



Stockholms
universitet

Downside Risk in Credit Markets: A Lower Partial Moments Extension of the CAPM

Oscar Sjöstrand

Masteruppsats 2026:9
Matematisk statistik
Juni 2026

www.math.su.se

Matematisk statistik
Matematiska institutionen
Stockholms universitet
106 91 Stockholm

Downside Risk in Credit Markets: A Lower Partial Moments Extension of the CAPM

Oscar Sjöstrand*

June 2026

Abstract

This thesis investigates whether a lower partial moment extension of the Capital Asset Pricing Model (LPM-CAPM) is theoretically and empirically applicable to credit markets. Credit instruments differ from equities through their asymmetric payoff structure: upside potential is limited, while downside losses may be substantial in adverse states. This motivates the use of downside risk measures rather than variance-based. Building on utility theory, lower partial moments, and downside beta pricing, the thesis derives a credit-oriented version of the LPM-CAPM in which expected excess returns are related to exposure to systematic downside market risk.

The empirical analysis is conducted using weekly excess returns on 30 total return bond indices from January 2000 to April 2026. A two-pass cross-sectional regression framework, is used to compare the standard CAPM with several LPM-CAPM specifications across different choices of the target return τ and downside sensitivity parameter α . To account for serial dependence and uncertainty from estimated betas, inference is based on a stationary bootstrap procedure.

The results provide partial support for the relevance of downside risk in credit markets. Several LPM-CAPM specifications are not rejected by the hypothesis tests, and some specifications achieve slightly higher average cross-sectional explanatory power than the standard CAPM. However, the improvement is modest, and the LPM-CAPM does not clearly dominate the CAPM. The CAPM remains a competitive benchmark, particularly under bootstrap inference. The findings therefore suggest that downside beta contains useful information for explaining credit-index returns, but that a single-factor downside-risk model is not sufficient to fully capture the cross-section of credit excess returns.

*Postal address: Mathematical Statistics, Stockholm University, SE-106 91, Sweden.
E-mail: oscar.sjostrand99@gmail.com. Supervisor: Daniel Ahlberg.

Acknowledgments

Although all of the ideas, original drafts and concepts are written by me, it is important to acknowledge I have used AI-based tool to assist in formulation, proof reading and generating some of the code used to achieve the results presented.

Contents

1	Introduction	7
2	Theory	10
2.1	Utility Theory and Risk Preferences	10
2.1.1	Dominance	12
2.1.2	Mean-Variance, a Special Case	13
2.2	CAPM	14
2.3	Credit Markets, Returns, and Downside Risk	19
2.3.1	Payoff Structure	19
2.3.2	Returns and Excess Returns	19
2.3.3	Total Return Indices (TRIV)	20
2.3.4	Applicability of the CAPM	21
2.4	Lower Partial Moments	22
2.5	A Credit Interpretation of the LPM-CAPM	24
2.5.1	Downside Covariance Pricing	26
2.5.2	Equilibrium Downside Beta Pricing	27
2.5.3	Connection to Repayment Risk	28
2.5.4	Pricing Implications	29
2.5.5	The choice of τ and α	29
2.5.6	From Single-Period Credit Claims to TRIV Returns	31
3	Empirical Study	33
3.1	Data	34
3.2	Dependence Structure and Inference	35
3.3	Empirical Methodology	36
3.3.1	First Pass: Beta Estimation	36
3.3.2	Second Pass: Fama–MacBeth Regressions	37
3.3.3	Dependence, Generated Regressors, and Limitations of Cross- Sectional Regression	38
3.3.4	Stationary Bootstrap	39
3.3.5	Hypothesis Tests	40
3.3.6	Model Evaluation	42

4	Results	44
4.1	Hypothesis Tests	44
4.2	Model Evaluation	46
4.3	Summary of Findings	47
5	Discussion	48
6	Appendix	49
	Bibliography	54

1 Introduction

The bond and credit markets are important parts of the world economy. They provide financing for governments and corporations and constitute essential components of institutional and private investment portfolios. In contrast to equities, credit instruments have a more asymmetric payoff structure: the upside is typically limited by promised coupon and principal payments, while downside losses may be substantial in the event of default or credit deterioration. This asymmetry suggests that downside risk is particularly important when studying credit returns.

Modern portfolio theory originates with the work of [Markowitz \(1952\)](#), who proposed variance as a measure of portfolio risk. In the mean–variance framework, investors choose portfolios by trading off expected return against return variance. This idea laid the foundation for equilibrium asset pricing models, most notably the Capital Asset Pricing Model (CAPM) introduced by [Sharpe \(1964\)](#). The CAPM implies that the expected excess return of an asset is proportional to its exposure to systematic market risk, measured by the asset’s market beta. Despite its simplicity and clear economic interpretation, the CAPM has been widely criticized for its limited empirical support. For example *The Capital Asset Pricing Model: Theory and Evidence* by [Fama and French \(2004\)](#) review the empirical failures of the CAPM. Nevertheless, it remains an important benchmark in both academic research and practical applications such as cost of capital estimation and performance evaluation.

A central limitation of the CAPM is its reliance on variance as the relevant measure of risk. Variance treats positive and negative deviations from the mean symmetrically, even though assumptions of risk aversion imply that investors are more concerned with losses, unexpected positive returns have the same negative impact. This limitation is especially relevant in credit markets, where return distributions are shaped by limited upside, default risk, spread widening, and other adverse market conditions. A variance-based model may therefore fail to capture the type of risk that is economically most important for credit investors.

Lower Partial Moments (LPMs) provide a natural alternative risk measure. Building on the mean-risk framework of [Fishburn \(1977\)](#), downside risk can be measured as

$$LPM_{\alpha}(\tau) = \mathbb{E} \left[(\tau - X)_{+}^{\alpha} \right],$$

where τ is a target return and $\alpha > 0$ determines the sensitivity to shortfalls below

the target. Unlike variance, the LPM only penalizes outcomes below the benchmark τ , making it well suited for settings where downside outcomes are more relevant than upside variation. This feature makes the LPM framework particularly appealing for credit markets.

Several authors have extended asset pricing theory beyond the mean–variance setting by incorporating downside risk. Early contributions by [Bawa \(1975\)](#) use lower partial moments in portfolio choice and derive pricing implications under asymmetric risk preferences. [Harlow and Rao \(1989\)](#) develop this line of research further by introducing downside beta and empirical methods for testing downside-risk asset pricing models. More recently, [Satchell \(2001\)](#) shows that, under generalized utility functions, equilibrium expected returns may be characterized by covariance with a downside risk factor rather than variance. These results suggest that CAPM-type pricing relations can be derived when risk is measured asymmetrically.

The main question of this thesis is whether an LPM-based extension of the CAPM is theoretically and empirically applicable to credit markets. More specifically, the thesis investigates whether exposure to systematic downside market risk helps explain the cross-section of corporate bond index excess returns, and whether such a downside-risk model improves upon the standard CAPM.

The contribution of the thesis is twofold. First, it extends the equity-oriented LPM-CAPM to a credit-oriented interpretation, by linking downside beta to the asymmetric payoff structure of credit instruments. Second, it evaluates the empirical relevance of this framework using credit index returns. In this way, the thesis connects the theoretical motivation for downside risk with an empirical asset pricing test in fixed income markets.

The empirical analysis is based on weekly excess returns for 30 total return bond indices from January 2000 to April 2026. The indices cover corporate bonds across different rating and maturity segments. A broad corporate bond index is used as the market proxy, and excess returns are computed relative to a short-term risk-free rate. The standard CAPM and several LPM-CAPM specifications are compared using a two-pass cross-sectional regression procedure based on the method presented in *Risk, Return, and Equilibrium: Empirical Tests* by [Fama and MacBeth \(1973\)](#).

In the first pass, asset-specific betas are estimated. In the second pass, cross-sectional regressions are used to test whether these betas explain average excess returns. Since estimated betas introduce additional uncertainty and financial returns may exhibit serial dependence, inference is conducted using a stationary bootstrap

procedure.

The empirical results are mixed. Some LPM-CAPM specifications are not rejected by the hypothesis tests and produce slightly higher average cross-sectional explanatory power than the standard CAPM. However, these improvements are small and should be interpreted cautiously, since the analysis does not establish that the LPM-CAPM has statistically significant incremental explanatory power beyond the CAPM. The CAPM therefore remains a competitive benchmark, particularly under stationary bootstrap inference. Overall, the results are consistent with the view that downside risk may be relevant for credit-index returns, but they do not provide strong evidence that a single-factor downside-risk model clearly improves upon, or replaces, the standard CAPM.

The remainder of the thesis is organized as follows. Section 2 presents the theoretical framework, including utility theory, the CAPM, lower partial moments, and the derivation of the LPM-CAPM in a credit setting. Section 3 describes the data and empirical methodology. Section 4 reports the empirical results. Section 5 discusses the findings and concludes the thesis.

2 Theory

Let us begin by clarifying some of the key terminology used in this paper. A portfolio is a collection of financial assets held by an investor, and the portfolio weights describe the proportion of the investors total wealth allocated to each asset. We distinguish between two broad categories of assets: equities and credits. Equities represent ownership claims, such as stocks, and do not provide guaranteed payments; their returns arise primarily from changes in market price and, when applicable, dividend distributions. Credits, in contrast, are debt instruments in which the investor lends capital to a borrower. These instruments typically promise a stream of coupon payments (interest) as well as the repayment of principal at maturity. As a result, while equities have theoretically unlimited upside and uncertain payoffs, credit instruments have a more structured payoff profile with a limited upside but exposure to downside risk through default. Risk-free return is in our case an asset with a guaranteed return. There are a multitude of different proxies for a risk-free return such as investing in short term treasury bonds.

2.1 Utility Theory and Risk Preferences

Expected wealth gain alone is generally not sufficient to rank investment opportunities. Two portfolios may have the same expected return while exposing the investor to very different distributions of outcomes. In particular, one portfolio may involve a substantial probability of large losses and spectacular gains whereas another delivers a more stable return with little downside. A theory of portfolio choice therefore requires a framework for comparing risky prospects, not only their expected monetary values. Utility theory provides such a framework by allowing preferences to be defined over final wealth and, more generally, over risky lotteries see [Varian \(1992\)](#) chapter 11.

Following [Varian \(1992\)](#) chapter 7, let $\mathcal{W} \subseteq \mathbb{R}$ denote the set of attainable wealth levels, where wealth should be interpreted as the investor's final (terminal) monetary wealth, and let $\succeq$ be a preference relation on $\mathcal{W}$. We assume that preferences are monotone, meaning that higher wealth is always weakly preferred to lower wealth. Consider the utility function

$$u : \mathcal{W} \rightarrow \mathbb{R}$$

represents the preference ordering if

$$W_1 \succeq W_2 \iff u(W_1) \geq u(W_2).$$

Thus, utility is not primarily a measure of wealth itself, but a device for ranking wealth outcomes according to the investor's preferences. Importantly, the shape of the utility function describes how the investor values additional units of wealth. For example, if the utility function is increasing, higher wealth levels are preferred to lower wealth levels. If the utility function is concave, the marginal utility of wealth is decreasing, meaning that an additional unit of wealth is valued less when the investor is already wealthier. In this way, the utility function summarizes both the ranking of wealth levels and how the value of additional wealth changes with the level of wealth.

In portfolio problems, terminal wealth is typically random. Let $\tilde{W}$ denote a random final wealth level defined on some probability space. Under the von Neumann–Morgenstern axioms see [von Neumann and Morgenstern \(1944\)](#), preferences over risky prospects can be represented by expected utility, so that the investor evaluates $\tilde{W}$ according to

$$U(\tilde{W}) = \mathbb{E}[u(\tilde{W})].$$

If $\tilde{W}$ has cumulative distribution function F_W , this may be written as

$$U(\tilde{W}) = \int_{\mathcal{W}} u(w) dF_W(w),$$

and, when $\tilde{W}$ has density f_W , equivalently as

$$U(\tilde{W}) = \int_{\mathcal{W}} u(w) f_W(w) dw.$$

Thus, the portfolio choice problem can be formulated as the maximization of expected utility rather than expected wealth alone. This formulation reflects the assumption that investors behave in accordance with expected utility theory, a standard framework in both microeconomic theory and asset pricing (see, e.g., [Cochrane \(2005\)](#)).

2.1.1 Dominance

Before specializing to particular risk measures, it is useful to clarify what it means for one portfolio to dominate another. In the most general sense, dominance is a statement about preference rankings over entire payoff distributions.

