

Empirical Asset Pricing

Francesco Franzoni

Swiss Finance Institute - USI Lugano

The Efficient Market Hypothesis

1. **Course Overview**
2. **Definition of Market Efficiency**
3. **Efficient Capital Markets II** (Fama, JF, 1991)
4. **Efficiently Inefficient** (Berk and Green, JPE, 2004; Garleanu and Pedersen, JF, 2018)
5. **Market Efficiency and Learning** (Martin and Nagel, 2022)

Relevant readings:

- Listen to Cochrane's short history of Asset Pricing at the Global Banking and Finance Conference – XLRI Delhi NCR 2026. YouTube <https://www.youtube.com/watch?v=MnTdy2HrADA>
- Campbell, Lo, and MacKinlay: chapter 1
- Fama, E. F., 1991, "Efficient Capital Markets II", *Journal of Finance*
- Garleanu and Pedersen, 2018, "Efficiently Inefficient Markets for Assets and Asset Management", *Journal of Finance*

1. Course Overview

The Central Question

- **What should a stock be worth?** In a No Arbitrage world, its price today should equal the present value of future cash flows. That present value requires a *discount rate* — the rate at which future dollars are translated into today's dollars, reflecting the riskiness of those cash flows
- **The discount rate varies with risk.** The object that captures this discounting is the **Stochastic Discount Factor (SDF)**, m : it is high in bad times (extra payoffs are most valuable) and low in good times. Every asset pricing model is a claim about m . The unifying condition is

$$p = E[mx]$$

where x is the next-period payoff, e.g. its dividend. Dividing each side by the price

$$1 = E[mR],$$

which implies that risk-adjusted returns are expected to be constant. In other words, no variable should predict risk-adjusted returns

- Lack of predictability of (risk-adjusted) returns is what you would expect in an efficient market

- However, **return predictability is present, in two forms**
 1. **Cross-sectionally** (the focus of this course): value stocks outperform growth stocks, past winners outperform past losers, profitable firms outperform unprofitable ones, etc. This can reflect either
 - (a) *risk premium* — these characteristics proxy for exposure to systematic risk, so the return is fair compensation (*no mispricing*)
 - (b) *mispricing* — these stocks are systematically misevaluated by investors
 2. **The same ambiguity arises in the time series**: a high price-dividend ratio predicts low future market returns, which can reflect a temporarily low discount rate (rational) or an overpriced market (mispricing)
- The first generation of researchers — after observing CAPM failures — defaulted to mispricing. Cochrane's warning: **predictability can reflect discount-rate variation, not necessarily mispricing**
- **The Joint Hypothesis Problem (JHP)**. Any test of efficiency simultaneously tests a model of m . A rejection may indict the market *or* the model. Fama (1991): “Efficiency and equilibrium are joined at the hip.” This is a *permanent structural constraint*, not a nuisance.

- **Cochrane's lesson** (AFA Presidential Address, 2011): “We have a lot of alphas (...the factor zoo). We have almost no understanding of why discount rates vary. Fix the second problem first.” Before concluding mispricing, exhaust the space of plausible SDFs.

Five topics around the EMH, Five Modules

The course follows a single chain of reasoning. **Each topic links to the next:**

- 1. Definition of the EMH. Fundamental issue in testing** *Module 1*
Any test requires specifying *correct* prices, i.e. a model of m . Every factor model is a different guess at m ; failures may reflect a misspecified model, not a mispriced market. The JHP means the EMH is never tested in isolation.
- 2. How do we test for deviations?** *Module 2*
Econometric approaches to asset pricing tests.
- 3. What are the deviations?** *Module 3*
Starting from a core set of deviations (value, momentum, profitability, low beta), the factor zoo has grown to 400+ published anomalies. There's an ongoing debate on how robust these are
- 4. Why are the deviations there?** *Module 4*
Three paradigms: *risk* (expand m , some anomalies dissolve), *mispricing* (investor errors), **institutional demand** (inelastic markets, non-fundamental flows)

