



American Economist
Journal of Economics Finance
and Global Policy

Forecasting Stock Prices in a Frontier Capital Market:

A Comparison of Naïve, Drift, ETS, ARIMA and ARIMA–GARCH Models for Listed Joint-Stock Companies in Uzbekistan

Muradova Dildora Abdusalimovna

Senior Lecturer, Department of Finance and Financial Technologies,
Tashkent State University of Economics

e-mail: dildoramuradova1985@gmail.com

Abstract

This study evaluates and compares the out-of-sample forecasting performance of classical and volatility-adjusted time-series models for the equity prices of eight joint-stock companies listed in Uzbekistan: HMKB, SQB, IPKY, ALKB, QZSM, UZTL, URTS and CBSK. Using daily closing prices from 9 September 2022 to 17 April 2026 (883 observations per stock with no missing values), five competing specifications are estimated and assessed: the Naïve (random-walk) benchmark, the Drift model, exponential smoothing (ETS), the autoregressive integrated moving-average model (ARIMA), and an ARIMA–GARCH model with Student-t innovations that explicitly captures volatility clustering. Models are identified on a three-year training window and evaluated on a subsequent six-month hold-out period using the mean absolute error (MAE), the root mean squared error (RMSE) and the mean absolute percentage error (MAPE); the model with the lowest hold-out RMSE is retained for each stock and re-estimated on the full sample to produce quarterly forecasts

for 2026Q3–2028Q4. The augmented Dickey–Fuller test indicates that all price series are integrated of order one, while Engle’s ARCH-LM test detects significant conditional heteroscedasticity in six of the eight return series. The central finding is that no single model dominates: ETS is preferred for three stocks, ARIMA–GARCH for three, and the Drift and ARIMA models for one each. Crucially, the volatility-adjusted model is selected only where ARCH effects are present and never where they are absent, yet simple benchmarks remain highly competitive and occasionally superior even for several volatile series. The forecasts imply predominantly upward trajectories, one flat path and one decline, and several trend-extrapolating forecasts are flagged as economically aggressive. The paper contributes rare model-comparison evidence from an under-researched frontier market and offers cautious, evidence-based guidance for investors, analysts and market-development policy.

Keywords: stock price forecasting; ARIMA; GARCH; exponential smoothing; forecast accuracy; volatility clustering; frontier markets; Uzbekistan capital market.

This work is Licensed under a Creative Commons Attribution 4.0 International License.

JEL classification: C22; C53; G15; G17.

Introduction

The ability to forecast equity prices with reasonable accuracy is of enduring interest to investors, portfolio managers, market regulators and academic researchers. Forecasts inform asset allocation, the construction of hedging strategies, the pricing of derivative instruments and the assessment of market efficiency. They also have a developmental dimension: in markets that are still maturing, credible quantitative tools can support the gradual deepening of price discovery and the cultivation of investor confidence. Yet forecasting financial prices is intrinsically difficult, because asset returns combine a large unpredictable component with subtle, time-varying dependence in both their conditional mean and their conditional variance.

These difficulties are amplified in emerging and frontier capital markets. Such markets are typically characterised by thin and infrequent trading, limited liquidity, pronounced price jumps, episodes of volatility

clustering, occasional structural breaks and a relatively small set of actively traded securities. Thin trading induces stale prices and irregular adjustment, while low liquidity can produce sharp, discontinuous movements that violate the smoothness assumptions underlying many statistical models. In this environment the relative merits of competing forecasting methods are not obvious a priori, and conclusions drawn from large, liquid developed markets need not carry over.

Uzbekistan's equity market is a compelling and largely unexplored empirical setting in which to study these questions. The market is comparatively young and concentrated, dominated by a modest number of listed joint-stock companies spanning banking, industry and utilities; trading is less continuous than in mature exchanges; and the empirical literature on the statistical behaviour and predictability of its share prices is sparse. These features make it both a challenging laboratory for forecasting models and a setting in which even descriptive, model-comparison evidence has practical value for a developing investor base.

Against this background, the present study asks whether more sophisticated, volatility-sensitive models deliver tangible forecasting gains over simple benchmarks for Uzbek-listed shares, or whether parsimonious methods remain competitive. Comparing the Naïve and Drift benchmarks with exponential smoothing (ETS), ARIMA and ARIMA–GARCH is informative precisely because the models embody progressively richer assumptions. The Naïve model treats the most recent price as the best guess of all future prices; the Drift model superimposes the average historical trend; ETS introduces adaptive level and trend components through exponential weighting; ARIMA models linear autocorrelation in the differenced series; and ARIMA–GARCH augments the conditional mean with an explicit, time-varying conditional variance, thereby accommodating the volatility clustering that is a stylised fact of financial returns. Establishing which of these mechanisms actually improves accuracy and for which securities is the empirical heart of the paper.

This study addresses five research questions. First, which forecasting model delivers the most accurate price forecasts for the selected Uzbek listed companies? Second, do more complex models such as ARIMA–GARCH outperform simpler alternatives (Naïve, Drift, ETS, ARIMA)? Third,

is forecasting performance heterogeneous across stocks? Fourth, what do the final 2026Q3–2028Q4 forecasts imply for the expected price dynamics of each security? Fifth, what practical implications follow for investors, portfolio managers, analysts and policymakers operating in a frontier market?

The paper makes three contributions. **First**, it provides systematic out-of-sample model-comparison evidence from Uzbekistan, a frontier equity market that is almost absent from the international forecasting literature, thereby extending the geographic reach of that literature to a new and policy-relevant setting. **Second**, it places simple benchmarks and advanced volatility-sensitive models within a single, transparent train–test framework, allowing a like-for-like assessment of whether the additional complexity of ETS, ARIMA and ARIMA–GARCH is empirically justified rather than merely assumed. **Third**, it links model choice to the underlying statistical properties of each series in particular the presence or absence of ARCH effects and reports forecasts together with an honest account of their diagnostic limitations, offering a methodologically careful template for applied forecasting in thin markets.