Let $\tilde{R}_F$ and $\tilde{R}_G$ denote the random returns of portfolios F and G , let $W_0 > 0$ be the initial wealth and the terminal wealth be given by

$$\tilde{W}_F = W_0(1 + \tilde{R}_F), \quad \tilde{W}_G = W_0(1 + \tilde{R}_G).$$

Let F and G be the cumulative distribution functions of $\tilde{W}_F$ and $\tilde{W}_G$. Portfolio F is said to *first-order stochastically dominate* see [Fishburn \(1977\)](#) portfolio G if

$$F(x) \leq G(x) \quad \text{for all } x,$$

with strict inequality for some x . In that case, every investor with an increasing utility function prefers F to G . If we include the notion of risk-aversion see [Varian \(1992\)](#) chapter 11.5 we get

$$\int_{-\infty}^x F(t) dt \leq \int_{-\infty}^x G(t) dt \quad \text{for all } x,$$

with strict inequality for some x , and $\mathbb{E}[\tilde{W}_F] \geq \mathbb{E}[\tilde{W}_G]$. Then F *second-order stochastically dominates* G , implying that every investor with an increasing and concave (risk-averse) utility function prefers F to G .

These dominance concepts provide a general preference-based ranking of risky prospects. In practice, however, portfolio theory often works with more structured *mean-risk* representations. Such a framework can be interpreted as approximations of stochastic dominance, where the full distribution is summarized by a reward functional $\mu(F)$ and a risk functional $\rho(F)$. Following the reasoning in [Fishburn \(1977\)](#), one may say that a portfolio F dominates portfolio G if it offers at least as much reward and no more risk, i.e.

$$\mu(F) \geq \mu(G), \quad \rho(F) \leq \rho(G), \tag{1}$$

with at least one strict inequality. Under appropriate choices of μ and ρ , this formulation is consistent with stochastic dominance and include mean–variance and

mean-LPM preferences as special cases.

2.1.2 Mean-Variance, a Special Case

Mean-variance analysis can be viewed as a simplification of the general framework presented in Section 2.1.1. Instead of considering the full distribution of outcomes, risk is summarized by the variance alone. This approximation is not valid in general, but becomes exact under specific assumptions on preferences or return distributions.

A notable exception arises under quadratic utility. Consider the utility function

$$u(W) = aW - \frac{b}{2}W^2,$$

where $a, b > 0$. For a random wealth $\tilde{W}$ with mean μ_W and variance σ_W^2 ,

$$\begin{aligned}\mathbb{E}[u(\tilde{W})] &= a\mathbb{E}[\tilde{W}] - \frac{b}{2}\mathbb{E}[\tilde{W}^2] \\ &= a\mu_W - \frac{b}{2}(\sigma_W^2 + \mu_W^2).\end{aligned}$$

Expected utility is therefore fully characterized by the first two moments of wealth. In this case, maximizing expected utility under quadratic preferences is equivalent to a mean-variance trade-off.

From the perspective of stochastic dominance, this implies that mean-variance preferences provide a tractable representation of a restricted class of utility functions. While second-order stochastic dominance compares entire distributions for all risk-averse investors, mean-variance analysis effectively approximates these preferences by summarizing risk through variance alone. This simplification is exact under quadratic utility, where investor preferences depend only on mean and variance. It also holds for certain classes of return distributions, such as elliptical distributions, where mean and variance fully describe the distribution and therefore suffice for comparing portfolios.

At the same time, quadratic utility has well-known limitations. Since

$$u'(W) = a - bW,$$

marginal utility eventually becomes negative for sufficiently large W . Quadratic utility is therefore only locally plausible. More importantly, mean-variance analysis treats positive and negative deviations from the mean symmetrically. This may be

restrictive in investment problems where downside outcomes are of primary concern.

Despite its limitations, the quadratic utility function and the mean–variance framework remain central tools in financial economics. Under the additional assumptions of frictionless markets, homogeneous beliefs, and the existence of a risk-free asset, this framework leads directly to the Capital Asset Pricing Model.

2.2 CAPM

The *Capital Asset Pricing Model* (CAPM), introduced by Sharpe (1964), is an asset-pricing model stating that the expected excess return of an asset, that is, the return in excess of the risk-free rate, is proportional to its systematic risk relative to the market portfolio. Hence, under the assumptions of the model, the CAPM determines the return that investors require for holding a given asset in equilibrium.

The assumptions of the model are:

- Quadratic utility (Mean–variance utility function),
- Risk-averse investors,
- Homogeneous expectations (investors share the same expectations about returns),
- Single period investment horizon (Investors make decisions over the same time horizon),
- Frictionless market (no taxes, fees or price-spreads),
- Existence of risk-free asset,
- All risky assets are tradable and perfectly divisible,
- Short selling is allowed,
- Competitive markets/ price-taking investors,
- Market equilibrium (total demand for risky assets equals the total supply).

These assumptions are highly idealized and are unlikely to hold exactly in real financial markets. Nevertheless, the CAPM remains important because it provides a simple and transparent benchmark for understanding the relation between risk and expected return. In particular, the model gives a clear economic interpretation

of systematic risk through market beta and shows why only risk that cannot be diversified away should be priced in equilibrium. For this reason, the CAPM is often used not as a literal description of markets, but as a reference model against which alternative asset-pricing models can be compared.

The standard formulation of the CAPM is

$$\mathbb{E}[R_i] - R_f = \beta_{iM}(\mathbb{E}[R_M] - R_f), \quad (2)$$

where R_i is the return on asset i , R_M is the return on the market portfolio, R_f is the risk-free return, and

$$\beta_{iM} = \frac{\text{Cov}(R_i, R_M)}{\text{Var}(R_M)} \quad (3)$$

measures the asset's systematic risk relative to the market.

The intuition behind the model is that each investor wants to maximize the utility of her risky portfolio. When only risky assets are available, the set of mean–variance efficient portfolios forms the efficient frontier. Since investors are assumed to have homogeneous expectations, they agree on the expected returns, variances, and covariances of all assets, and therefore face the same efficient frontier seen in Figure 1.

When a risk-free asset is introduced, investors can combine the risk-free asset with any risky portfolio. This generates straight lines in mean–variance space, all originating from the risk-free return. The optimal risky portfolio is the tangency portfolio, which gives the steepest such line and is represented by M in the figure. Under the assumptions of the CAPM, all investors face the same investment opportunities and therefore choose this same risky portfolio, differing only in how much they allocate between the risk-free asset and the tangency portfolio. In equilibrium, aggregate demand for risky assets must equal aggregate supply, so the tangency portfolio coincides with the market portfolio.

Let us formally derive CAPM taking inspiration from the derivation in [Cochrane \(2005\)](#) section 9.1.2. Let

- $R \in \mathbb{R}^n$ denote the vector of risky assets returns,
- $\mu = \mathbb{E}[R]$ expected returns,
- Σ the covariance matrix,

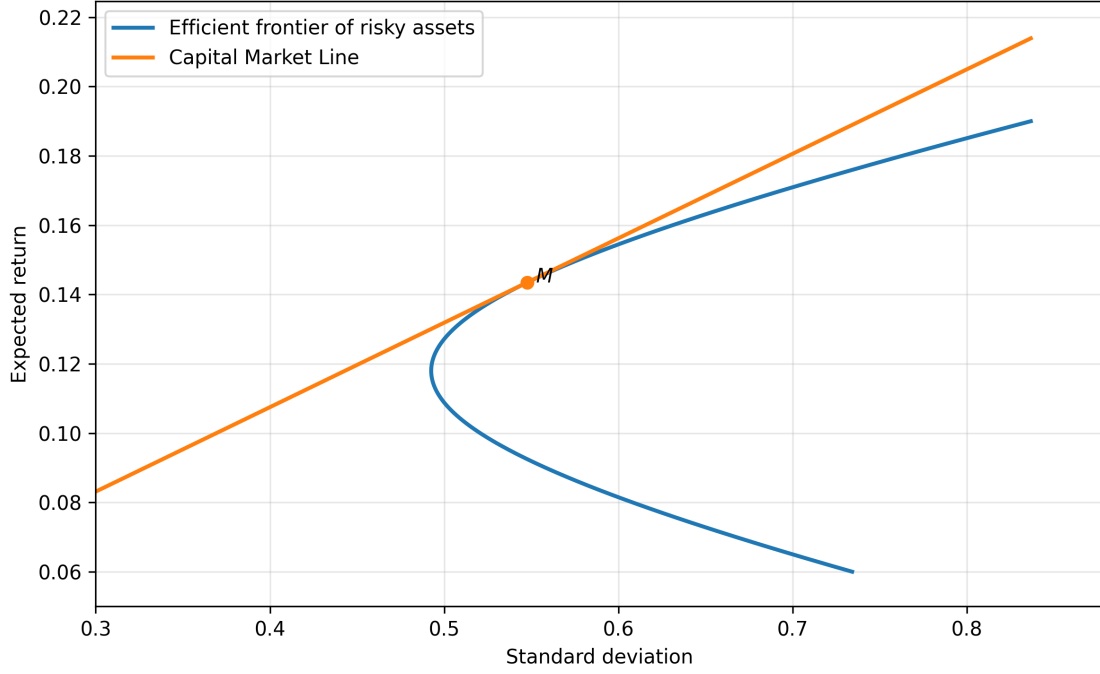


Figure 1: Mean–variance diagram illustrating the CAPM. The curved line is the efficient frontier of risky assets, and the straight line is the Capital Market Line. The tangency point M represents the optimal risky portfolio, which coincides with the market portfolio in equilibrium.

- R_f the risk free return.

Let $\omega \in \mathbb{R}^n$ denote the vector of portfolio weights invested in the risky assets. The return on the risky portfolio is then

$$R_p = \omega^T R.$$

If the investor can also invest in the risk-free asset, the relevant object is the portfolio excess return. Let

$$r = R - R_f \mathbf{1},$$

where $\mathbf{1}$ is an n -dimensional vector of ones. Then the expected excess return on the risky portfolio is

$$\mathbb{E}[\omega^T r] = \omega^T (\mu - R_f \mathbf{1}).$$

The variance of the risky portfolio return is

$$\text{Var}(\omega^T R) = \omega^T \Sigma \omega.$$

To see this explicitly, note that

$$\begin{aligned} \text{Var}(\omega^T R) &= \mathbb{E} \left[\left(\omega^T (R - \mu) \right)^2 \right] \\ &= \mathbb{E} \left[\omega^T (R - \mu) (R - \mu)^T \omega \right] \\ &= \omega^T \mathbb{E} \left[(R - \mu) (R - \mu)^T \right] \omega \\ &= \omega^T \Sigma \omega. \end{aligned}$$

Thus, the covariance matrix enters through the variance of the portfolio return. Since expected utility under quadratic preferences depends on mean and variance, the investor's portfolio choice can be represented as a mean–variance trade-off. We therefore write the optimization problem as

$$\max_{\omega} \quad \omega^T (\mu - R_f \mathbf{1}) - \frac{\gamma}{2} \omega^T \Sigma \omega,$$

where $\gamma > 0$ is the investor's coefficient of risk aversion.

For notational simplicity, we write $\mu - R_f$ for $\mu - R_f \mathbf{1}$. Differentiation gives us

$$\frac{\partial}{\partial \omega} \left(\omega^T (\mu - R_f) - \frac{\gamma}{2} \omega^T \Sigma \omega \right) = (\mu - R_f) - \gamma \Sigma \omega,$$

since Σ is symmetric. Setting the derivative to zero and rearranging gives us the optimal weights

$$\omega^* = \frac{1}{\gamma} \Sigma^{-1} (\mu - R_f). \tag{4}$$

This implies that, under the assumptions of the CAPM (quadratic utility or normally distributed returns, frictionless markets, homogeneous expectations, and short selling), every investor's optimal complete portfolio is a combination of the risk-free asset and a single risky portfolio. This is the two-fund separation theorem introduced in *Liquidity Preference as Behavior Towards Risk* by [Tobin \(1958\)](#): regardless of their level of risk aversion, all investors hold the same risky portfolio, differing only in how much they invest in it relative to the risk-free asset.

Remark 2.1. Equation (4) allows individual components of ω^* to be negative. Such negative positions correspond to short positions in the corresponding risky assets. This is a standard assumption in the unrestricted CAPM derivation, but it may be restrictive in credit markets.

Let ω^* denote the vector of weights of this common risky portfolio. From the mean–variance optimization problem, the optimal risky weights for an investor with risk aversion γ are given by

$$\omega^* = \frac{1}{\gamma} \Sigma^{-1}(\mu - R_f).$$

Thus, each investor’s demand for risky assets is proportional to the vector $\Sigma^{-1}(\mu - R_f)$.

In equilibrium, the aggregate demand for risky assets must equal their aggregate supply. Since all investors hold portfolios that are scalar multiples of the same vector, the market portfolio must also be proportional to this vector. Let ω_M denote the vector of market portfolio weights, defined by the relative market capitalizations of all risky assets. Then it follows that

$$\omega_M \propto \Sigma^{-1}(\mu - R_f).$$

This implies that there exists a scalar $\lambda > 0$ such that

$$\Sigma^{-1}(\mu - R_f) = \lambda \omega_M.$$

Multiplying both sides by Σ gives

$$\mu - R_f = \lambda \Sigma \omega_M.$$

The i th element of this expression yields

$$\mathbb{E}[R_i] - R_f = \lambda \text{Cov}(R_i, R_M),$$

where $R_M = \omega_M^T R$ denotes the return on the market portfolio.