5. **Why are they not arbitrated away?**

Module 5

Capital flows *away* from the best opportunities at the worst times (Shleifer and Vishny, 1997). Noise-trader risk (DSSW, 1990), liquidity spirals (Brunnermeier and Pedersen, 2009), and slow-moving capital (Duffie, 2010) allow deviations to persist

2. The Efficient Market Hypothesis (EMH)

Historical Perspective

- **Bachelier (1900):** first formal treatment of price randomness — efficiency of the Paris Bourse
- **Samuelson (1965):** in an informationally efficient market, *price changes should be unpredictable*
- **Fama (1970):** a market is efficient if *prices fully reflect all available information*
- **Malkiel (1992):** “...A capital market is said to be efficient if it *fully and correctly reflects all relevant information* in determining security prices. Formally, the market is said to be efficient with respect to some information set... if *security prices would be unaffected by revealing that information to all participants*. Moreover, efficiency with respect to an information set... implies that it is *impossible to make economic profits* by trading on the basis of [that information]...”

Three Ideas

1. **Prices fully reflect all available information** — suggestive but *empirically useless*
2. **Prices would not move if information were revealed to all participants** (because they have already exploited it) — a useful thought experiment, but still *not directly testable*
3. **Abnormal profits cannot be made by trading on the information** — the *truly operational* concept; this is what we test

Strategies for Testing Efficiency

The third idea suggests two empirical strategies:

1. **Observe profits of professional market participants.** If they earn superior risk-adjusted returns, markets are not efficient (e.g. mutual fund managers).
 - *Problem:* the information they use is unobservable
 - One may conclude efficiency when the truth is lack of skill
 - Alternatively, per **Berk and Green (2004)**: skill exists, but decreasing returns to scale plus managers' market power drive after-fee abnormal profits to zero
2. **Evaluate hypothetical strategies on an explicitly specified information set.** Do they earn superior returns?
 - Requires specifying: (i) the information set and (ii) a model for risk
 - Trading costs must be accounted for

To implement strategy 2, one must specify the information set. Different choices yield different *forms of efficiency* (Roberts, 1967):

1. **Weak-form Efficiency:** the information set contains only the *history of prices or returns*
2. **Semistrong-form Efficiency:** the information set contains *all publicly available information*
3. **Strong-form Efficiency:** the information set contains *all information known to any market participant*, including private information

Impossibility of Perfect Efficiency

- **Grossman and Stiglitz (1980):** if information is costly to gather, the market *cannot* be perfectly efficient
- Consider an investment firm that pays millions to set up a research division. It will only do so if it can profit from trading on that research
- Therefore, *some degree of inefficiency must exist* to compensate informed investors for their information costs
- If no profit were possible, no one would gather information — and the market would break down
- **Implication:** the empirically relevant question is not *whether* markets are inefficient, but *how much* — and whether the degree of inefficiency exceeds reasonable transaction costs

Abnormal Returns

- A key ingredient in any efficiency test is the definition of *superior*, or **abnormal**, returns:

$$\text{Abnormal Return} = \text{Realized Return} - \text{Normal Return}$$

- **Normal returns** represent fair compensation for risk and require specifying a model (e.g. CAPM, APT):

$$AR_{t+1} = R_{t+1} - E^M(R_{t+1})$$

where the superscript M denotes model-dependence

- **The EMH as a testable hypothesis:**

$$H_0 : E(AR_{t+1} \mid I_t) = 0$$

- If abnormal returns are predictable using information in I_t , market efficiency is rejected — *jointly with the model for risk*

Joint-Hypothesis Problem (JHP)

- **Fama (1970, 1991):** any empirical test of the EMH contains a *joint hypothesis*:
 1. Markets are efficient
 2. The model for risk (i.e. for normal returns) is correctly specified
- A rejection therefore *cannot* be unambiguously attributed to market inefficiency — it may reflect model misspecification
- **Consequence:** market efficiency, strictly speaking, can never be definitively rejected
- **Framing the rational vs. behavioral debate:** the JHP is precisely what keeps the debate alive. A *rational* researcher interprets any anomaly as evidence of a missing risk factor; a *behavioral* researcher interprets it as evidence of investor error

