iteration. Similarly, in Algorithm 1, $\hat{d}_A^{I+1} = P_A(\hat{d}_B^I)$. Then, by design, we have

$$g(\hat{d}_A^I, D_B) \geq g(D_A, \hat{d}_B^I) \geq g(\hat{d}_A^{I+1}, D_B),$$

that is, $g(\hat{d}_A^I, D_B)$ is nonincreasing in I . Therefore, when $T^I(\hat{d}_A)$ converges as $I \rightarrow \infty$, $g(\hat{d}_A^I, D_B)$ must converge to the minimum squared distance between D_A and D_B , which is precisely $g(A, B)$. ■

References

- Adler, M., and B. Dumas, 1983, "International Portfolio Choice and Corporation Finance: A Synthesis," *Journal of Finance* 38, 925-964.
- Black, F., 1974, "International Capital Market Equilibrium with Investment Barriers," *Journal of Financial Economics*, 1, 337-352.
- Black, F., and M. Scholes, 1973, "The Pricing of Options and Corporate Liabilities," *Journal of Political Economy*, 81, 637-659.
- Braun, P., 1991, "Asset Pricing and Capital Investment: Theory and Evidence," working paper, Northwestern University.
- Browder, F., 1965, "Nonexpansive Nonlinear Operators in a Banach Space," *Proceedings of the National Academy of Sciences*, 54, 1041-1044.
- Chamberlain, G., and M. Rothschild, 1983, "Arbitrage, Factor Structure, and Mean-Variance Analysis on Large Asset Markets," *Econometrica*, 51, 1281-1304.
- Chen, Z., and P. Knez, 1994, "A Pricing Operator-Based Testing Foundation for a Class of Factor Pricing Models," *Mathematical Finance*, 4, 121-141.
- Cho, D., C. Eun, and L. Senbet, 1986, "International Arbitrage Pricing Theory: An Empirical Investigation," *Journal of Finance*, 41, 313-329.
- Goebel, K., and W. Kirk, 1990, *Topics in Fixed Point Theory*, Cambridge University Press, New York.

Goldburg, M., and R. Marks III, 1985, "Signal Sythesis in the Presence of an Inconsistent Set of Constraints," *IEEE Transactions on Circuits and Systems*, CAS-32, 647-663.

Gultekin, M., N. Gultekin, and A. Penati, 1989, "Capital Controls and International Market Segmentation: The Evidence from the Japanese and American Stock Markets," *Journal of Finance*, 44, 849-869.

Hansen, L, J. Heaton, and E. Luttmer, 1993, "Econometric Evaluation of Asset Pricing Models," working paper, University of Chicago.

Hansen, L, and R. Jagannathan, 1991. "Implications of Security Market Data for Models of Dynamic Economies," *Journal of Political Economy*, 99, 225-262.

Hansen, and R. Jagannathan, 1992, "Assessing Specification Errors In Stochastic Discount Factor Models." working paper, University of Chicago.

Hansen, L., and S. Richard, 1987, "The Role of Conditioning Information In Deducing Testable Restrictions Implied by Dynamic Asset Pricing Models," *Econometrica*, 55, 587-613.

Harrison, J., and D. Kreps, 1979, "Martingales and Arbitrage in Multiperiod Securities Markets," *Journal of Economic Theory*, 20, 381-408.

Jorion, P., and E. Schwartz, 1986, "Integration vs. Segmentation in the Canadian Stock Market," *Journal of Finance*, 41, 603-614.

Knez, P., 1991, "Pricing Money Market Securities with Stochastic Discount Factors," working paper, University of Wisconsin-Madison.

Korajczyk, R., and C. Viallet, 1990, "An Empirical Investigation of International Asset Pricing," *Review of Financial Studies*, 3, 353-385.

Luenberger, D., 1969, *Optimization by Vector Space Methods*, John Wiley & Sons, New York.

Ross, S., 1976, "The Arbitrage Theory of Capital Asset Pricing," *Journal of Economic Theory*, 12, 341-360.

Ross, S., 1978. "A Simple Approach to the Valuation of Risky Streams," *Journal of Business*, 51, 453-475.

Snow, K., 1991, "Diagnosing Asset Pricing Models Using the Distribution of Asset Returns," *Journal of Finance*, 46, 955-983.

Solnik, B., 1933, "International Arbitrage Pricing Theory," *Journal of Finance*, 38, 449-457.

Stehle, R., 1977, "An Empirical Test of the Alternative Hypotheses of National and International Pricing of Risky Assets," *Journal of Finance*, 32, 493-502.

Stulz, R., 1981, "A Model of International Asset Pricing." *Journal of Financial Economics*, 9, 383-406.

Wheatley, S., 1988, "Some Tests of International Equity Integration," *Journal of Financial Economics*, 21, 177-222.