Applying the same relation to the market portfolio itself gives

$$\mathbb{E}[R_M] - R_f = \lambda \text{Var}(R_M).$$

Solving for λ and substituting back yields

$$\mathbb{E}[R_i] - R_f = \frac{\text{Cov}(R_i, R_M)}{\text{Var}(R_M)} (\mathbb{E}[R_M] - R_f).$$

Defining

$$\beta_i = \frac{\text{Cov}(R_i, R_M)}{\text{Var}(R_M)},$$

we obtain the Capital Asset Pricing Model,

$$\mathbb{E}[R_i] - R_f = \beta_i (\mathbb{E}[R_M] - R_f).$$

2.3 Credit Markets, Returns, and Downside Risk

Credit markets differ fundamentally from equity markets in both payoff structure and risk characteristics. While equities offer unbounded upside, the maximum payoff of a credit instrument is typically capped by the contractual coupon and principal payments. At the same time, downside losses may be substantial if the borrower defaults or the market value of the instrument deteriorates. This asymmetry suggests that, when applying CAPM-style models to credit returns, a symmetric risk measure such as variance may be incomplete. A downside-risk framework is therefore more natural in this setting, since it focuses directly on losses or shortfalls relative to a relevant benchmark.

2.3.1 Payoff Structure

Consider a one-period credit instrument i with promised payoff $\bar{X}_i$. The realized payoff can be expressed as

$$X_i = q_i \bar{X}_i, \quad 0 < q_i \leq 1, \quad (5)$$

where q_i denotes the repayment fraction. A value of $q_i = 1$ corresponds to full repayment, while $q_i < 1$ reflects credit losses arising from default or restructuring. The randomness in X_i is therefore entirely driven by the stochastic nature of q_i .

2.3.2 Returns and Excess Returns

To connect the payoff structure of a credit instrument to the empirical analysis, we now introduce returns over time. Consider an investor who holds a credit instrument

for one period, from time $t - 1$ to time t . The return over this period consists of the change in the market price of the instrument and any coupon payment received during the holding period.

Let $P_{i,t}$ denote the market price of credit instrument i at time t , and let $C_{i,t}$ denote the coupon payment received during the period from $t - 1$ to t . The realized gross holding-period return is then given by

$$1 + R_{i,t} = \frac{P_{i,t} + C_{i,t}}{P_{i,t-1}}. \quad (6)$$

Here, $P_{i,t-1}$ is the price paid at the beginning of the period, while $P_{i,t} + C_{i,t}$ is the value of the investment at the end of the period, including both the new market price and any coupon income.

The corresponding excess return is defined as

$$r_{i,t}^e = R_{i,t} - R_{f,t}, \quad (7)$$

where $R_{f,t}$ is the risk-free return over the same period.

Excess returns are central in asset pricing because they measure the compensation for bearing risk relative to a risk-free benchmark. In the stylized credit setting, this compensation is linked to systematic variation in repayment risk q_i . In the empirical analysis, however, observed credit-index returns also reflect price changes, coupon income, and other market conditions.

2.3.3 Total Return Indices (TRIV)

In empirical applications, credit markets are often represented by total return indices values (TRIV), which measure the cumulative return from holding a diversified portfolio of credit instruments, including both price changes and reinvested coupon payments.

Formally, let I_t denote a total return index. Returns are computed as

$$R_t = \frac{I_t}{I_{t-1}} - 1. \quad (8)$$

TRIV indices provide a sufficient proxy for tradable credit portfolios. The advantages are the duration and easy access of the data. However, since they aggregate across many individual bonds, implicitly diversifying away individual default risk,

the indices preserving exposure to systematic credit risk and the asymmetry in returns.

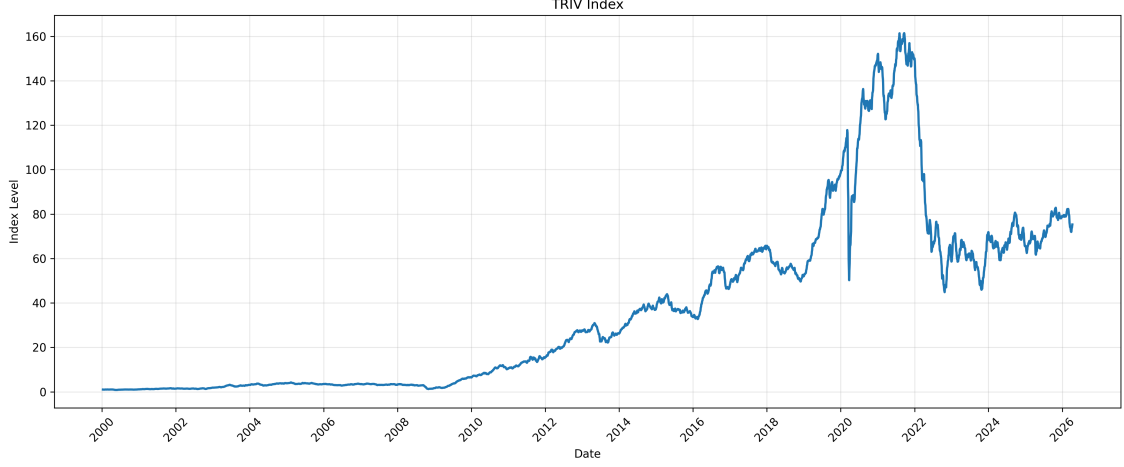


Figure 2: A visualization of a Total Return Bond Index Value (TRIV).

2.3.4 Applicability of the CAPM

Since TRIV indices provide return series for diversified credit portfolios, where we can make the same assumptions as for equities except the utility function, the standard CAPM can be applied directly to credit index returns. The model implies

$$\mathbb{E}[r_i^e] = \beta_i \mathbb{E}[r_m^e], \quad (9)$$

where

$$\beta_i = \frac{\text{Cov}(r_i^e, r_m^e)}{\text{Var}(r_m^e)}. \quad (10)$$

Thus, from a statistical perspective, the CAPM is well defined for credit index returns.

However, we have to note that since the default risk for individual bonds is virtually diversified away. The TRIV is not a perfect representation of the asymmetry in credit returns described in previous sections. At the same time, the index is still composed of credit instruments whose upside is limited by contractual coupon and principal payments, while adverse credit-market conditions may generate large negative returns. Therefore, downside movements in TRIV returns should be interpreted

as reduced-form adverse credit-market states rather than as direct observations of individual repayment losses.

This motivates the use of a downside-risk framework, where systematic risk is measured by co-movement with adverse market states rather than by covariance with total market variation.

2.4 Lower Partial Moments

Let us now introduce the theory surrounding lower partial moments.

Definition 2.2. Let X be a real-valued random return with cumulative distribution function F . For a target return $\tau \in \mathbb{R}$ and order $\alpha > 0$, the *Lower Partial Moment* is defined as

$$\begin{aligned} LPM_\alpha(\tau) &= \int_{-\infty}^{\tau} (\tau - x)^\alpha dF(x) \\ &= \mathbb{E}[(\tau - X)_+^\alpha], \end{aligned}$$

where $(x)_+ = \max(0, x)$.

The lower partial moment measures the magnitude of shortfalls relative to the benchmark τ . Only outcomes for which $X < \tau$ contribute to the measure, while returns above the target do not increase the perceived risk.

Let us now verify that the lower partial moment is well defined as a finite expectation.

Proposition 2.3. *Let X be a random variable and $\alpha > 0$. If $\mathbb{E}[|X|^\alpha] < \infty$, then the lower partial moment $LPM_\alpha(\tau)$ exists and is finite for any $\tau \in \mathbb{R}$.*

Proof. Since $(\tau - X)_+ \leq |\tau - X|$, we have

$$(\tau - X)_+^\alpha \leq |\tau - X|^\alpha. \tag{11}$$

Using the generalized triangle inequality

$$|a + b|^\alpha \leq C_\alpha(|a|^\alpha + |b|^\alpha) \tag{12}$$

for some constant $C_\alpha > 0$, it follows that

$$|\tau - X|^\alpha \leq C_\alpha(|\tau|^\alpha + |X|^\alpha). \quad (13)$$

Taking expectations yields

$$LPM_\alpha(\tau) = \mathbb{E}[(\tau - X)_+^\alpha] \leq C_\alpha(|\tau|^\alpha + \mathbb{E}[|X|^\alpha]). \quad (14)$$

Since $\mathbb{E}[|X|^\alpha] < \infty$, the lower partial moment is finite. $\square$

The parameter α determines the sensitivity to downside deviations. Important special cases include

- $\alpha = 1$: expected shortfall below τ ,
- $\alpha = 2$: target semivariance.

To connect the LPM with the mean-risk framework introduced in the Section 2.1.2, we follow [Fishburn \(1977\)](#) and consider the so-called α, τ -utility function.

Definition 2.4 (Fishburn's α, τ -utility). Let $\tau \in \mathbb{R}$ be a target level, $\alpha > 0$, and $k > 0$. The *Fishburn α, τ -utility* is defined by

$$\begin{cases} u_{\alpha, \tau}(x) = x & \text{for } x \geq \tau, \\ u_{\alpha, \tau}(x) = x - k(\tau - x)_+^\alpha & \text{for } x \leq \tau. \end{cases} \quad (15)$$

This utility function introduces an explicit asymmetry in preferences. Outcomes above the target τ contribute linearly to utility, while outcomes below the target are penalized according to the magnitude of the shortfall. The parameter α determines how strongly large downside deviations are penalized, while k controls the overall degree of risk aversion.

Taking expectations yields

$$\mathbb{E}[u_{\alpha, \tau}(X)] = \mathbb{E}[X] - k \mathbb{E}[(\tau - X)_+^\alpha] = \mathbb{E}[X] - k LPM_\alpha(\tau). \quad (16)$$

Hence, maximizing expected Fishburn utility is equivalent to a mean-LPM criterion. This provides a direct utility-theoretic argument for lower partial moments, analogous to how quadratic utility leads to mean-variance preferences.

The key advantage of Fishburn’s specification is that it allows investors to distinguish between upside and downside risk. While traditional mean–variance analysis treats positive and negative deviations symmetrically, the LPM framework focuses exclusively on downside outcomes relative to a target level. This feature is particularly relevant in credit markets, where investors are primarily concerned with the risk of losses rather than unusually high gains.

It is precisely this asymmetry that makes lower partial moments a natural building block for equilibrium asset pricing models such as the LPM-CAPM developed below.

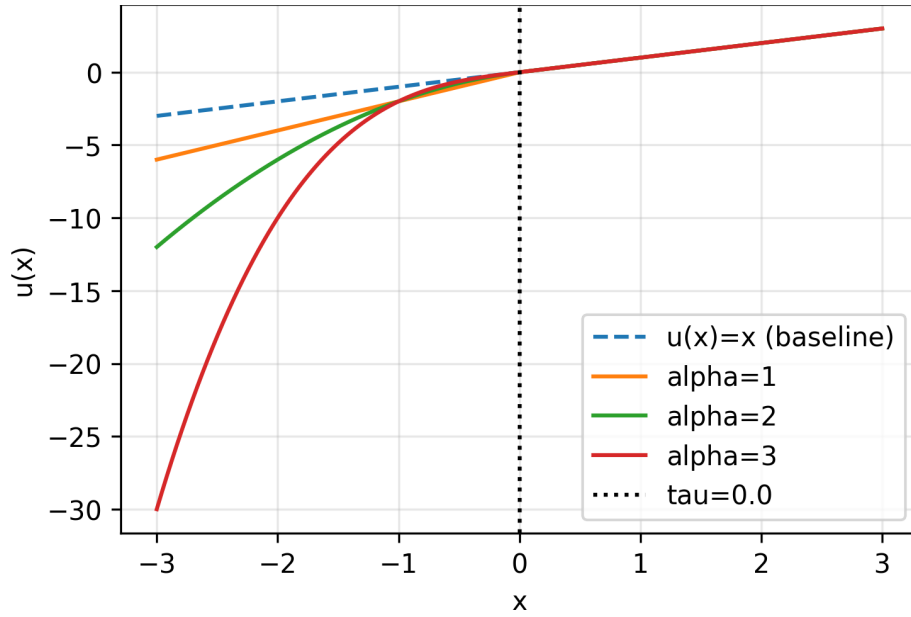


Figure 3: Fishburn α, τ -utility for different values of α . The dashed line represents the linear baseline $u(x) = x$. For outcomes above the target $\tau = 0$, utility is linear, while outcomes below the target are penalized increasingly as α rises, reflecting stronger sensitivity to downside risk.

2.5 A Credit Interpretation of the LPM-CAPM

The purpose of this section is to derive a pricing relation for credit assets based on downside risk, analogous to the classical CAPM. In particular, we aim to obtain an expression of the form

$$\mathbb{E}[r_i^e] = \beta_i^{LPM} \mathbb{E}[r_M^e], \quad (17)$$

where the beta is defined with respect to a downside risk measure rather than variance.

The derivation builds on the lower partial moment framework of [Fishburn \(1977\)](#), the portfolio-choice and pricing implications developed by [Bawa \(1975\)](#), and the downside-beta asset-pricing formulation of [Harlow and Rao \(1989\)](#). The argument follows the same equilibrium logic as the CAPM derivation in Section 2.2, but replaces the mean–variance objective with a mean–LPM objective.