The remainder of the paper is organised as follows. Section 2 reviews the relevant literature. Section 3 describes the data and sets out the methodology, including all model and accuracy formulae. Section 4 reports the empirical results stock by stock. Section 5 presents the 2026Q3–2028Q4 forecasts and discusses them. Section 6 examines robustness and the diagnostic properties of the volatility models. Section 7 draws out practical implications, and Section 8 concludes.

2. Literature Review

2.1 Market efficiency and the limits of predictability

The intellectual starting point for any forecasting exercise in finance is the efficient market hypothesis. Fama (1970), building on his earlier analysis of the behaviour of stock-market prices (Fama, 1965), argued that in an informationally efficient market prices fully reflect available information, so that successive price changes are largely unpredictable and approximate a random walk. Under the weak form of the hypothesis, past prices contain no exploitable information about future prices, which provides a strong rationale

for treating the Naïve random-walk forecast as the natural benchmark against which all other models must prove their worth. Subsequent work qualified this view: Pesaran and Timmermann (1995) documented statistically detectable and economically meaningful predictability in U.S. stock returns conditional on publicly available variables, while Welch and Goyal (2008) cautioned that much in-sample predictability fails to survive out-of-sample evaluation. This tension between detectable structure and fragile out-of-sample gains motivates the careful hold-out design adopted here.

2.2 Linear time-series models in financial markets

The modern practice of univariate time-series forecasting was codified by Box and Jenkins (1970), whose autoregressive integrated moving-average (ARIMA) framework provides a systematic identification–estimation–diagnostic methodology for stationary and difference-stationary series. ARIMA models capture linear autocorrelation parsimoniously and remain a standard benchmark in applied forecasting (Granger and Newbold, 1986; Brooks, 2019). Exponential smoothing constitutes a complementary tradition; its modern, fully statistical formulation as a family of innovations state-space models is developed by Hyndman, Koehler, Ord and Snyder (2008) and is operationalised through automatic model-selection algorithms (Hyndman and Khandakar, 2008). The accessible synthesis in Hyndman and Athanasopoulos (2021) has made the ETS and ARIMA families, together with the Naïve and Drift benchmarks, the canonical toolkit for comparative forecasting studies precisely the toolkit deployed in this paper.

2.3 Volatility clustering and GARCH-type models

A defining feature of financial returns, first emphasised by Mandelbrot (1963), is that large changes tend to be followed by large changes and small by small volatility clustering accompanied by heavy-tailed, leptokurtic return distributions. Engle (1982) introduced the autoregressive conditional heteroscedasticity (ARCH) model to capture time-varying conditional variance, and Bollerslev (1986) generalised it to the parsimonious GARCH framework that has since become standard. Because empirical returns are fat-tailed, Bollerslev (1987) proposed estimating GARCH models under a Student-t distribution, an approach adopted in the volatility models analysed here. The persistence of conditional variance, often close to unity, is

documented by Engle and Bollerslev (1986), while asymmetric responses to good and bad news are modelled by Nelson (1991). Combining an ARIMA (or ARFIMA) conditional mean with a GARCH conditional variance yields the ARIMA–GARCH model, which can in principle sharpen interval forecasts and, where return dependence interacts with volatility dynamics, point forecasts as well. Comprehensive reviews of the broader volatility-forecasting literature are provided by Poon and Granger (2003) and Andersen, Bollerslev, Diebold and Labys (2003); the stylised facts of asset returns are catalogued by Cont (2001) and the textbook treatment by Tsay (2010).

2.4 Forecast evaluation and the comparison of accuracy

Rigorous model comparison requires well-chosen accuracy criteria and an out-of-sample protocol. The mean absolute error, the root mean squared error and the mean absolute percentage error remain the most widely used scale-dependent and scale-independent measures, and their properties and pitfalls are analysed by Hyndman and Koehler (2006). Diebold and Mariano (1995) provide a formal test of equal predictive accuracy, and the large-scale forecasting competitions surveyed by Makridakis and Hibon (2000) and Makridakis, Spiliotis and Assimakopoulos (2018) repeatedly deliver a sobering lesson: statistically sophisticated methods do not automatically outperform simple ones, and parsimonious benchmarks are frequently hard to beat out of sample. This evidence frames the central question of the present study and counsels against assuming that ARIMA–GARCH must dominate.

2.5 Forecasting in emerging and frontier markets

Emerging and frontier equity markets exhibit higher and more persistent volatility, partial segmentation and evolving informational efficiency, as documented by Harvey (1995), Bekaert and Harvey (1997) and the survey of Kearney (2012). These features make volatility modelling potentially valuable but also complicate estimation, because short samples, thin trading and structural change strain the assumptions of both linear and conditional-variance models. Hybrid and machine-learning approaches have been proposed to capture nonlinearity (Zhang, 2003; Atsalakis and Valavanis, 2009), but the comparative performance of classical statistical models in genuinely frontier settings remains under-documented. Empirical

evidence specifically on the forecasting of Uzbek-listed equities is, to the best of our knowledge, essentially absent from the international literature [local empirical references on the Uzbekistan capital market citation needed; to be supplied by the author]. The present study helps to fill this gap.

3. Data and Methodology

3.1 Data

The empirical analysis uses daily closing prices for eight joint-stock companies listed in Uzbekistan, identified by their tickers HMKB, SQB, IPKY, ALKB, QZSM, UZTL, URTS and CBSK. For each security the sample comprises 883 daily observations spanning 9 September 2022 to 17 April 2026, with no missing values. The data are recorded at daily (trading-day) frequency, as confirmed by the predominantly one-day spacing between consecutive dates in the cleaned data set. The price series are expressed in their quoted nominal units; the precise unit of quotation (for example, Uzbek so'm per share) and the source exchange should be stated explicitly by the author in the final version, as this information is not contained in the underlying output files.

Descriptive statistics reveal substantial heterogeneity in price level and range across the eight stocks, which is itself relevant for model comparison because scale-dependent accuracy measures cannot be compared directly across securities. Table 1 summarises the number of observations, the minimum, maximum and mean price, and the first and last observed prices for each stock.