3. Efficient Capital Markets?

- **Fama (1970), Efficient Capital Markets I:** EMH holds unambiguously — the evidence strongly supports efficiency
- **Fama and French (1991):** first major negative results for CAPM and evidence of return predictability force a reassessment
- **The rational response:** these are not true rejections of EMH, but *expressions of the JHP* — we need better models of risk, not weaker notions of efficiency
- **Key insight:** despite the JHP making empirical conclusions ambiguous, the exercise is not irrelevant — we have learned an enormous amount about the properties of asset prices

1. Tests for Predictability

- (a) *Time-Series*: can past returns or macro variables predict future index returns?
- (b) *Cross-Sectional*: can stock characteristics predict the cross-section of returns?
 - Cross-sectional tests are one of the **central topics of this course**
 - The JHP is especially severe here: any characteristic-based anomaly may simply proxy for a missing risk factor

2. Event Studies (tests of semi-strong efficiency)

- Measure returns around the release of public information
- Short-window studies require minimal risk adjustment $\Rightarrow$ *cleanest available test of EMH*

3. Tests for Private Information (tests of strong efficiency)

- (a) Mutual and hedge fund performance

(b) Insider trading

- JHP applies, but *strong-form efficiency is not expected to hold* given Grossman–Stiglitz: informed traders must earn a return on their information costs

Quick Review of Results

1. Tests of Predictability

a) Time-Series Predictability

- Goal: test $E(AR_{t+1} | I_t) = 0$ for $t = 1, \dots, T$, where I_t is past returns or other public information

Results:

- **Short horizons** (daily, weekly, monthly):
 - Near-zero predictability for market indices; RW hypothesis holds
 - Some predictability for small stocks via infrequent trading (positive cross-autocorrelation, Lo–MacKinlay 1990) and *bid-ask bounce* (negative autocorrelation)
 - Does not survive reasonable transaction costs; not evidence of irrational pricing
- **Long horizons** (multi-year):
 - Negative return autocorrelation over 3–5 years (Fama–French, 1988) and long-horizon index returns are predictable by valuation ratios (D/P, E/P, etc.; Fama–French 1989)

- Two competing explanations:
 1. *Irrational pricing*: mean-reverting sentiment, $p_t = p^* + s_t$, where s_t is slow-moving (Summers, 1986). Excitement is followed by bursting bubbles and pessimism
 2. *Rational pricing*: time-varying risk aversion, linked to the business cycle, generates time-varying expected returns; a positive shock to expected returns lowers the price today, implying negative returns followed by higher future returns
- **Clear manifestation of the JHP**: both stories fit the facts

b) Cross-Sectional Predictability

- Test $E(AR_{t+1}^i | I_t) = 0$ for $i = 1, \dots, N$: do stock characteristics (size, B/M, past returns, volatility, profitability...) predict *differential* returns?
- Under CAPM, in the time-series regression

$$R_{t+1}^i - R_f = \alpha_i + \beta_i(R_{t+1}^m - R_f) + \varepsilon_{t+1}^i$$

the null is:

$$H_0 : \alpha_i = 0 \quad i = 1, \dots, N$$

- This is **indistinguishable from a test of CAPM** itself — the JHP at its starkest

2. Event Studies

- **Fama, Fisher, Jensen, and Roll (1969):** seminal paper on event studies
- Event studies measure average returns around the release of public information, over short (days) or long (up to one year) windows
- **Short-run studies** (daily): given average daily market return of $\approx 0.04\%$, risk adjustment is negligible $\Rightarrow$ *cleanest test of market efficiency*
- **Long-run studies:** suffer from JHP and statistical difficulties (cross-sectional return correlations)
- **Fama et al. (1969):** information is fully incorporated at the time of release — evidence *for* efficiency
- **Post-Earnings Announcement Drift** (Bernard–Thomas, 1990):
 - Positive earnings surprises trigger *positive price drift*, and vice versa
 - Consistent with *underreaction* to information

- PEAD has diminished in recent years — markets have become more efficient with respect to this signal
- Event studies return when we discuss *index addition/deletion effects*

3. Tests for Private Information

a) Insider Trading

- Prices begin rising several days *before* positive announcements (and vice versa)
- Not fully anticipated — prices do not reach post-announcement levels until the event
- **Evidence against strong-form efficiency**

b) Professional Portfolio Managers

- **Question:** do managers generate abnormal returns using their private information?
- **Approach:** test $H_0 : \alpha_i = 0$ in:

$$E(R_{t+1}^i) - R_f = \alpha_i + \beta_i^1 (E(R_{t+1}^m) - R_f) + \beta_i^2 F_2 + \beta_i^3 F_3 + \dots$$

- **Note:** $\alpha > 0$ implies skill *and* rejects strong efficiency — but the JHP applies: what is the right model for risk?