Thus, the market assumptions used in the CAPM setting are retained: investors have homogeneous expectations, markets are frictionless, a risk-free asset is available, and investors can choose interior portfolio positions. Under these assumptions, all investors face the same investment opportunity set and the optimal risky portfolio must coincide with the market portfolio in equilibrium.

As in the classical CAPM derivation, we first consider the optimal risky portfolio chosen by a representative investor. Under homogeneous expectations and identical investment opportunities, all investors choose the same risky portfolio up to scale. Market clearing then implies that the aggregate risky portfolio must equal the aggregate supply of risky assets. Hence, in equilibrium, the representative investor’s optimal risky portfolio coincides with the market portfolio.

The key difference from the classical CAPM is the measure of risk. Instead of penalizing the variance of portfolio returns, investors are assumed to penalize downside shortfalls relative to a target return τ . Let r_i denote the excess return on asset i , and let

$$r_p = \sum_{i=1}^n \omega_i r_i \tag{18}$$

denote the excess return on a risky portfolio with weights $\omega_1, \dots, \omega_n$. We assume that investors evaluate portfolios according to the mean–LPM objective

$$V(\omega) = \mathbb{E}[r_p] - k \mathbb{E} \left[(\tau - r_p)_+^\alpha \right], \quad \alpha > 1, \quad k > 0. \tag{19}$$

Here, τ is the target return, α determines the sensitivity to downside shortfalls, and k controls the overall weight placed on downside risk. The restriction $\alpha > 1$ ensures that the downside penalty is sufficiently regular for the first-order condition used below.

In the credit setting, this objective is economically natural because credit payoffs

are asymmetric. Upside outcomes are limited by contractual payments, while adverse outcomes may involve substantial losses. The mean-LPM objective therefore captures the idea that investors care not only about expected excess return, but also about the tendency of returns to fall below a relevant benchmark.

2.5.1 Downside Covariance Pricing

We first characterize the implications of the LPM criterion for optimal portfolio choice.

Lemma 2.5 (Downside covariance pricing). *Suppose investors maximize $V(w)$ with interior optimal risky exposures. Then, in equilibrium where the optimal risky portfolio coincides with the market portfolio M , expected excess returns satisfy*

$$\mathbb{E}[r_i] = \lambda_D \text{Cov}(r_i, D_M), \quad (20)$$

where

$$D_M := (\tau - r_M)_+^{\alpha-1}, \quad \lambda_D := \frac{-k\alpha}{1 + k\alpha\mathbb{E}[D_M]}.$$

Since D_M is large in adverse market states, assets with low returns in those states tend to have negative covariance with D_M . The negative price of downside covariance therefore implies positive compensation for exposure to systematic downside states.

Proof. Consider a representative investor choosing risky portfolio weights ω . The excess return on the investor's risky portfolio is

$$r_p = \sum_{j=1}^n \omega_j r_j. \quad (21)$$

For an interior optimum, the first-order condition with respect to ω_i is

$$\frac{\partial V(w)}{\partial \omega_i} = \mathbb{E}[r_i] + k\alpha\mathbb{E}[(\tau - r_p)_+^{\alpha-1} r_i] = 0. \quad (22)$$

Thus,

$$\mathbb{E}[r_i] = -k\alpha\mathbb{E}[(\tau - r_p)_+^{\alpha-1} r_i]. \quad (23)$$

Under homogeneous expectations, all investors face the same investment opportunity set and solve the same risky-portfolio problem. Therefore, their optimal risky portfolios are proportional to the same vector of risky asset weights. In equilibrium, aggregate demand for risky assets must equal aggregate supply. It follows that this common risky portfolio must coincide with the market portfolio. Hence, in equilibrium,

$$r_p = r_M. \quad (24)$$

Defining $D_M = (\tau - r_M)_+^{\alpha-1}$ yields

$$\mathbb{E}[r_i] = -k\alpha\mathbb{E}[D_M r_i]. \quad (25)$$

Using $\mathbb{E}[D_M r_i] = \text{Cov}(r_i, D_M) + \mathbb{E}[r_i]\mathbb{E}[D_M]$ and rearranging gives the result. $\square$

The lemma shows that expected returns are determined by covariance with a downside state variable D_M , which assigns greater weight to adverse market outcomes. In credit markets, such states correspond to periods of widespread credit deterioration.

2.5.2 Equilibrium Downside Beta Pricing

We now obtain the equilibrium pricing relation. Given the equilibrium result above, we can express expected returns in terms of downside risk exposure. In particular, defining the downside beta as

$$\beta_i^{LPM} = \frac{\text{Cov}(r_i, D_M)}{\text{Cov}(r_M, D_M)}, \quad (26)$$

where $D_M = (\tau - r_M)_+^{\alpha-1}$. By applying the lemma to both asset i and the market portfolio M and dividing the two, we obtain the following pricing relation:

$$\mathbb{E}[r_i] = \beta_i^{LPM} \mathbb{E}[r_M]. \quad (27)$$

This expression constitutes the LPM-CAPM for credit assets. It mirrors the classical CAPM, but replaces variance-based risk with downside risk relative to a

target return. This modification is particularly relevant in credit markets, where returns are asymmetric due to limited upside and potential default losses.

2.5.3 Connection to Repayment Risk

To specialize the model to credit markets, consider a credit claim i with promised payoff $\bar{X}_i$ and repayment fraction $q_i \in (0, 1]$. Let P_i be the price the investor bought the contract claim. The realized return satisfies

$$R_i = \frac{q_i \bar{X}_i}{P_i} - 1, \quad r_i = R_i - R_f. \quad (28)$$

The repayment factor q_i captures the fundamental source of uncertainty in credit instruments. In adverse market conditions, repayment rates tend to decline, linking credit losses directly to downside market states.

Proposition 2.6 (Repayment-factor representation of downside beta). *Suppose $\bar{X}_i$ is non-random. Then*

$$\text{Cov}(r_i, D_M) = \frac{\bar{X}_i}{P_i} \text{Cov}(q_i, D_M), \quad (29)$$

and therefore

$$\beta_{iM}^{LPM} = \frac{\bar{X}_i}{P_i} \frac{\text{Cov}(q_i, D_M)}{\text{Cov}(r_M, D_M)}. \quad (30)$$

Proof. From

$$r_i = \frac{q_i \bar{X}_i}{P_i} - (1 + R_f), \quad (31)$$

and since constants do not affect covariance, the result follows immediately. $\square$

Remark 2.7. The assumption that $\bar{X}_i$ is non-random is not always true in reality due to possible floating coupon contracts or embedded options.

This representation highlights that systematic risk in credit markets is fundamentally driven by the co-movement between repayment rates and downside market conditions. In particular, assets whose repayment fractions deteriorate more strongly in adverse states command higher expected excess returns.

Corollary 2.8 (Idiosyncratic default risk is not priced). *If $\text{Cov}(q_i, D_M) = 0$, then $\beta_{iM}^{LPM} = 0$ and hence*

$$\mathbb{E}[r_i] = 0. \quad (32)$$

This result shows that only systematic variation in repayment risk is priced in equilibrium, while purely idiosyncratic default risk does not command a risk premium.

2.5.4 Pricing Implications

Using $\mathbb{E}[R_i] = R_f + \mathbb{E}[r_i]$, the price of a credit claim satisfies

$$P_i = \frac{\mathbb{E}[q_i \bar{X}_i]}{1 + R_f + \beta_{iM}^{LPM} \mathbb{E}[r_M]}. \quad (33)$$

This expression should be interpreted as an equilibrium expected-return representation rather than a full structural credit pricing model. The term $\beta_{iM}^{LPM} \mathbb{E}[r_M]$ captures compensation for exposure to systematic downside credit risk.

2.5.5 The choice of τ and α

Recall Fishburn's α, τ -utility,

$$u_{\alpha, \tau}(x) = x - k(\tau - x)_+^\alpha, \quad k > 0.$$

The parameters τ and α determine what is treated as a downside outcome and how strongly such outcomes reduce investor utility. The target τ defines the return level below which outcomes are penalized, while α determines the sensitivity to the size of the shortfall.

A useful way to interpret τ is as the investor's required benchmark return. If τ is chosen too high, many ordinary return realizations will be treated as downside outcomes. For a sufficiently downside-risk-averse investor, this may imply that almost any risky asset appears unattractive, since even moderate underperformance is penalized heavily. At the other extreme, if τ is chosen too low, only very severe losses are classified as downside outcomes. Then the lower partial moment may become close to zero for most assets, and the model gives little weight to ordinary credit risk.

Real-world investors may differ in what they regard as a sufficient return and in how strongly they penalize downside outcomes. In an equilibrium asset-pricing model, however, these heterogeneous preferences are summarized by market-level parameters. We therefore interpret τ and α as representative, or market-implied, parameters: τ reflects the return level below which outcomes are treated as downside states, while α captures the market's sensitivity to the severity of such shortfalls.

This interpretation is consistent with [Harlow and Rao \(1989\)](#), who argue that the target return in a mean-LPM framework need not be fixed at the risk-free rate, but may instead be treated as a parameter to be investigated empirically. In the empirical analysis, we therefore evaluate several values of τ , chosen relative to the distribution of excess returns in the data.

The parameter α determines the curvature of the downside penalty. If $\alpha = 1$, shortfalls are penalized linearly. If $\alpha < 1$, the penalty function is concave in the size of the shortfall, meaning that the marginal penalty for additional losses decreases as losses become larger. Such a specification does not represent downside risk aversion in the usual sense and may instead imply risk-seeking behavior over losses. For the LPM-CAPM derived above, we therefore restrict attention to $\alpha > 1$. This condition ensures that the downside penalty is convex in losses and that the downside market factor

$$D_M = (\tau - r_M)_+^{\alpha-1}$$

is well defined in the first-order condition. Values of α close to one place relatively more emphasis on whether the market return falls below the target, whereas larger values of α give more weight to the severity of large downside realizations.

An argument is made in [Satchell \(2001, Remark 1, p. 166\)](#) that a natural restriction in raw-return units is

$$0 < \tau < R_f.$$

Economically, this means that the target should be positive but not more demanding than the risk-free return. This avoids two undesirable extremes: a target so high that almost all risky investments are rejected, and a target so low that downside risk becomes irrelevant. In the empirical analysis, however, this restriction is used as guidance rather than as a binding constraint. Since returns are measured in excess-return units, positive values of τ allow us to test stricter downside benchmarks than the risk-free rate.

The choice of τ must also be interpreted relative to the return variable used in

the empirical analysis. Since our empirical tests are conducted using excess returns,

$$r_{i,t}^e = R_{i,t} - R_{f,t},$$

the downside condition is

$$r_{i,t}^e < \tau.$$

In raw-return units, this is equivalent to

$$R_{i,t} < R_{f,t} + \tau.$$

Thus, $\tau = 0$ corresponds to using the risk-free return as the benchmark. A positive value of τ means that the asset must exceed the risk-free return by at least τ in order not to be classified as a downside outcome, while a negative value of τ allows some underperformance relative to the risk-free rate before the outcome is classified as downside.

Figure 4 illustrates this intuition. The histogram shows weekly excess returns for all total return bond indices. Most observations are clustered close to zero, but the distribution exhibits a clear left tail. Choosing τ close to zero therefore captures economically relevant downside realizations without focusing exclusively on extreme crashes. The empirical analysis evaluates several combinations of τ and α in order to examine whether credit returns are better explained by moderate downside risk or by more severe downside market states.

2.5.6 From Single-Period Credit Claims to TRIV Returns

The theoretical LPM-CAPM derived in Section 2.5 is formulated for a single-period credit claim whose payoff uncertainty is represented by the repayment fraction q_i . In the empirical analysis, however, the observable assets are total return index values (TRIV), which represent diversified portfolios of corporate bonds. The empirical implementation therefore requires an interpretation of the theoretical model at the index-return level.

We treat TRIV returns as a reduced-form empirical counterpart of the single-period credit return. Rather than modeling the repayment fraction q_i directly, we use observed excess returns on bond indices. Under this interpretation, the theoretical

pricing relation

$$\mathbb{E}[r_i] = \beta_i^{LPM} \mathbb{E}[r_M] \quad (34)$$

is applied to portfolio returns, where β_i^{LPM} measures the sensitivity of index i to systematic downside movements in the credit market.

The downside market state is represented by

$$D_{M,t} = (\tau - r_{M,t}^e)^{\alpha-1}_+, \quad (35)$$

which is large when the market excess return falls below the target level τ . The empirical downside beta is then defined as

$$\beta_i^{LPM} = \frac{\text{Cov}(r_{i,t}^e, D_{M,t})}{\text{Cov}(r_{M,t}^e, D_{M,t})}. \quad (36)$$

This beta is interpreted as a reduced-form measure of systematic downside exposure. It does not isolate a specific structural channel such as repayment risk, spread risk, interest-rate risk, or liquidity risk. Instead, it measures whether the total return of index i tends to be low in the same states in which the market return is below the target.