Table 1. Description of the analysed stocks (daily closing prices, 2022-09-09 to 2026-04-17)

Ticker	Obs.	Missing	Min	Max	Mean	First	Last
HMKB	883	0	10.50	59.98	27.72	11.67	56.00
SQB	883	0	8.600	39.90	11.84	11.49	36.00
IPKY	883	0	10.04	426.00	98.91	11.93	143.99
ALKB	883	0	0.370	1.080	0.611	0.590	1.080
QZSM	883	0	708.00	3,890.0	2,103.9	3,555.0	1,449.0
UZTL	883	0	3,480.0	9,400.0	6,270.5	7,000.0	7,039.9

Ticker	Obs.	Missing	Min	Max	Mean	First	Last
URTS	883	0	2,800.0	8,720.0	4,058.1	3,800.0	8,720.0
CBSK	883	0	1.000	3.800	2.172	1.220	3.360

Source: authors' computations. Prices in quoted nominal units (unit to be confirmed by the author). "First" = 2022-09-09; "Last" = 2026-04-17.

3.2 Train–test design and forecast horizon

Following standard practice for out-of-sample evaluation (Hyndman and Athanasopoulos, 2021), the sample is partitioned chronologically. Models are identified and estimated on a three-year training window covering 9 September 2022 to 8 September 2025 (732 observations) and are then evaluated on a six-month hold-out period running from 9 September 2025 to 6 March 2026 (124 observations). A short residual tail of 27 observations (10 March 2026 to 17 April 2026) completes the full sample and is retained when models are re-estimated to generate forward forecasts. For each stock the model with the lowest hold-out RMSE is selected as the final specification and re-estimated on the full sample; quarterly forecasts are then produced for the ten quarters spanning 2026Q3 to 2028Q4.

It is important to distinguish two estimation stages. At the selection stage, each model family is fitted to the training window and judged on the hold-out period. At the forecasting stage, the selected family is re-estimated on the complete sample so that the forecasts exploit all available information. Because automatic identification is re-run on the longer sample, the precise orders or components chosen at the two stages can differ even though the model family is unchanged; this is a normal feature of the procedure rather than an inconsistency, and the specific differences are documented in Section 5 and in the Author Notes.

3.3 Stationarity and ARCH testing

Prior to estimation, the order of integration of each series is assessed with the augmented Dickey–Fuller (ADF) test (Dickey and Fuller, 1979), applied to both the price level and the log return. Conditional heteroscedasticity in the return series is examined with Engle's (1982) Lagrange-multiplier (ARCH-LM) test. The results, reported in Table 2, are unambiguous on the first question and informative on the second. For every stock the null of a unit root cannot be rejected in the price level (p-values

from 0.225 for CBSK to 0.990 for URTS), whereas it is decisively rejected for the log returns (p-values at the 0.01 reporting floor throughout). All eight price series are therefore treated as integrated of order one, which justifies first-differencing within the ARIMA specifications and the modelling of returns within the ARIMA–GARCH framework.

Table 2. Unit-root and ARCH diagnostics

Ticker	ADF (level) stat.	ADF (level) p	ADF (ret.) stat.	ADF (ret.) p	ARCH- LM p	ARCH?
HMKB	-0.864	0.956	-10.05	0.010	5.57e-5	Yes
SQB	-0.978	0.942	-7.032	0.010	2.26e-25	Yes
IPKY	-1.794	0.666	-10.75	0.010	1.000	No
ALKB	-0.750	0.966	-9.493	0.010	9.73e-23	Yes
QZSM	-0.645	0.975	-10.06	0.010	6.89e-8	Yes
UZTL	-0.561	0.979	-9.472	0.010	3.32e-8	Yes
URTS	1.497	0.990	-9.836	0.010	0.312	No
CBSK	-2.834	0.225	-10.06	0.010	4.87e-25	Yes

Source: authors' computations. ADF p-values are reported by the test routine with a lower bound of 0.01. "ARCH?" indicates rejection of the no-ARCH null at the 5% level. ARCH-LM statistics: 40.66 (HMKB), 147.04 (SQB), 0.46 (IPKY), 133.99 (ALKB), 57.33 (QZSM), 59.08 (UZTL), 13.83 (URTS), 145.39 (CBSK).

The ARCH-LM test reveals significant volatility clustering in six of the eight return series HMKB, SQB, ALKB, QZSM, UZTL and CBSK at conventional levels, frequently with extremely small p-values. For two stocks, IPKY ($p \approx 1.000$) and URTS ($p = 0.312$), the test provides no evidence of ARCH effects. This pattern is consequential: it implies that an explicit conditional-variance model is statistically warranted for some, but not all, of the securities, and it anticipates the finding in Section 4 that the volatility-adjusted model is never selected for the two stocks without ARCH effects.

3.4 Forecasting models

Five forecasting models are estimated and compared. Their specifications are stated below; a compact summary is provided in Table 3.

3.4.1 Naïve model

The Naïve forecast sets every future value equal to the most recently observed price:

$$\hat{y}_{t+h|t} = y_t \quad (1)$$

This is the random-walk benchmark implied by the weak-form efficient market hypothesis. It is deliberately uninformative about future dynamics and serves as the minimal standard that any richer model should surpass.

3.4.2 Drift model

The Drift model extends the Naïve forecast by adding the average historical change over the estimation sample, thereby projecting the prevailing linear trend:

$$\hat{y}_{t+h|t} = y_t + h((y_t - y_1) / (t - 1)) \quad (2)$$

The drift term equals the slope of the line connecting the first and last observations of the estimation window, so the model is equivalent to a random walk with an estimated deterministic trend.