Results:

- **Jensen (1968):** net-of-fee returns are $\approx 1\%$ below benchmark; gross returns ≈ 0 — no evidence of private information
- **Ippolito (1989):** $+0.83\%$ above benchmark — but this disappears with a multifactor model (adding size and B/M)

- **Carhart (1997):** zero outperformance once momentum is controlled for — managers exploit existing anomalies, not private information
- **Passive investing implications:** these findings spurred the rise of the passive fund industry
- **Recent evidence of heterogeneous skill:** Kacperczyk–Sialm–Zheng (2005, 2006), Cremers–Petajisto (2009), Kosowski–Timmermann–Wermers–White (2006) find that *some* managers persistently outperform

4. Efficiently Inefficient

Efficient Market for Asset Managers

- **Berk and Green (2004, JPE):** a neoclassical model of active fund management with two key features:
 1. Investors *learn* about manager skill from past performance
 2. *Decreasing returns to fund scale*
- Manager skill is scarce; investors are abundant $\Rightarrow$ managers hold *monopoly power* over fees
- **Equilibrium outcome:**
 - Managers generate *positive before-fee alphas*
 - But *after-fee alphas are zero* — managers extract all rents
- **Flow dynamics:** rational investors chase past performance, but flows do not predict future performance **out-of-sample** because after-fee alphas are zero in expectation

- **In sample**, luck inflates perceived skill above true skill $\Rightarrow$ manager and investors alike overestimate skill $\Rightarrow$ active portfolio is oversized relative to true ability $\Rightarrow$ decreasing returns to scale harm performance $\Rightarrow$ realized returns are negative on average after high-flow periods.

In-sample predictability

- **Efficiency interpretation:** performance is equalized across managers after fees and is zero — an *efficient allocation to asset managers*

Friction in the Search for Asset Managers

- **Garleanu and Pedersen (2018, JF):** extend Grossman–Stiglitz by adding a *search cost* c borne by investors to find skilled managers (due diligence), alongside the *information cost* k borne by managers
- **Equilibrium conditions:**
 - Marginal investor indifferent between passive investing and paying c to find a skilled manager
⇒ **after-fee alpha** $= c > 0$
 - Marginal manager breaks even: fee revenue covers information cost k
 - These conditions jointly pin down the equilibrium level of price inefficiency η — the *efficiently inefficient* fixed point
- **Key implication:** observed outperformance is not “free money” — it is the equilibrium compensation for search frictions

Contrasts with benchmarks:

- **vs. Berk–Green (2004):** competitive capital flows drive after-fee alpha to *zero*; here, after-fee alpha is *positive* due to search frictions
- **vs. strong EMH:** no manager outperforms at all

Predictions consistent with stylized facts:

1. Some managers *persistently outperform* after fees (Kosowski et al. 2006) — contradicts both Berk–Green and strong EMH
2. Investors who bear higher search costs are more likely to access outperforming managers (Evans–Fahlenbrach 2012: institutional share classes outperform)
3. *Larger* investors outperform smaller ones (Gerakos–Linnainmaa–Morse 2016) — consistent with a fixed search cost c affordable only above a minimum AUM threshold
4. Anomalies *persist longer* in assets managed by specialists for whom search costs are higher (private equity, convertible bonds)
5. Anomalies are *larger in more complex markets* (equity vs. bonds) where fewer managers bear the information cost k
6. *Higher fees* in markets with more mispricing (hedge funds vs. mutual funds) — managers capture surplus where returns are higher

5. Market Efficiency and Learning

Setting aside the JHP, consider risk-neutral investors. The statement “prices fully reflect all available information” has different implications depending on the model:

1. **Rational Expectations:** investors *know* all parameters of the cash-flow model
 - Implication: returns are not predictable *in sample or out of sample*
 - The econometrician knows no more than investors; out-of-sample testing loses power without gaining interpretability (Cochrane 2008; Campbell–Thompson 2008)
2. **Learning (Bayesian):** investors *learn* about parameters using information up to time t
 - Implication: returns are predictable *in sample* but *not out of sample*
 - The econometrician uses data up to $T > t$, so has had more observations to learn the parameters — but this does not mean investors were irrational
 - **Key implication:** out-of-sample tests are essential not only to control for data mining and p -hacking, but also to distinguish genuine inefficiency from the artefact of *learning*

A Simple Example

Based on **Lewellen and Shanken (2002)**.

A stock pays dividends with unknown mean $\bar{d}$:

$$d_t = \bar{d} + \varepsilon_t$$

1. **Diffuse prior** (investors know nothing about $\bar{d}$):

- Investors estimate $\bar{d}_t = \frac{1}{t} \sum_{s=1}^t d_s$
- Consider a large realization of d_t , purely due to noise, which raises $\hat{\bar{d}}_t$ and the stock price p_t
- In the future, prices must *correct downwards* because the noise dissipates $\Rightarrow$ **in-sample mean reversion**
- But investors at time t used all available information efficiently; from their perspective, returns were *not* predictable

2. Strong prior on low dividends:

- A high d_t leads to *underreaction* as investors discount the signal relative to their prior
- Prices *drift upwards* slowly as beliefs update $\Rightarrow$ **in-sample momentum**
- Eventually, investors learn about $\bar{d}$ and predictability disappears
- **Conclusion:** in-sample predictability can arise under *full rationality*. Investors were using all available information optimally — they simply could not have known what the econometrician knows after the fact

Martin and Nagel (2022) – Setup and Main Result

- **Motivation:** big data has expanded the number of potential predictors J (social media, satellite images, etc.) to the point where J is comparable to the number of stocks N
- **Setup:** investors learn the unknown mapping from characteristics to expected cash flows:

$$\Delta y_t = Xg + e_t$$

where X is $N \times J$, $e_t \sim \mathcal{N}(0, I)$, and the unknown coefficient $g \sim \mathcal{N}(0, \frac{\theta}{J}I)$ is estimated via Bayes' rule

- **The J/N problem:** when J/N is large, the data cannot separately identify all J coefficients $\Rightarrow$ investors must *shrink* estimates toward zero (Bayesian regularization)
- **Main result:** in-sample cross-sectional predictability — the “factor zoo” (Cochrane, 2011) — arises in equilibrium *under full rationality and market efficiency*, with no mispricing or data snooping required
- Since investors use all available information optimally in real time, **out-of-sample predictability is zero**

Martin and Nagel (2022) – Two Sources of In-Sample Predictability

- **Source 1: Shrinkage-induced underreaction**

- Suppose true $g = 1$ for profitability; investors estimate $\hat{g} = 0.6$ due to parameter uncertainty
- Prices underweight profitability $\Rightarrow$ high-profitability stocks are *systematically underpriced*
- When cash flows are realized, prices correct $\Rightarrow$ above-normal returns in-sample for profitable companies
- The econometrician, using the full sample, detects this pattern — but real-time investors could not have done better

- **Source 2: Noise-fitting**

- At each t , past cash-flow growth noise $\bar{e}_t$ happens, by chance, to line up with *some* combination of the J characteristics. Investors incorporate this spurious pattern into prices
- When next-period cash flows arrive independently of past noise, prices correct $\Rightarrow$ returns are predictable *in sample*