This interpretation allows the theoretical LPM-CAPM to be taken to the data without requiring individual bond-level repayment information. The theoretical repayment factor q_i is replaced by observed total return dynamics, while the central asset-pricing implication is preserved: assets with greater exposure to systematic downside market states should command higher expected excess returns. The empirical tests therefore examine whether this reduced-form downside beta helps explain the cross-section of credit index excess returns.

3 Empirical Study

The theoretical framework establishes that an asset pricing relation analogous to the CAPM can be derived using downside risk, as measured by lower partial moments. The purpose of this section is to examine whether such a relationship is supported by the data.

In particular, the LPM-CAPM implies a linear relation between the expected excess return of an asset and its exposure to systematic downside risk. To evaluate this implication, we test whether cross-sectional differences in average excess returns can be explained by LPM-based betas.

The empirical analysis is conducted using a two-pass cross-sectional regression procedure based on the method in [Fama and MacBeth \(1973\)](#). This framework allows us to assess whether the key theoretical restrictions hold in the data, namely whether the intercept is zero and whether the estimated price of risk is consistent with the expected excess return on the market portfolio.

In addition to hypothesis testing, we evaluate model performance using pricing error measures such as the mean squared error (MSE) and cross-sectional R^2 . These metrics provide complementary evidence on how well the model explains the cross-section of returns.

A well-known issue in beta-pricing models is that estimated betas are measured with error see *On the Estimation of Beta-Pricing Models* by [Shanken \(1992\)](#). This leads to an errors-in-variables problem, which affects both the estimated risk premia and their statistical inference. In particular, standard Fama–MacBeth procedures tend to understate the true sampling variability of the estimated coefficients, resulting in overstated t -statistics.

Furthermore, financial return data often exhibit serial dependence, arising from both market dynamics and data construction. If not properly accounted for, this may lead to biased inference.

To address these issues, we employ a stationary bootstrap procedure, which preserves the dependence structure of the data and allows for more reliable statistical inference.

The empirical analysis implements a reduced-form version of the LPM-CAPM, where total return indices (TRIV) are used as observable proxies for the theoretical payoff structure. Consequently, the estimated downside beta reflects the overall exposure of credit returns to systematic downside market states, rather than a single

structural risk component.

3.1 Data

The empirical sample consists of weekly excess returns on 30 total return bond indices observed over a 26-year period, from January 2000 to April 2026. The underlying bond index series are provided by ICE Data Indices, LLC and correspond to ICE BofA total return bond indices ([ICE Data Indices, LLC, 2025](#)). The data were downloaded using the FRED API ([Federal Reserve Bank of St. Louis, 2026](#)). The weekly returns are calculated by

$$R_{i,t} = \frac{P_{i,t}}{P_{i,t-5}} - 1$$

to reduce high-frequency noise. We use a Friday-to-Friday structure: for each Friday, or the closest available trading day before the weekend, we compute the return relative to the corresponding observation in the previous week.

Total return index values (TRIV) measure the full economic return from holding a portfolio of bonds, including both coupon payments and price changes. Since the asset pricing models are formulated in terms of excess returns, a proxy for the risk-free rate is required. We use the Federal Funds Rate (DFF) as a proxy (converted to fit weekly returns $\frac{DFF_t}{100.52}$), as it represents a short-term benchmark rate with minimal credit risk in U.S. dollar-denominated markets. Since DFF is an overnight interbank rate rather than a directly traded asset, it should be interpreted as a benchmark rate rather than as the return on an investable risk-free asset.

The indices can be grouped into four broad categories:

1. **Investment Grade Corporate Bonds:** high credit quality corporate debt (AAA–BBB).
2. **Maturity Segmented Corporate Bonds:** portfolios grouped by remaining maturity.
3. **High Yield Bonds:** lower-rated corporate debt (BB and below).
4. **Emerging Market Bonds:** bond portfolios exposed to developing economies.

A broad credit market index (BAMLCC0A0CMTRIV) is used as a proxy for the market portfolio. This market proxy is excluded from the test assets, leaving 30 indices in the cross-section.

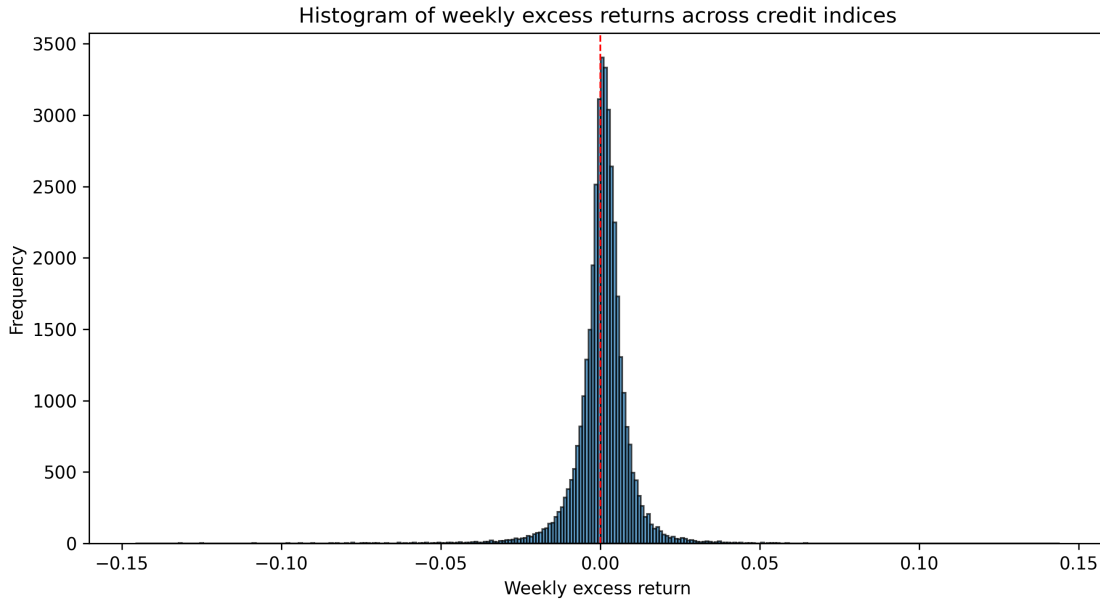


Figure 4: Histogram of weekly excess returns for all total return bond index. The distribution is centered near zero but exhibits negative skewness and a pronounced left tail, highlighting the importance of downside risk in credit markets.

3.2 Dependence Structure and Inference

The empirical framework relies on a two-pass cross-sectional regression, where risk premia are estimated over time. Standard inference in this setting assumes that the sequence of estimated coefficients is independently distributed. However, this assumption is unlikely to hold for the present data.

Financial return series are known to exhibit time dependence, including auto-correlation and volatility clustering. These features imply that the estimated risk premia inherit serial dependence, leading to biased standard errors and unreliable hypothesis tests if standard methods are applied. To account for these issues, a stationary bootstrap procedure is employed.

3.3 Empirical Methodology

To evaluate the empirical relevance of the LPM-CAPM, we employ a two-pass cross-sectional regression procedure in the spirit of [Fama and MacBeth \(1973\)](#). We compare the standard CAPM against a grid of LPM-CAPM specifications with varying target return τ and downside order α .

3.3.1 First Pass: Beta Estimation

For each asset i and week t , define excess returns

$$r_{i,t}^e = R_{i,t} - R_{f,t}, \quad r_{m,t}^e = R_{m,t} - R_{f,t}.$$

The CAPM beta is estimated using sample moments:

$$\hat{\beta}_i^{CAPM} = \frac{\widehat{\text{Cov}}(r_{i,t}^e, r_{m,t}^e)}{\widehat{\text{Var}}(r_m^e)}, \quad (37)$$

For the LPM-CAPM, define the downside transformed market factor

$$D_{m,t} = (\tau - r_{m,t}^e)_+^{\alpha-1}, \quad (38)$$

where τ is the target return and $\alpha > 1$ determines downside sensitivity. The downside beta estimate is then

$$\hat{\beta}_i^{LPM} = \frac{\widehat{\text{Cov}}(r_{i,t}^e, D_{m,t})}{\widehat{\text{Cov}}(r_{m,t}^e, D_{m,t})}, \quad (39)$$

To allow for time variation and avoid look-ahead bias, betas are estimated using non-overlapping windows of 52 weekly return observations so that we estimate a beta over the full year.

3.3.2 Second Pass: Fama–MacBeth Regressions

The second-pass regression is the empirical counterpart of the pricing relations derived in the theory section. Under the CAPM, the model implies

$$\mathbb{E}[r_i^e] = \beta_i^{CAPM} \mathbb{E}[r_m^e],$$

while the LPM-CAPM in equation (27) implies the analogous relation

$$\mathbb{E}[r_i^e] = \beta_i^{LPM} \mathbb{E}[r_m^e].$$

Thus, in both models, expected excess returns should be linearly related to the corresponding beta. Since expected returns are not directly observed, the Fama–MacBeth procedure replaces the unobserved conditional expected return at each date by the realized excess return. The realized return can be interpreted as

$$r_{i,t}^e = \mathbb{E}_{t-1}[r_{i,t}^e] + u_{i,t},$$

where $u_{i,t}$ is an unexpected return innovation. The cross-sectional regression therefore uses realized excess returns as noisy observations of expected excess returns. We do not impose a specific distributional assumption on $u_{i,t}$; instead, statistical inference is based on the stationary bootstrap described later.

At each observation t , we estimate the cross-sectional regression

$$r_{i,t}^e = \gamma_{0,t} + \gamma_{1,t} \hat{\beta}_{i,t} + \varepsilon_{i,t}. \quad (40)$$

The error term $\varepsilon_{i,t}$ captures the part of the realized excess return that is not explained by the estimated beta in the cross-section at time t . It therefore includes both pricing errors and unexpected return shocks. Since the realized returns are noisy observations of expected returns, the second-pass coefficients $\gamma_{0,t}$ and $\gamma_{1,t}$ are estimated with sampling error. We do not rely on normality or independence assumptions for $\varepsilon_{i,t}$ when conducting inference; instead, the sampling distribution of the average coefficients is approximated using the stationary bootstrap.

Since expected excess returns are unobserved, the second-pass regressions use realized excess returns as noisy observations of conditional expected returns. Averaging the estimated coefficients over time then estimates the unconditional cross-sectional pricing relation.

The coefficients are then averaged over time:

$$\bar{\gamma}_k = \frac{1}{T} \sum_{t=1}^T \gamma_{k,t}, \quad k = 0, 1. \quad (41)$$

Under both CAPM and LPM-CAPM, the theoretical restrictions are

$$\gamma_0 = 0, \quad \gamma_1 = \mathbb{E}[r_m^e]. \quad (42)$$

These restrictions follow directly from the theoretical asset pricing model. Under the CAPM, expected excess returns are proportional to market beta:

$$\mathbb{E}[r_i^e] = \beta_i \mathbb{E}[r_m^e].$$

The LPM-CAPM yields an analogous relation with downside beta. When this relation is tested in a cross-sectional regression, it implies that the intercept must be zero and that the slope coefficient equals the expected excess market return.

3.3.3 Dependence, Generated Regressors, and Limitations of Cross-Sectional Regression

The two-pass procedure is a standard way to test beta-pricing relations see [Cochrane \(2005, Ch.12\)](#), but inference is complicated by several features of the data and estimation procedure. First, the second-pass regressions use estimated betas rather than observed risk exposures. As shown by [Shanken \(1992\)](#), this creates an errors-in-variables problem: treating estimated betas as known regressors may understate the sampling uncertainty of the estimated prices of risk.

Second, financial return data may exhibit serial dependence, volatility clustering, and fat tails. As a result, the sequence of estimated cross-sectional coefficients may not be independent over time, and inference based on normality or independence assumptions may be unreliable.

These issues do not invalidate the two-pass regression framework, but they imply that conventional standard errors should be interpreted cautiously. We therefore use a stationary bootstrap, described below, to approximate the sampling distribution of the estimated parameters.

3.3.4 Stationary Bootstrap

We use the stationary bootstrap introduced by [Politis and Romano \(1994\)](#). The method resamples blocks of consecutive observations with random lengths, thereby preserving the temporal dependence structure of the data.

Let p denote the probability of starting a new block, implying an expected block length of $1/p$. A bootstrap sample is constructed by:

1. Selecting a random starting observation,
2. Continuing to the next observation with probability $1 - p$,
3. Restarting at a randomly selected observation with probability p ,
4. Repeating until a sample of length T is obtained.

A common rule of thumb is to set $p = T^{-1/3}$.

The bootstrap is applied along the time dimension of the full return panel, such that the cross-sectional structure is preserved at each point in time. For each bootstrap replication, the entire estimation procedure is repeated. Specifically, a resampled dataset is first generated, after which risk exposures are estimated in a first step, followed by cross-sectional regressions in a second step. The resulting estimates are then averaged over time. This procedure is repeated $B = 1000$ times, yielding an empirical distribution of the estimated parameters.