3.4.3 Exponential smoothing (ETS)

The ETS family expresses the series through unobserved level and, where appropriate, trend components that are updated by exponentially weighted smoothing (Hyndman, Koehler, Ord and Snyder, 2008). For a local-trend specification the components evolve as

$$\ell_t = \alpha y_t + (1 - \alpha)(\ell_{t-1} + b_{t-1}) \quad (3a)$$

$$b_t = \beta(\ell_t - \ell_{t-1}) + (1 - \beta)b_{t-1} \quad (3b)$$

where α and β are the level and trend smoothing parameters. The error, trend and seasonal components may each be additive (A), multiplicative (M) or absent (N), giving specifications such as ETS(A,N,N), ETS(M,N,N) and ETS(M,A,N). The actual forms identified for each stock reported in Sections 4 and 5 were selected automatically by information criteria; the additive-error, no-trend case collapses to simple exponential

smoothing, while multiplicative-error variants accommodate variance that scales with the level.

3.4.4 ARIMA model

The autoregressive integrated moving-average model of order (p, d, q) is written compactly using the lag operator L as

$$\varphi(L)(1 - L)^d y_t = c + \theta(L) \varepsilon_t \quad (4)$$

where $\varphi(L) = 1 - \varphi_1 L - \dots - \varphi_p L^p$ is the autoregressive polynomial, $\theta(L) = 1 + \theta_1 L + \dots + \theta_q L^q$ is the moving-average polynomial, d is the order of differencing, c is an optional constant (drift), and ε_t is a white-noise innovation. Because the ADF tests place every price series in the I(1) class, d = 1 throughout. The specific orders identified for each stock are reported in Section 4.

3.4.5 ARIMA–GARCH model

The ARIMA–GARCH model couples a conditional-mean equation of ARMA/ARFIMA type with a conditional-variance equation of GARCH type, estimated jointly. The mean equation for the return is

$$r_t = \mu + \sum \varphi_i r_{t-i} + \sum \theta_j \varepsilon_{t-j} + \varepsilon_t \quad (5)$$

the innovation is decomposed into a standardised shock and a time-varying conditional standard deviation,

$$\varepsilon_t = z_t \sqrt{h_t} \quad (6)$$

and the conditional variance follows the parsimonious sGARCH(1,1) process

$$h_t = \omega + \alpha_1 \varepsilon_{t-1}^2 + \beta_1 h_{t-1} \quad (7)$$

where $\omega > 0$, $\alpha_1 \geq 0$ and $\beta_1 \geq 0$. The coefficient α_1 measures the short-run impact of a recent shock on current volatility, β_1 measures volatility persistence, and the sum $\alpha_1 + \beta_1$ governs how slowly volatility shocks decay; values close to one imply highly persistent, near-integrated volatility. In line with the leptokurtosis of financial returns (Bollerslev, 1987), all ARIMA–GARCH models in this study are estimated under a Student-t

distribution for z_t , whose estimated degrees-of-freedom (shape) parameter quantifies tail thickness. Estimation is by maximum likelihood.

3.4.6 Forecast accuracy metrics

Out-of-sample accuracy is measured on the six-month hold-out period using three complementary criteria. The mean absolute error and the root mean squared error are scale-dependent,

$$MAE = \left(\frac{1}{n}\right) \sum |y_t - \hat{y}_t| \quad (8)$$

$$RMSE = \sqrt{\left[\left(\frac{1}{n}\right) \sum (y_t - \hat{y}_t)^2\right]} \quad (9)$$

while the mean absolute percentage error is scale-independent and therefore comparable across stocks of very different price levels

$$MAPE = (100/n) \sum |(y_t - \hat{y}_t) / y_t| \quad (10)$$

For all three measures, lower values denote greater accuracy. RMSE penalises large errors more heavily than MAE because deviations are squared, which makes it sensitive to the occasional sharp price moves that characterise thin markets; for this reason, and to maintain a single consistent rule, RMSE is used as the primary criterion for model selection, with MAE and MAPE reported alongside as supporting evidence.

Table 3. Summary of the five forecasting models

Model	Specification	What it captures
Naïve	$\hat{y}_{t+h} = y_t$ (random walk)	Weak-form benchmark; last price carried forward
Drift	$\hat{y}_{t+h} = y_t + h \cdot (y_t - y_1) / (t - 1)$	Random walk with estimated linear trend
ETS	Exponentially weighted level/trend (A/M/N components)	Adaptive level and local trend; smoothing parameters α, β

Model	Specification	What it captures
ARIMA(p,1,q)	$\varphi(L)(1-L)y_t = c + \theta(L)\varepsilon_t$	Linear autocorrelation in the differenced (return) series
ARIMA–GARCH	ARFIMA mean + sGARCH(1,1) variance, Student-t	Conditional mean plus time-varying conditional variance (volatility clustering, fat tails)

4. Empirical Results

This section reports the model-identification outcomes and the out-of-sample accuracy comparison, followed by a stock-by-stock discussion. Table 4 records the model orders selected on the training window for the three parametric families. Table 5 presents the full accuracy comparison across all five models and eight stocks, sorted within each stock by RMSE. Figure 1 places the eight price histories on a common normalised scale, and Figure 2 visualises the relative accuracy of the competing models.

Table 4. Model orders identified on the training window (2022-09-09 to 2025-09-08)

Ticker	ARIMA(p,d,q)	ETS form	ARIMA–GARCH mean
HMKB	ARIMA(0,1,2)	ETS(M,A,N)	sGARCH(1,1)+ARFIMA(1,0,1)
SQB	ARIMA(3,1,2)	ETS(M,N,N)	sGARCH(1,1)+ARFIMA(1,0,3)
IPKY	ARIMA(2,1,2)	ETS(M,A,N)	sGARCH(1,1)+ARFIMA(2,0,2)
ALKB	ARIMA(1,1,2)	ETS(M,N,N)	sGARCH(1,1)+ARFIMA(0,0,1)
QZSM	ARIMA(4,1,3)+drift	ETS(M,N,N)	sGARCH(1,1)+ARFIMA(2,0,3)
UZTL	ARIMA(0,1,2)	ETS(A,N,N)	sGARCH(1,1)+ARFIMA(0,0,2)
URTS	ARIMA(0,1,1)	ETS(M,N,N)	sGARCH(1,1)+ARFIMA(3,0,2)
CBSK	ARIMA(1,1,1)	ETS(A,N,N)	sGARCH(1,1)+ARFIMA(2,0,3)

Source: authors' computations. All ARIMA models use $d = 1$, consistent with the $I(1)$ finding; all ARIMA–GARCH models use a Student-t sGARCH(1,1) variance equation. These are the training-stage specifications used for the

accuracy comparison; final forecasting models are re-estimated on the full sample (Section 5).