- At $t + 1$, a *different* spurious pattern forms with a *different* combination of characteristics. The specific pattern changes every period — but the tendency to fit noise is *always present* when J is large
- Crucially, this does not make any *single* characteristic (e.g. B/M) persistently predict returns. Rather, it inflates in-sample predictability *across the board*: the econometrician's joint test finds that returns are correlated with *some* linear combination of characteristics in every period, even though no single one is consistently responsible
- **Example.** Suppose $J = 100$ characteristics are available. Over periods $1, \dots, t$, past noise happened to align with a particular weighted mix of, say, size, past returns, and investment. Prices reflect this mix, so when dividends arrive and the noise does not repeat, that *specific mix* earns abnormal returns in the econometrician's sample. Over periods $t + 1, \dots, 2t$, a *completely different* mix — say, volatility and accruals — lines up with noise, and again corrects. The econometrician, pooling all T periods, runs a joint test and finds that returns are *significantly* correlated with the characteristics: every subperiod contributed a different spurious pattern, but together they make the joint test reject strongly. No single characteristic drives the result; the effect is spread across the 100 variables

- **Why J/N matters:** larger $J \Rightarrow$ more shrinkage (Source 1) and more opportunities for noise-fitting (Source 2) $\Rightarrow$ in-sample predictability grows with J/N , even as out-of-sample predictability stays zero
- **Bottom line:** in-sample predictability *cannot* distinguish rational learning from genuine mispricing. Only out-of-sample evidence can.

An Example: Pension Plan Funding and Stock Market Efficiency

- **Franzoni and Marín (2006, JF):** firms with severely *underfunded* defined-benefit pension plans earn significantly lower risk-adjusted returns for up to five years — an apparent violation of semi-strong efficiency
- **The mechanism:** under SFAS No. 87, pension liabilities were disclosed only in *footnotes*, not on the balance sheet $\Rightarrow$ investors failed to process the information
- **Fundamental channels:**
 - Amortisation of pension losses flows through the income statement, depressing future earnings
 - Mandatory ERISA contributions drain cash from operations
- When these effects materialise, *negative earnings surprises* trigger downward price adjustment — consistent with *underreaction to complex, obscure information*
- **Portfolio evidence:** long-short strategy on pension funding ratio generates:

- CAPM alpha $\approx -1\%$ /month for the most underfunded decile
 - Four-factor (Carhart) alpha $\approx -0.6\%$ /month (significant)
 - Dynamic strategy: 1.51% /month ($\sim 18\%$ annualised), Sharpe ratio 0.26 over 1989–2004
- **Learning interpretation:** consistent with Martin–Nagel Source 1 — investors underweighted an opaque predictor; as information became more salient (SFAS No. 158 in 2006 moved liabilities to the face of the balance sheet), the anomaly substantially diminished

Out-of-Sample Evidence: Do Anomalies Survive Publication?

- **McLean and Pontiff (2016, JF)**: study 97 cross-sectional return predictors and decompose returns across three windows:
 - **In-sample**: average long-short return of 0.58%/month
 - **Post-sample, pre-publication**: returns decline by 26% — the authors impute this reduction to data-mining bias. It is also consistent with the learning explanation (Martin and Nagel, 2022)
 - **Post-publication**: returns decline by 58% — anomalies retain only $\approx 42\%$ of original returns
- The *additional* 32 percentage point decline **after** publication reflects **investor learning from academic research**: post-publication anomaly stocks show increased trading volume, short interest, and return variance
- **Interpretation via learning**: in-sample predictability is consistent with Martin–Nagel Source 1. As accounting reforms and publication made signals more salient, learning accelerated and anomalies weakened
- **But anomalies do not disappear entirely**: anomalies with residual post-publication returns display arbitrage costs — idiosyncratic risk is strong for these stocks

Conclusion

In-sample evidence of return predictability may reflect:

1. **Market inefficiency** — irrationality or limited rationality (*behavioral finance*) combined with limits to arbitrage
2. **Model misspecification** — wrong risk model (*rational finance* à la Fama–French); the JHP means we can never rule this out
3. **Parameter uncertainty and Bayesian learning** — rational investors using all available information efficiently, generating in-sample predictability the econometrician detects only in hindsight

The critical role of out-of-sample testing:

- Distinguishes hypothesis 3 (learning) from hypotheses 1 and 2 (inefficiency or wrong model)
- Controls for data mining and p -hacking
- Is the standard to which the field has converged following McLean–Pontiff (2016), Harvey–Liu–Zhu (2016), and Hou–Xue–Zhang (2020)