The stationary bootstrap is not an analytical solution to the errors-in-variables problem discussed by [Shanken \(1992\)](#). The issue remains that the betas used in the second-pass regressions are estimated rather than observed. That is, the estimated beta can be written as

$$\hat{\beta}_i = \beta_i + \eta_i,$$

where η_i denotes first-pass estimation error. If $\hat{\beta}_i$ is treated as a fixed regressor in the second pass, the uncertainty from η_i is not fully reflected in the inference.

The bootstrap mitigates this concern by re-estimating the betas in each bootstrap replication. For each resampled return panel, the first-pass beta becomes

$$\hat{\beta}_i^{*(b)} = \beta_i + \eta_i^{*(b)},$$

where the estimation error varies across bootstrap samples. The second-pass regressions are then estimated using these bootstrap-specific betas. As a result, the

bootstrap distribution of the estimated pricing parameters reflects variation arising both from realized return shocks and from first-pass beta estimation error. Thus, the bootstrap does not remove the errors-in-variables problem, but it allows part of the beta-estimation uncertainty to enter the empirical sampling distribution. This leads to more cautious inference than treating the original estimated betas as known regressors.

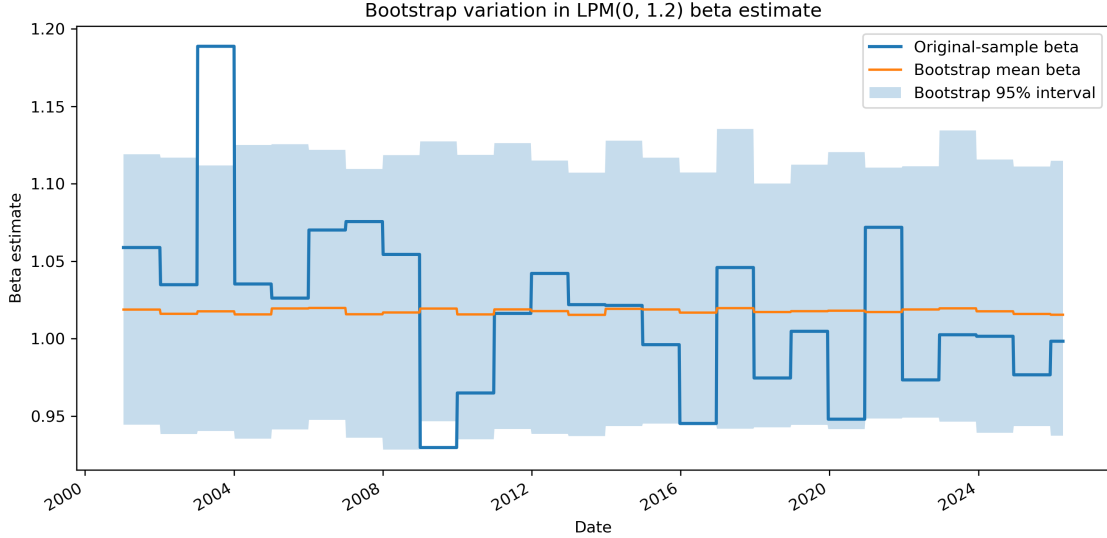


Figure 5: The figure should be interpreted as a visualization of estimation uncertainty rather than as direct evidence of the true beta error, since the true beta is unobserved. The dispersion of the bootstrap beta estimates indicates how sensitive the estimated risk exposure is to resampling of the return panel.

3.3.5 Hypothesis Tests

The hypothesis tests are based on the bootstrap distribution of the estimated average risk premia. This is a non-parametric inference procedure: rather than assuming that the estimated coefficients are normally distributed, we approximate their sampling distribution by repeatedly applying the full estimation procedure to stationary bootstrap samples. This approach is motivated by standard bootstrap theory for dependent data, where block-based resampling methods are used to preserve serial dependence in time series observations.

Bootstrap p-values are computed using a two-sided centered bootstrap test following [Hall \(1992, Sec. 3.12\)](#). For a parameter θ with null value θ_0 , let $\hat{\theta}^{*(1)}, \dots, \hat{\theta}^{*(B)}$ denote the bootstrap estimates and let $\hat{\theta}$ be the original sample estimate. The test

compares the observed deviation from the null, $|\hat{\theta} - \theta_0|$, with the empirical distribution of bootstrap deviations around the original estimate, $|\hat{\theta}^{*(b)} - \hat{\theta}|$. Centering the bootstrap distribution around $\hat{\theta}$ is intended to approximate the sampling variation of the estimator without imposing a parametric normal approximation.

The p-value is then computed as

$$p = \frac{1}{B} \sum_{b=1}^B \mathbf{1} \left\{ |\hat{\theta}^{*(b)} - \hat{\theta}| \geq |\hat{\theta} - \theta_0| \right\}.$$

For the intercept test, $\theta = \gamma_0$ and $\theta_0 = 0$. For the slope restriction, $\theta = \gamma_1$ and $\theta_0 = \bar{r}_m$, where $\bar{r}_m$ denotes the sample average excess return on the market proxy.

In addition to testing the two restrictions separately, we conduct a joint Wald test of the full pricing restriction. A Wald test measures the distance between the estimated parameter vector and the value imposed under the null, scaled by the covariance matrix of the estimator. In a conventional Wald test, this statistic is compared with an asymptotic χ^2 distribution. Here, however, we use the same quadratic form but obtain the reference distribution from the stationary bootstrap. This follows the bootstrap testing logic discussed by Horowitz (2001, Sec. 3.3), where bootstrap critical values are used when first-order asymptotic approximations may be unreliable.

The null hypothesis is

$$H_0 : \begin{pmatrix} \gamma_0 \\ \gamma_1 \end{pmatrix} = \begin{pmatrix} 0 \\ \bar{r}_m \end{pmatrix},$$

which tests whether the intercept is zero and whether the estimated price of risk equals the average excess return on the market portfolio simultaneously.

Let

$$\hat{\boldsymbol{\theta}} = \begin{pmatrix} \hat{\gamma}_0 \\ \hat{\gamma}_1 \end{pmatrix}, \quad \boldsymbol{\theta}_0 = \begin{pmatrix} 0 \\ \bar{r}_m \end{pmatrix}.$$

The observed Wald statistic is defined as

$$W_{\text{obs}} = (\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}_0)' \hat{\Sigma}^{-1} (\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}_0),$$

where $\hat{\Sigma}$ is the bootstrap covariance matrix of $(\hat{\gamma}_0, \hat{\gamma}_1)'$.

To make the joint test consistent with the centered bootstrap inference used for the individual restrictions, the p-value is computed from the centered bootstrap

distribution of the Wald statistic. For each bootstrap replication, define

$$\hat{\boldsymbol{\theta}}^{*(b)} = \begin{pmatrix} \hat{\gamma}_0^{*(b)} \\ \hat{\gamma}_1^{*(b)} \end{pmatrix}.$$

The centered bootstrap Wald statistic is then

$$W^{*(b)} = (\hat{\boldsymbol{\theta}}^{*(b)} - \hat{\boldsymbol{\theta}})' \hat{\Sigma}^{-1} (\hat{\boldsymbol{\theta}}^{*(b)} - \hat{\boldsymbol{\theta}}).$$

The joint bootstrap p-value is computed as

$$p_{\text{joint}} = \frac{1}{B} \sum_{b=1}^B \mathbf{1} \{ W^{*(b)} \geq W_{\text{obs}} \}.$$

The individual tests evaluate each pricing restriction separately, while the joint Wald test evaluates whether both restrictions hold simultaneously. Since the joint test accounts for the covariance between the estimated intercept and slope, its conclusion need not coincide with the two individual tests. Because the cross-section contains 30 indices, the results should be interpreted as evidence for this specific credit-index universe rather than as a universal asset-pricing result.

3.3.6 Model Evaluation

A model may fail to be rejected statistically without necessarily fitting the data well. We therefore supplement hypothesis testing with model comparison measures based on pricing performance.

For each model we compute:

- Mean Squared Pricing Error (MSE),
- Average cross-sectional R^2 ,
- γ_0 ,
- $\gamma_1 - \mathbb{E}[r_m^e]$.

The mean squared pricing error (MSE) is defined as

$$\text{MSE} = \frac{1}{NT} \sum_{t=1}^T \sum_{i=1}^N \hat{\varepsilon}_{i,t}^2, \tag{43}$$

where $\hat{\varepsilon}_{i,t}$ are the residuals from the cross-sectional regression.

The cross-sectional R^2 at time t is defined as

$$R_t^2 = 1 - \frac{\sum_i \hat{\varepsilon}_{i,t}^2}{\sum_i (r_{i,t}^e - \bar{r}_t)^2}, \quad (44)$$

and the reported value is the time average of R_t^2 . The reported MSE averages squared pricing errors across both assets and time. The reported R^2 is the time-series average of the cross-sectional R_t^2 .

In the empirical comparison, we evaluate the CAPM together with a grid of LPM-CAPM specifications. The order parameter is chosen as $\alpha \in \{1.2, 2, 3\}$ in order to span different degrees of downside-risk aversion. Values close to one place relatively more weight on whether returns fall below the target, while larger values place greater weight on the severity of downside shortfalls. The target parameter is chosen as $\tau \in \{-0.02, -0.01, 0, 0.01, 0.02\}$, based on the empirical distribution of weekly excess returns. Since the return variable is measured in excess-return units, $\tau = 0$ corresponds to using the risk-free rate as the benchmark in raw-return units. Choosing substantially larger values of τ would classify almost all observations as downside outcomes, while substantially lower values would make the downside factor active only in rare extreme observations. The selected grid therefore allows us to examine whether credit-index returns are better explained by broad downside states close to the risk-free benchmark or by more severe downside market realizations.

4 Results

This section presents the empirical results from the two-pass cross-sectional regressions. The full set of results for the original sample and the stationary bootstrap procedure are reported in the Appendix. To make the main findings easier to interpret, Table 1, Table 2, and Table 3 summarize selected model specifications.

All estimated pricing coefficients are reported in weekly decimal return units and should therefore be interpreted relative to the scale of the underlying excess returns. In the sample, the average weekly excess return across the credit indices $\sum_{i=1}^N \bar{r}_i^e = 0.000708$, while the corresponding average weekly excess return on the market proxy is $\hat{r}_m^e = 0.000588$. Thus, an estimated intercept of, for example, 0.0005 corresponds to approximately 0.05% per week. Figure 4 illustrates the distribution of the weekly excess returns across the credit indices and provides context for the magnitudes reported below.

The hypothesis tests evaluate whether the theoretical pricing restrictions implied by the CAPM and LPM-CAPM can be rejected. In this setting, rejecting the intercept restriction $\gamma_0 = 0$ indicates that the proposed risk premium does not fully explain the average cross-sectional differences in excess returns. Similarly, rejecting the slope restriction $\gamma_1 = \bar{r}_m$ indicates that the estimated price of risk is not consistent with the average excess return on the market portfolio. The joint test evaluates whether both restrictions hold simultaneously. Hence, the tests are not simply measures of statistical fit, but tests of whether the proposed beta-pricing relation is sufficient to explain the cross-section of credit returns.

4.1 Hypothesis Tests

Table 1 reports the original-sample (non-bootstrapped data) hypothesis tests for selected models. The CAPM is included as the benchmark. The LPM specifications are chosen to represent different outcomes across the hypothesis tests and model evaluation measures.

Table 1: Original Sample Hypothesis Tests for Selected Models

Model	$\hat{\gamma}_0$	$p(\gamma_0 = 0)$	$\hat{\gamma}_1 - \bar{r}_m$	$p(\gamma_1 = \bar{r}_m)$	p_{joint}
CAPM	0.000459	0.132228	-0.000198	0.435530	0.321670
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000750	0.029706	-0.000540	0.040989	0.067121
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000261	0.388697	-0.000016	0.954082	0.593178
LPM ($\tau = -0.02, \alpha = 3.0$)	-0.000268	0.485180	0.001036	0.178022	0.396743

In the original sample, the CAPM is not rejected by either individual restriction or by the joint test. The intercept p-value is 0.132228, the slope p-value is 0.435530, and the joint p-value is 0.321670. Thus, the original-sample results do not provide strong evidence against the CAPM pricing restrictions.

The LPM specification with $\tau = 0.00$ and $\alpha = 1.2$ gives a different result. Both individual restrictions are rejected at the five percent level, with $p(\gamma_0 = 0) = 0.029706$ and $p(\gamma_1 = \bar{r}_m) = 0.040989$. However, the joint p-value is 0.067121, which is above the five percent level but close to rejection. This indicates evidence against the individual restrictions, but weaker evidence against the full set of pricing restrictions when tested jointly.

For the LPM specification with $\tau = 0.01$ and $\alpha = 3.0$, the original-sample p-values are much larger. The intercept p-value is 0.388697, the slope p-value is 0.954082, and the joint p-value is 0.593178. This specification is therefore not rejected by the original-sample tests. The low-target specification with $\tau = -0.02$ and $\alpha = 3.0$ is also not rejected, although its slope estimate is above the average market excess return.

Table 2 reports the corresponding centered bootstrap hypothesis tests.