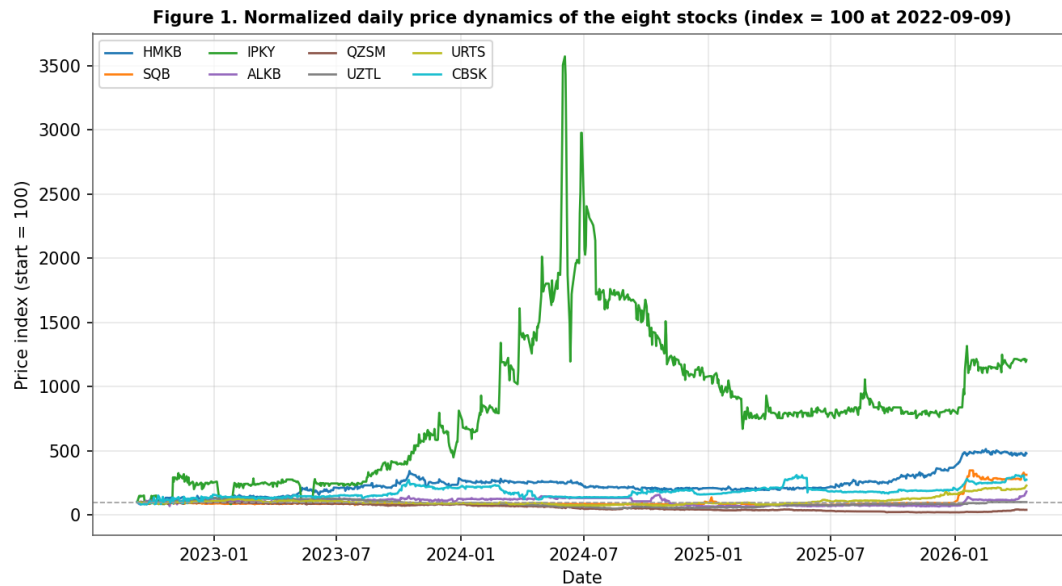


Figure 1. Normalised daily price dynamics of the eight stocks

Each series is indexed to 100 at the start of the sample (9 September 2022). The figure shows marked cross-sectional heterogeneity in cumulative price behaviour, with several banking and industrial names rising strongly while QZSM declines over the window. Source: authors' computations from the cleaned daily data.

4.1 Out-of-sample accuracy comparison

Table 5 reports the MAE, RMSE and MAPE for every model and stock on the six-month hold-out period. Because price levels differ by several orders of magnitude across securities, the absolute error magnitudes are not comparable across rows; the meaningful comparison is within each stock, and the scale-free MAPE provides the cross-stock perspective. The model achieving the lowest RMSE within each stock is shown in bold and is the specification carried forward to the forecasting stage.

Table 5. Out-of-sample accuracy comparison (six-month hold-out; sorted by RMSE within each stock)

Ticker	Model	MAE	RMSE	MAPE (%)
HMKB	ETS	11.61	14.86	23.36
	Drift	11.92	15.18	24.03
	ARIMA-GARCH	12.77	16.19	25.79
	ARIMA	13.23	16.72	26.77
	Naive	13.44	16.91	27.25
SQB	ARIMA-GARCH	7.506	12.85	24.49
	Naive	7.489	12.86	24.33
	Drift	7.561	12.94	24.71
	ETS	7.623	12.98	25.12
	ARIMA	8.081	13.30	28.26
IPKY	ETS	13.42	14.96	12.31
	Drift	13.51	16.27	11.68
	ARIMA	15.26	19.97	12.65
	Naive	13.94	21.29	10.92
	ARIMA-GARCH	14.82	24.24	11.25
ALKB	ETS	0.094	0.149	14.48
	ARIMA	0.093	0.149	14.40
	ARIMA-GARCH	0.093	0.150	14.37
	Naive	0.093	0.155	13.95
	Drift	0.100	0.167	14.91

Ticker	Model	MAE	RMSE	MAPE (%)
QZSM	ARIMA-GARCH	113.39	130.28	14.11
	ARIMA	118.97	170.02	13.66
	Drift	120.14	171.33	13.80
	Naive	164.74	179.11	20.72
	ETS	169.36	184.03	21.30
UZTL	ARIMA	284.12	427.80	4.38
	ETS	284.29	428.05	4.38
	Naive	286.71	432.80	4.42
	Drift	345.66	512.49	5.32
	ARIMA-GARCH	625.74	825.36	9.73
URTS	Drift	1,579.5	2,023.3	23.44
	ARIMA	1,613.9	2,063.8	23.97
	ETS	1,613.9	2,063.8	23.97
	Naive	1,619.1	2,068.6	24.06
	ARIMA-GARCH	1,650.4	2,105.1	24.54
CBSK	ARIMA-GARCH	0.280	0.422	9.58
	Drift	0.285	0.427	9.76
	ARIMA	0.362	0.511	12.57
	ETS	0.362	0.511	12.58
	Naive	0.364	0.513	12.65

Source: authors' computations Bold rows mark the lowest-RMSE model retained for forecasting. MAE and RMSE are in each stock's quoted price unit and are not comparable across stocks; MAPE is scale-free.

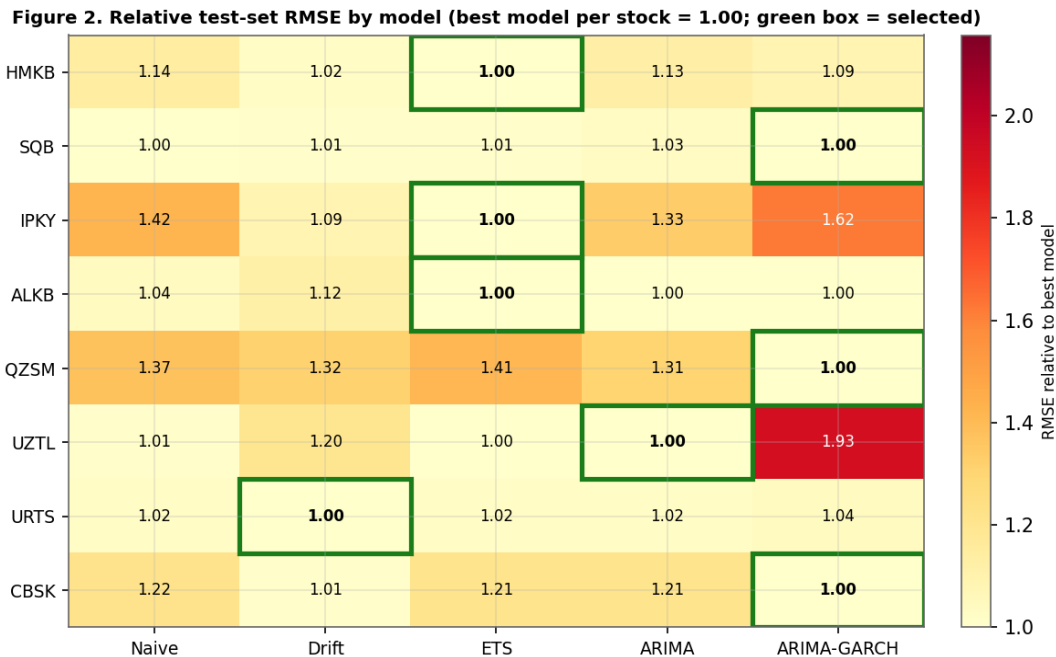


Figure 2. Relative out-of-sample RMSE by model

Each cell reports a model's hold-out RMSE divided by the best (lowest) RMSE for that stock, so the best model equals 1.00 (green box). Values near 1.00 indicate near-equivalent accuracy; larger values indicate worse relative performance. For several stocks (notably ALKB) the simple and advanced models are virtually indistinguishable, whereas ARIMA–GARCH is clearly worse than benchmarks for IPKY and UZTL. Source: authors' computations.

4.2 Stock-by-stock results

4.2.1 ALKB

ALKB is the lowest-priced security in the sample and displays significant ARCH effects. Exponential smoothing is the most accurate model on the hold-out period, with RMSE 0.149, but the result is essentially a four-way tie: ARIMA (RMSE 0.149), ARIMA–GARCH (RMSE 0.150) and the Naïve benchmark (RMSE 0.155) differ from ETS only at the third or fourth decimal place. For ALKB, therefore, the explicit modelling of volatility yields no practical advantage over the simplest possible benchmark—a clear illustration that statistically significant ARCH effects need not translate into superior point forecasts.

4.2.2 CBSK

CBSK exhibits strong ARCH effects, and here the volatility-adjusted model is genuinely preferred: ARIMA–GARCH attains the lowest RMSE (0.422) and the lowest MAPE (9.58%), narrowly ahead of the Drift model (RMSE 0.427). The remaining models ARIMA, ETS and Naïve are noticeably less accurate, clustering around an RMSE of 0.51. CBSK is thus one of the cases in which the conditional-variance machinery earns its keep, although the margin over the trend-based Drift benchmark is modest.

4.2.3 HMKB

HMKB, a mid-priced banking share with significant ARCH effects, is best forecast by ETS (RMSE 14.86), followed closely by Drift (RMSE 15.18). Notably, the ARIMA–GARCH model (RMSE 16.19) is outperformed by both simpler trend models despite the presence of volatility clustering. The adaptive local trend of ETS evidently captures HMKB's drift more effectively than either the linear-trend Drift model or the volatility model.

4.2.4 IPKY

IPKY is the most volatile series in level terms, with a price range from about 10 to 426, yet its returns show no significant ARCH effects (ARCH-LM $p \approx 1.000$). Consistent with this, ETS is the best model (RMSE 14.96) and ARIMA–GARCH is the worst of the five (RMSE 24.24), substantially inferior even to the Naïve benchmark. IPKY is the clearest cautionary case in the sample: imposing a conditional-variance model on a series without ARCH effects degrades, rather than improves, forecast accuracy.

4.2.5 QZSM

QZSM, the only security to fall over the sample, displays significant ARCH effects, and the ARIMA–GARCH model delivers a clear improvement in RMSE (130.28) over the best basic model, ARIMA (RMSE 170.02). Interestingly, the ranking depends on the criterion: on MAPE the basic ARIMA model is marginally better (13.66% versus 14.11%). Because RMSE is the primary criterion and penalises the large errors to which a declining, volatile series is prone, ARIMA–GARCH is selected; this is a case in which volatility modelling demonstrably sharpens squared-error performance.

4.2.6 SQB

SQB is a banking share with very strong ARCH effects. The comparison is exceptionally close at the top: ARIMA–GARCH records the lowest RMSE (12.85) by the narrowest of margins over the Naïve benchmark

(RMSE 12.86), even though the Naïve model has a marginally lower MAE (7.489 versus 7.506) and MAPE. The selection of ARIMA–GARCH for SQB therefore rests entirely on the RMSE criterion and should be regarded as effectively a statistical dead heat with the random walk an important nuance for honest interpretation.

4.2.7 URTS

URTS is a high-priced, strongly trending series whose returns show no significant ARCH effects (ARCH-LM $p = 0.312$). The Drift model is the most accurate (RMSE 2,023.3), ahead of ARIMA and ETS, while ARIMA–GARCH is again the worst performer (RMSE 2,105.1). The dominance of the simple linear-trend model is intuitive for a series with a pronounced, persistent upward drift and no volatility clustering to exploit.