Table 2: Bootstrap Hypothesis Tests for Selected Models

Model	$\hat{\gamma}_0$	$p(\gamma_0 = 0)$	$\hat{\gamma}_1 - \bar{r}_m$	$p(\gamma_1 = \bar{r}_m)$	p_{joint}
CAPM	0.000459	0.250000	-0.000198	0.682000	0.417000
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000750	0.034000	-0.000540	0.074000	0.131000
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000261	0.655000	-0.000016	0.989000	0.552000
LPM ($\tau = -0.02, \alpha = 3.0$)	-0.000268	0.865462	0.001036	0.730924	0.811245

The bootstrap results give weaker evidence against the pricing restrictions than the original-sample tests. For the CAPM, none of the restrictions are rejected. The intercept p-value increases to 0.250000, the slope p-value to 0.682000, and the joint p-value to 0.417000. This suggests that, once sampling uncertainty is evaluated using the centered bootstrap distribution, the data do not provide evidence against the CAPM restrictions.

The LPM specification with $\tau = 0.00$ and $\alpha = 1.2$ still rejects the intercept restriction at the five percent level, with $p(\gamma_0 = 0) = 0.034000$. However, the slope restriction is no longer rejected at the five percent level, and the joint test gives $p_{\text{joint}} = 0.131000$. This means that the intercept result should be interpreted as evidence of a possible pricing error, but not as a rejection of the full set of pricing restrictions when they are tested jointly.

The LPM specification with $\tau = 0.01$ and $\alpha = 3.0$ is not rejected by the bootstrap tests. Its slope estimate is very close to the average excess market return, with $\hat{\gamma}_1 - \bar{r}_m = -0.000016$, and the slope p-value is 0.989000. The joint p-value is 0.552000, indicating no evidence against the simultaneous restrictions. The low-target specification with $\tau = -0.02$ and $\alpha = 3.0$ also has large p-values, including a joint p-value of 0.811245.

The confidence intervals reported in the Appendix provide additional information about the precision of these tests. For example, the bootstrap confidence interval for γ_0 in the LPM specification with $\tau = 0.00$ and $\alpha = 1.2$ is $[-0.000020, 0.001323]$, which includes zero even though the centered bootstrap p-value is below five percent. Since zero lies close to the lower boundary of the interval, this result should be interpreted as marginal evidence against the intercept restriction rather than a decisive rejection. For the low-target specification with $\tau = -0.02$ and $\alpha = 3.0$, the corresponding interval is much wider, $[-0.000637, 0.001797]$, indicating substantial uncertainty. Therefore, the large p-values for this specification should not be interpreted as strong confirmation of the model.

4.2 Model Evaluation

Table 3 reports the model evaluation measures for the same selected specifications. These measures are interpreted separately from the hypothesis tests. The hypothesis tests evaluate whether the theoretical pricing restrictions can be rejected, while the model evaluation measures describe how much cross-sectional variation the model explains.

Table 3: Model Performance Comparison for Selected Models

Sample	Model	Mean MSE	Mean R^2
Original	CAPM	0.000021	0.369725
Original	LPM ($\tau = 0.00, \alpha = 1.2$)	0.000021	0.379451
Original	LPM ($\tau = 0.01, \alpha = 3.0$)	0.000021	0.372337
Original	LPM ($\tau = -0.02, \alpha = 3.0$)	0.000024	0.248438
Bootstrap	CAPM	0.000023	0.337494
Bootstrap	LPM ($\tau = 0.00, \alpha = 1.2$)	0.000023	0.347674
Bootstrap	LPM ($\tau = 0.01, \alpha = 3.0$)	0.000023	0.339436
Bootstrap	LPM ($\tau = -0.02, \alpha = 3.0$)	0.000027	0.222318

The CAPM and the two LPM specifications with non-negative targets have very similar pricing errors. In the bootstrap results, all three have a reported mean

MSE of 0.000023. The LPM specification with $\tau = 0.00$ and $\alpha = 1.2$ has the highest reported average R^2 among the selected bootstrap models, equal to 0.347674, compared with 0.337494 for the CAPM and 0.339436 for the LPM specification with $\tau = 0.01$ and $\alpha = 3.0$. The differences are small, indicating that the gain in explanatory power from the LPM specifications is limited.

The low-target specification with $\tau = -0.02$ and $\alpha = 3.0$ performs worse according to the model evaluation measures. In the bootstrap results, it has a higher mean MSE of 0.000027 and a lower average R^2 of 0.222318. This is also the case in the original sample, where its average R^2 is 0.248438. Hence, although this specification is not rejected by the hypothesis tests, it explains less of the cross-sectional variation in excess returns.

4.3 Summary of Findings

The results provide mixed evidence. The CAPM is not rejected in either the original-sample tests or the centered bootstrap tests. Several LPM-CAPM specifications are also not rejected, especially specifications with higher values of α and non-negative values of τ .

At the same time, the model evaluation measures show only small differences between the CAPM and the selected LPM specifications with non-negative targets. The LPM specification with $\tau = 0.00$ and $\alpha = 1.2$ has a slightly higher average R^2 , but it also shows evidence of a non-zero intercept. Conversely, the low-target specification with $\tau = -0.02$ and $\alpha = 3.0$ is difficult to reject statistically, but has substantially weaker explanatory performance.

The main result is therefore cautious. The evidence suggests that downside beta may contain some information for explaining credit-index returns, but the improvement over the CAPM is modest. Moreover, the hypothesis tests and model evaluation measures do not always point in the same direction. This motivates a more detailed discussion of the interpretation, limitations, and economic meaning of the results in the next section.

5 Discussion

The theoretical framework motivates the LPM-CAPM by arguing that downside risk should be especially relevant in credit markets, where upside potential is limited while losses can be severe. However, the empiric results do not prove the theory correct, since the pricing restrictions are not uniformly satisfied and the improvement over the CAPM is modest.

One reason for this mixed evidence may be the use of total return indices. TRIV data are appropriate for measuring realized investor returns, since they include coupon income and price changes, but they are not a pure measure of credit repayment risk. The indices capture aggregate market sentiment about credit risk, rather than repayment risk directly. In the empirical analysis, downside movements in TRIV returns are therefore interpreted as reduced-form adverse credit-market states.

The mixed results may also reflect the strict assumptions behind single-factor beta-pricing models. Both the CAPM and LPM-CAPM assume that one systematic risk exposure is sufficient to explain expected excess returns. The models have some constraining assumptions of the market which are improbable to achieve. A non-zero intercept should therefore not necessarily be interpreted as evidence that downside risk is irrelevant, but rather that the proposed risk premium does not fully capture the cross-section on its own.

The comparison between the original-sample tests and the bootstrap tests shows that inference is sensitive to how uncertainty is treated. The bootstrap generally gives weaker evidence against the pricing restrictions, which suggests that the original tests may understate sampling uncertainty. This is consistent with the concern raised by [Shanken \(1992\)](#): because betas are estimated in the first pass, the second-pass risk-premium estimates are subject to generated-regressor uncertainty.

At the same time, the stationary bootstrap should not be interpreted as a complete solution to the Shanken problem. It accounts for serial dependence and re-estimates the full two-pass procedure across resampled panels, but it is not the same as an analytical Shanken correction. The bootstrap results therefore provide more cautious inference, rather than proof that the pricing restrictions are correct.

The empirical comparison between the CAPM and LPM-CAPM is mixed. The CAPM remains a strong benchmark, since it is not rejected by the hypothesis tests and performs similarly to the selected LPM specifications.

The choice of τ and α clearly affects the results. Specifications with $\tau = 0.00$ appear to capture broad negative market states, while very low targets such as $\tau = -0.02$ perform worse in terms of explanatory power. This suggests that, for diversified weekly credit indices, general downside market movements may be more informative than only extreme downside events.

The role of α seems less important. At least we do not observe a similar separation in performance as for positive and negative τ .

Further research could test the LPM-CAPM on more disaggregated credit data, such as individual corporate bonds. This would be closer to the theoretical payoff structure and could better capture default and repayment risk. It would also be useful to separate total returns into duration, spread, coupon, and default-related components.

Future work could also extend the model by adding additional credit-risk factors. A multifactor framework including duration risk, credit-spread risk, liquidity risk, and downside market risk could show whether LPM beta adds explanatory power beyond standard credit variables. Another promising direction is to estimate τ and α from the data or allow them to vary over time with market conditions.

Overall, the LPM-CAPM does not clearly dominate the CAPM in this empirical setting. Nevertheless, the analysis indicates that downside beta may contain economically relevant information for credit-index returns. The main contribution of the thesis is therefore not to establish the LPM-CAPM as a superior replacement for the CAPM, but to show how downside-risk pricing can be adapted to a credit-market setting and evaluated empirically using total return bond indices. The results suggest that downside risk is relevant, but that a single-factor downside-risk model is too restrictive to fully explain the cross-section of credit excess returns. Future research should therefore examine whether LPM-based measures add incremental explanatory power in multifactor credit-pricing models and in more disaggregated bond-level data.

6 Appendix

Table 4: Index names used in the Empirical Section

Category	Index Name	FRED Series Code
Investment Grade by Rating	AAA US Corporate	BAMLCC0A1AAATRIV
Investment Grade by Rating	AA US Corporate	BAMLCC0A2AATRIV
Investment Grade by Rating	A US Corporate	BAMLCC0A3ATRIV
Investment Grade by Rating	BBB US Corporate	BAMLCC0A4BBBTRIV
Investment Grade by Maturity	US Corporate 1–3Y	BAMLCC1A013YTRIV
Investment Grade by Maturity	US Corporate 3–5Y	BAMLCC2A035YTRIV
Investment Grade by Maturity	US Corporate 5–7Y	BAMLCC3A057YTRIV
Investment Grade by Maturity	US Corporate 7–10Y	BAMLCC4A0710YTRIV
Investment Grade by Maturity	US Corporate 10–15Y	BAMLCC7A01015YTRIV
Investment Grade by Maturity	US Corporate 15+Y	BAMLCC8A015PYTRIV
High Yield	US High Yield	BAMLHYH0A0HYM2TRIV
High Yield	BB US High Yield	BAMLHYH0A1BBTRIV
High Yield	B US High Yield	BAMLHYH0A2BTRIV
High Yield	CCC and Lower US High Yield	BAMLHYH0A3CMTRIV
Euro High Yield	Euro High Yield	BAMLHE00EHYITRIV
Emerging Markets Corporate	EM Corporate Plus	BAMLEMCBPITRIV
Emerging Markets Corporate	Asia EM Corporate Plus	BAMLEMRAACRPIASIATRIV
Emerging Markets Corporate	EMEA EM Corporate Plus	BAMLEMRECRPIEMEATRIV
Emerging Markets Corporate	Latin America EM Corporate Plus	BAMLEMRLCRPILATRIV
Emerging Markets Corporate	Euro EM Corporate Plus	BAMLEMEBCRPIETRIV
Emerging Markets Corporate	US EM Corporate Plus	BAMLEMUBCRPIUSTRIV
Emerging Markets Corporate	High Grade EM Corporate Plus	BAMLEMIBHGCRPITRIV
Emerging Markets Corporate	High Yield EM Corporate Plus	BAMLEMHBHYCRPITRIV
Emerging Markets by Rating	AAA–A EM Corporate Plus	BAMLEM1BRRAAA2ACRPITRIV
Emerging Markets by Rating	BBB EM Corporate Plus	BAMLEM2BRRBBBCRPITRIV
Emerging Markets by Rating	BB EM Corporate Plus	BAMLEM3BRRBBBCRPITRIV
Emerging Markets by Rating	B and Lower EM Corporate Plus	BAMLEM4BRRBLCRPITRIV
Emerging Markets by Rating	Crossover EM Corporate Plus	BAMLEM5BCOCPITRIV
Emerging Markets by Sector	Private Financial EM Corporate Plus	BAMLEMFSFCRPITRIV
Emerging Markets by Sector	Public Sector EM Corporate Plus	BAMLEMPBPUBSICRPITRIV
Market Proxy	ICE BofA US Corporate Master	BAMLCC0A0CMTRIV
Risk-free Rate	Effective Federal Funds Rate	DFE

Table 5: Bootstrap Hypothesis Test for γ_0

Model	$\hat{\gamma}_0$	CI Low	CI High	$p(\gamma_0 = 0)$
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000750	-0.000020	0.001323	0.034000
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000261	-0.000034	0.001300	0.655000
LPM ($\tau = 0.02, \alpha = 3.0$)	0.000279	-0.000041	0.001306	0.603000
LPM ($\tau = 0.01, \alpha = 2.0$)	0.000419	-0.000059	0.001309	0.317000
LPM ($\tau = 0.00, \alpha = 2.0$)	0.000332	-0.000034	0.001286	0.490000
LPM ($\tau = 0.00, \alpha = 3.0$)	0.000263	-0.000011	0.001275	0.635000
LPM ($\tau = 0.02, \alpha = 2.0$)	0.000444	-0.000053	0.001295	0.281000
CAPM	0.000459	-0.000048	0.001290	0.250000
LPM ($\tau = 0.02, \alpha = 1.2$)	0.000593	-0.000018	0.001251	0.067000
LPM ($\tau = 0.01, \alpha = 1.2$)	0.000664	-0.000044	0.001255	0.045000
LPM ($\tau = -0.01, \alpha = 1.2$)	0.000745	-0.000076	0.001267	0.039000
LPM ($\tau = -0.01, \alpha = 2.0$)	0.000345	-0.000071	0.001284	0.442000
LPM ($\tau = -0.01, \alpha = 3.0$)	0.000768	-0.000133	0.001390	0.066000
LPM ($\tau = -0.02, \alpha = 2.0$)	0.000199	-0.000404	0.001756	0.763000
LPM ($\tau = -0.02, \alpha = 1.2$)	0.000209	-0.000401	0.001772	0.754000
LPM ($\tau = -0.02, \alpha = 3.0$)	-0.000268	-0.000637	0.001797	0.865462