4.2.8 UZTL

UZTL displays significant ARCH effects, yet the basic ARIMA model is preferred (RMSE 427.80), narrowly ahead of ETS (RMSE 428.05) and the Naïve benchmark. Strikingly, ARIMA–GARCH is by far the worst performer for UZTL (RMSE 825.36, roughly twice the best model's error, with a MAPE of 9.73% against around 4.4% for the leading models). Here the presence of ARCH effects coexists with poor out-of-sample behaviour of the volatility model, a reminder that good in-sample variance dynamics do not guarantee good point forecasts.

4.3 Synthesis: does complexity pay?

Taken together, the eight cases yield a nuanced answer to the paper's central question. Under the RMSE criterion the final model is ETS for three stocks (ALKB, HMKB, IPKY), ARIMA–GARCH for three (CBSK, QZSM, SQB), Drift for one (URTS) and ARIMA for one (UZTL). When the volatility model is excluded from the candidate set, the best basic models are ETS for three stocks (ALKB, HMKB, IPKY), ARIMA for two (QZSM, UZTL), Drift for two (CBSK, URTS) and Naïve for one (SQB). Two regularities stand out. First, model performance is markedly heterogeneous across securities, so a one-size-fits-all model would be inappropriate. Second, the ARIMA–GARCH model is selected only for stocks that exhibit significant ARCH effects and is never selected for the two stocks (IPKY, URTS) without them; indeed, for those two it is the worst performer. Yet the converse does not hold: ARCH effects are present in ALKB, HMKB and UZTL as well, and for these the

simpler ETS or ARIMA models prevail, while even where ARIMA–GARCH wins its margin over the best simple alternative is slim. The evidence thus supports a measured conclusion: volatility modelling can help, but only where volatility clustering is genuinely present, and even then simple benchmarks remain strong competitors. This echoes the broader forecasting-competition literature (Makridakis and Hibon, 2000; Makridakis et al., 2018).

5. Forecasting Results and Discussion

The selected model for each stock was re-estimated on the full sample and used to generate quarterly forecasts for 2026Q3–2028Q4. Table 6 summarises the final forecasting models and their key estimated parameters; Table 7 reports the forecast at the start and end of the horizon together with the implied direction and cumulative change. Figure 3 displays the ten-quarter forecast path for each stock on its own scale.

Table 6. Final forecasting models (re-estimated on the full sample) and key parameters

Ticker	Final model	Estimated parameters / notes
HMKB	ETS(M,A,N)	$\alpha = 0.6441, \beta = 0.0023, \sigma = 0.0388$; AIC = 6012.68; in-sample MAPE 2.59%
SQB	ARIMA-GARCH	sGARCH(1,1)+ARFIMA(1,0,2), Student-t ($v \approx 2.80$); $\alpha_1 = 0.4303, \beta_1 = 0.5687$ ($\alpha + \beta = 0.999$); logL=-2004.05; AIC=4.563
IPKY	ETS(M,A,N)	$\alpha = 0.9999, \beta = 0.0020, \sigma = 0.1102$; AIC=9766.12; in-sample MAPE 5.70%
ALKB	ETS(M,A,N)	$\alpha = 0.5239, \beta = 0.0041, \sigma = 0.0549$; AIC=-50.65; in-sample MAPE 3.56%
QZSM	ARIMA-GARCH	sGARCH(1,1)+ARFIMA(3,0,2), Student-t ($v \approx 3.19$); $\alpha_1 = 0.1762, \beta_1 = 0.8003$ ($\alpha + \beta = 0.976$); logL=-2269.82; AIC=5.170
UZTL	ARIMA(0,1,2)	$\theta_1 = -0.2386, \theta_2 = -0.1427; \sigma^2 = 27,673$; AIC=11,528.4; no drift $\Rightarrow$ flat forecast; in-sample MAPE 1.97%

Ticker	Final model	Estimated parameters / notes
URTS	Drift (RW + drift)	daily drift slope positive; 80%/95% forecast intervals reported
CBSK	ARIMA-GARCH	sGARCH(1,1)+ARFIMA(2,0,3), Student-t ($\nu \approx 3.10$); $\alpha_1 = 0.4446$, $\beta_1 = 0.5544$ ($\alpha + \beta = 0.999$); $\log L = -2228.21$; AIC=5.075

Source: authors' computations. α , β are ETS smoothing parameters; α_1 , β_1 are GARCH parameters; ν is the Student-t shape (degrees of freedom). Where the full-sample model order differs from the training-window order in Table 4 (QZSM, SQB, and ALKB's ETS form), this reflects automatic re-identification on the longer sample; see the Author Notes.

Figure 3. Final quarterly price forecast paths, 2026Q3-2028Q4 (each panel on its own scale)

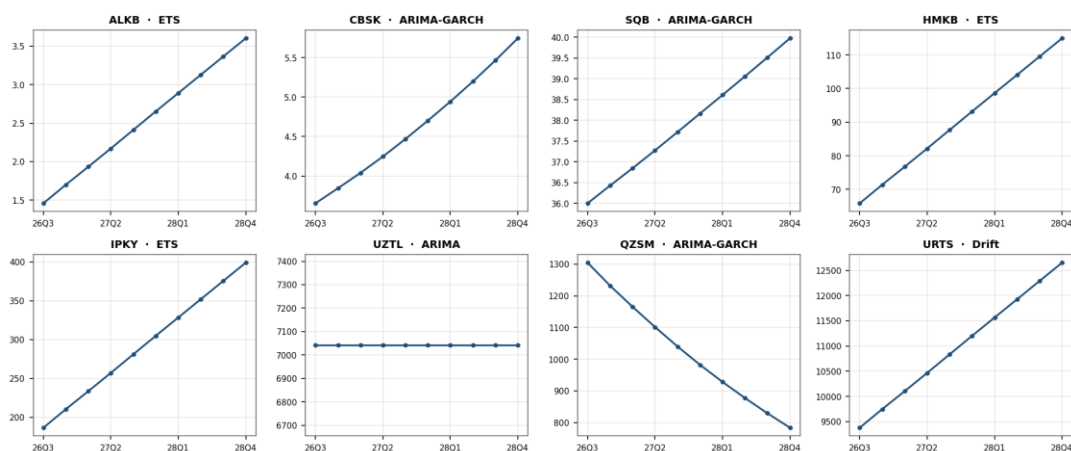


Figure 3. Final quarterly price-forecast paths, 2026Q3–2028Q4

Each panel shows the selected model's quarterly point forecast for one stock, plotted on its own vertical scale because price levels differ widely. The trajectories are predominantly upward; QZSM declines; and UZTL is flat because its ARIMA(0,1,2) specification contains no drift term. Source: authors' computations (quarterly_final_forecast_2026Q3_2028Q4.xlsx).