Table 6: Bootstrap Hypothesis Test for $\gamma_1 = \bar{r}_m$ and Joint Wald Test

Model	$\hat{\gamma}_1$	$\hat{\gamma}_1 - \bar{r}_m$	$p(\gamma_1 = \bar{r}_m)$	W_{joint}	p_{joint}
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000048	-0.000540	0.074000	6.287296	0.131000
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000572	-0.000016	0.989000	7.501771	0.552000
LPM ($\tau = 0.02, \alpha = 3.0$)	0.000562	-0.000026	0.978000	7.988257	0.538000
LPM ($\tau = 0.01, \alpha = 2.0$)	0.000427	-0.000161	0.791000	7.878453	0.423000
LPM ($\tau = 0.00, \alpha = 2.0$)	0.000528	-0.000060	0.949000	8.612471	0.470000
LPM ($\tau = 0.00, \alpha = 3.0$)	0.000550	-0.000038	0.971000	5.658771	0.564000
LPM ($\tau = 0.02, \alpha = 2.0$)	0.000405	-0.000183	0.726000	7.608576	0.418000
CAPM	0.000390	-0.000198	0.682000	7.546600	0.417000
LPM ($\tau = 0.02, \alpha = 1.2$)	0.000247	-0.000341	0.305000	6.547127	0.285000
LPM ($\tau = 0.01, \alpha = 1.2$)	0.000144	-0.000444	0.119000	5.586603	0.200000
LPM ($\tau = -0.01, \alpha = 1.2$)	0.000199	-0.000389	0.216000	8.413171	0.224000
LPM ($\tau = -0.01, \alpha = 2.0$)	0.000600	0.000012	0.994000	7.901641	0.566000
LPM ($\tau = -0.01, \alpha = 3.0$)	0.000643	0.000055	0.960000	11.493352	0.416000
LPM ($\tau = -0.02, \alpha = 2.0$)	0.001277	0.000689	0.826000	3.740713	0.743000
LPM ($\tau = -0.02, \alpha = 1.2$)	0.001330	0.000742	0.821000	4.167850	0.738000
LPM ($\tau = -0.02, \alpha = 3.0$)	0.001624	0.001036	0.730924	3.056798	0.811245

Table 7: Bootstrap Model Performance Comparison

Model	Mean MSE	Mean R^2
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000023	0.347674
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000023	0.339436
LPM ($\tau = 0.02, \alpha = 3.0$)	0.000023	0.339588
LPM ($\tau = 0.01, \alpha = 2.0$)	0.000023	0.341743
LPM ($\tau = 0.00, \alpha = 2.0$)	0.000023	0.337044
LPM ($\tau = 0.00, \alpha = 3.0$)	0.000023	0.329730
LPM ($\tau = 0.02, \alpha = 2.0$)	0.000023	0.337739
CAPM	0.000023	0.337494
LPM ($\tau = 0.02, \alpha = 1.2$)	0.000023	0.328185
LPM ($\tau = 0.01, \alpha = 1.2$)	0.000023	0.327804
LPM ($\tau = -0.01, \alpha = 1.2$)	0.000024	0.313905
LPM ($\tau = -0.01, \alpha = 2.0$)	0.000024	0.307628
LPM ($\tau = -0.01, \alpha = 3.0$)	0.000024	0.287012
LPM ($\tau = -0.02, \alpha = 2.0$)	0.000025	0.270525
LPM ($\tau = -0.02, \alpha = 1.2$)	0.000025	0.271943
LPM ($\tau = -0.02, \alpha = 3.0$)	0.000027	0.222318

Table 8: Original Sample Hypothesis Test for γ_0

Model	$\hat{\gamma}_0$	CI Low	CI High	$p(\gamma_0 = 0)$
LPM ($\tau = -0.01, \alpha = 3.0$)	0.000768	0.000155	0.001381	0.014060
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000750	0.000074	0.001426	0.029706
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000261	-0.000332	0.000854	0.388697
LPM ($\tau = 0.00, \alpha = 3.0$)	0.000263	-0.000324	0.000851	0.379445
LPM ($\tau = 0.02, \alpha = 3.0$)	0.000279	-0.000313	0.000871	0.355403
LPM ($\tau = 0.00, \alpha = 2.0$)	0.000332	-0.000254	0.000919	0.266701
LPM ($\tau = 0.01, \alpha = 2.0$)	0.000419	-0.000176	0.001014	0.167353
CAPM	0.000459	-0.000139	0.001057	0.132228
LPM ($\tau = 0.02, \alpha = 2.0$)	0.000444	-0.000153	0.001040	0.144672
LPM ($\tau = 0.02, \alpha = 1.2$)	0.000593	-0.000038	0.001224	0.065361
LPM ($\tau = 0.01, \alpha = 1.2$)	0.000664	0.000015	0.001312	0.044914
LPM ($\tau = -0.01, \alpha = 2.0$)	0.000345	-0.000281	0.000971	0.279785
LPM ($\tau = -0.01, \alpha = 1.2$)	0.000745	0.000067	0.001424	0.031265
LPM ($\tau = -0.02, \alpha = 3.0$)	-0.000268	-0.001022	0.000485	0.485180
LPM ($\tau = -0.02, \alpha = 2.0$)	0.000199	-0.000718	0.001117	0.670390
LPM ($\tau = -0.02, \alpha = 1.2$)	0.000209	-0.000712	0.001130	0.656517

Table 9: Original Sample Hypothesis Test for $\gamma_1 = \bar{r}_m$ and Joint Wald Test

Model	$\hat{\gamma}_1$	$\hat{\gamma}_1 - \bar{r}_m$	CI Low	CI High	$p(\gamma_1 = \bar{r}_m)$	W_{joint}	p_{joint}
LPM ($\tau = -0.01, \alpha = 3.0$)	0.000643	0.000055	-0.000656	0.000767	0.879109	9.814939	0.007391
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000048	-0.000540	-0.001057	-0.000022	0.040989	5.402518	0.067121
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000572	-0.000016	-0.000570	0.000537	0.954082	1.044521	0.593178
LPM ($\tau = 0.00, \alpha = 3.0$)	0.000550	-0.000038	-0.000578	0.000501	0.889851	0.979280	0.612847
LPM ($\tau = 0.02, \alpha = 3.0$)	0.000562	-0.000026	-0.000577	0.000524	0.925488	1.143535	0.564527
LPM ($\tau = 0.00, \alpha = 2.0$)	0.000528	-0.000060	-0.000598	0.000478	0.826262	1.461885	0.481455
LPM ($\tau = 0.01, \alpha = 2.0$)	0.000427	-0.000161	-0.000675	0.000353	0.539645	1.936780	0.379694
CAPM	0.000390	-0.000198	-0.000697	0.000300	0.435530	2.268457	0.321670
LPM ($\tau = 0.02, \alpha = 2.0$)	0.000405	-0.000183	-0.000685	0.000319	0.475382	2.136966	0.343529
LPM ($\tau = 0.02, \alpha = 1.2$)	0.000247	-0.000341	-0.000812	0.000131	0.156984	3.578166	0.167113
LPM ($\tau = 0.01, \alpha = 1.2$)	0.000144	-0.000444	-0.000924	0.000037	0.070315	4.574195	0.101561
LPM ($\tau = -0.01, \alpha = 2.0$)	0.000600	0.000012	-0.000559	0.000583	0.967848	1.741722	0.418591
LPM ($\tau = -0.01, \alpha = 1.2$)	0.000199	-0.000389	-0.000897	0.000119	0.133113	4.768025	0.092180
LPM ($\tau = -0.02, \alpha = 3.0$)	0.001624	0.001036	-0.000472	0.002544	0.178022	1.848934	0.396743
LPM ($\tau = -0.02, \alpha = 2.0$)	0.001277	0.000689	-0.000424	0.001802	0.225145	3.245369	0.197368
LPM ($\tau = -0.02, \alpha = 1.2$)	0.001330	0.000742	-0.000409	0.001893	0.206239	3.463852	0.176943

Table 10: Original Sample Model Performance Comparison

Model	Mean MSE	Mean R^2
LPM ($\tau = -0.01, \alpha = 3.0$)	0.000020	0.318814
LPM ($\tau = 0.00, \alpha = 1.2$)	0.000021	0.379451
LPM ($\tau = 0.01, \alpha = 3.0$)	0.000021	0.372337
LPM ($\tau = 0.00, \alpha = 3.0$)	0.000021	0.364294
LPM ($\tau = 0.02, \alpha = 3.0$)	0.000021	0.371870
LPM ($\tau = 0.00, \alpha = 2.0$)	0.000021	0.366817
LPM ($\tau = 0.01, \alpha = 2.0$)	0.000021	0.371743
CAPM	0.000021	0.369725
LPM ($\tau = 0.02, \alpha = 2.0$)	0.000021	0.368953
LPM ($\tau = 0.02, \alpha = 1.2$)	0.000022	0.364269
LPM ($\tau = 0.01, \alpha = 1.2$)	0.000022	0.370026
LPM ($\tau = -0.01, \alpha = 2.0$)	0.000022	0.354776
LPM ($\tau = -0.01, \alpha = 1.2$)	0.000023	0.347830
LPM ($\tau = -0.02, \alpha = 3.0$)	0.000024	0.248438
LPM ($\tau = -0.02, \alpha = 2.0$)	0.000026	0.317423
LPM ($\tau = -0.02, \alpha = 1.2$)	0.000027	0.316661

Bibliography

- Bawa, V. S. (1975). Optimal rules for ordering uncertain prospects. *Journal of Financial Economics*, 2(1):95–121.
- Cochrane, J. H. (2005). *Asset Pricing*. Princeton University Press, Princeton, NJ, revised edition edition.
- Fama, E. F. and French, K. R. (2004). The capital asset pricing model: Theory and evidence. *Journal of Economic Perspectives*, 18(3):25–46.
- Fama, E. F. and MacBeth, J. D. (1973). Risk, return, and equilibrium: Empirical tests. *Journal of Political Economy*, 81(3):607–636.
- Federal Reserve Bank of St. Louis (2026). FRED API Documentation. Federal Reserve Economic Data. Accessed: 2026-05-06.
- Fishburn, P. C. (1977). Mean-risk analysis with risk associated with below-target returns. *The American Economic Review*, 67(2):116–126.
- Hall, P. (1992). *The Bootstrap and Edgeworth Expansion*. Springer Series in Statistics. Springer, New York.
- Harlow, W. V. and Rao, R. K. S. (1989). Asset pricing in a generalized mean-lower partial moment framework: Theory and evidence. *Journal of Financial and Quantitative Analysis*, 24(3):285–311.
- Horowitz, J. L. (2001). The bootstrap. In Heckman, J. J. and Leamer, E., editors, *Handbook of Econometrics*, volume 5, pages 3159–3228. Elsevier.
- ICE Data Indices, LLC (2025). ICE Bond Index Methodologies. Accessed: 2026-05-06.
- Markowitz, H. (1952). Portfolio selection. *The Journal of Finance*, 7(1):77–91.
- Politis, D. N. and Romano, J. P. (1994). The stationary bootstrap. *Journal of the American Statistical Association*, 89(428):1303–1313.
- Satchell, S. E. (2001). Lower partial-moment capital asset pricing models: A re-examination. In Sortino, F. A. and Satchell, S. E., editors, *Managing Downside Risk in Financial Markets: Theory, Practice and Implementation*, pages 156–168. Butterworth-Heinemann, Oxford.

- Shanken, J. (1992). On the estimation of beta-pricing models. *The Review of Financial Studies*, 5(1):1–33.
- Sharpe, W. F. (1964). Capital asset prices: A theory of market equilibrium under conditions of risk. *The Journal of Finance*, 19(3):425–442.
- Tobin, J. (1958). Liquidity preference as behavior towards risk. *The Review of Economic Studies*, 25(2):65–86.
- Varian, H. R. (1992). *Microeconomic Analysis*. W. W. Norton & Company, 3rd edition.
- von Neumann, J. and Morgenstern, O. (1944). *Theory of Games and Economic Behavior: 60th Anniversary Commemorative Edition*. Princeton University Press, Princeton, NJ. With an introduction by Harold W. Kuhn and an afterword by Ariel Rubinstein.