Table 7. Final forecast results for 2026Q3–2028Q4

Ticker	Final model	2026Q3	2028Q4	Δ over horizon	Direction
HMKB	ETS	65.86	114.95	+74.5%	Upward

Ticker	Final model	2026Q3	2028Q4	Δ over horizon	Direction
SQB	ARIMA-GARCH	36.00	39.97	+11.0%	Upward
IPKY	ETS	186.66	398.91	+113.7%	Upward
ALKB	ETS	1.459	3.603	+147.0%	Upward
QZSM	ARIMA-GARCH	1,304.3	783.94	-39.9%	Downward
UZTL	ARIMA	7,041.6	7,041.6	+0.0%	Flat
URTS	Drift	9,378.2	12,652.7	+34.9%	Upward
CBSK	ARIMA-GARCH	3.655	5.744	+57.2%	Upward

Source: authors' computations. 2026Q3 and 2028Q4 are the first and last quarterly point forecasts; Δ is the cumulative percentage change between them, in the quoted price unit of each stock.

5.1 Forecast trajectories by stock

Six of the eight stocks are forecast to rise over the horizon, one is broadly flat, and one declines. ALKB is projected to increase from about 1.459 in 2026Q3 to 3.603 in 2028Q4 under its ETS(M,A,N) specification, an upward path of roughly 147%. CBSK rises from 3.655 to 5.744 (about +57%) under ARIMA–GARCH. HMKB increases from 65.86 to 114.95 (about +75%) and IPKY from 186.66 to 398.91 (about +114%), both under ETS(M,A,N). SQB shows a mild upward drift from 36.00 to 39.97 (about +11%), and URTS continues its linear ascent from 9,378.2 to 12,652.7 (about +35%) under the Drift model.

Two cases warrant particular comment. QZSM is the sole declining forecast, falling from 1,304.3 in 2026Q3 to 783.94 in 2028Q4 (about -40%); the negative autoregressive structure of its full-sample ARFIMA(3,0,2) mean produces mean-reverting dynamics that, given recent levels, imply continued downward pressure, and the 2026Q3 forecast already sits below the last observed price. UZTL is essentially flat at 7,041.6 across all ten quarters; this is a direct consequence of its selected ARIMA(0,1,2) specification, which

contains no drift term and therefore collapses to a constant (random-walk) forecast equal to the final estimated level. In effect, the model chosen for UZTL behaves like a Naïve forecast over the horizon, a point worth making explicit to readers.

5.2 Cautions on the trend-extrapolating forecasts

Several of the upward forecasts should be interpreted with care. The ETS(M,A,N) and Drift specifications project the locally estimated trend forward linearly, so over a ten-quarter horizon they can imply large cumulative changes. The most aggressive cases are IPKY and ALKB, each forecast to more than double (approximately +114% and +147% respectively), with HMKB also rising substantially (+75%). Such trajectories reflect the mechanical extrapolation of a recent trend rather than any structural forecast of company fundamentals, and they ignore the possibility of trend reversal, mean reversion or regime change that is common in thin markets. These forecasts are therefore best read as conditional, model-based projections of the prevailing momentum, not as point predictions of realised prices; the author is advised to verify the underlying price scale and to consider damping the trend or reporting the forecasts with explicit uncertainty bands in the final version. Among the volatility models, the QZSM ARIMA–GARCH forecast inherits the high estimated volatility persistence documented in Section 6, so its central path is surrounded by wide implied intervals.

5.3 Discussion

The forecasting results reinforce the heterogeneity theme of Section 4 and connect it to the economics of a frontier market. Different stocks require different models because they differ in the speed of price adjustment, the strength and stability of trends, the intensity of volatility clustering and the degree of company-specific risk all of which are shaped by liquidity and market microstructure. Trending, liquid names with no ARCH effects (URTS) are best served by a simple drift; series with strong, exploitable volatility dynamics (CBSK, QZSM, SQB) benefit from explicit conditional-variance modelling; and series whose movement is dominated by an adaptive local trend (ALKB, HMKB, IPKY) are captured well by exponential smoothing. Importantly, the evidence does not support the proposition that advanced models are universally superior. ARIMA–GARCH wins only where it

should where ARCH effects are present and even then by narrow margins, while it can be markedly inferior when applied to series without volatility clustering. For practitioners in Uzbekistan's market, the practical message is that model choice should be guided by the statistical properties of each security, with simple, robust benchmarks retained as a disciplined point of comparison.

6. Robustness and Diagnostic Analysis

This section scrutinises the three ARIMA–GARCH models that were selected as final specifications (for CBSK, QZSM and SQB), drawing on the full battery of diagnostics reported by the estimation routine. The aim is to assess whether the fitted volatility models are statistically adequate and to disclose their limitations honestly. Table 8 collects the key diagnostics; Figure 4 plots the estimated conditional volatility for these three stocks over the hold-out window.

Table 8. Diagnostic summary for the final ARIMA–GARCH models

Diagnostic	CBSK	QZSM	SQB
Mean specification	ARFIMA(2,0,3)	ARFIMA(3,0,2)	ARFIMA(1,0,2)
Variance model	sGARCH(1,1)	sGARCH(1,1)	sGARCH(1,1)
Distribution (Student-t ν)	$t, \nu \approx 3.10$	$t, \nu \approx 3.19$	$t, \nu \approx 2.80$
ARCH coeff. α_1	0.4446	0.1762	0.4303
GARCH coeff. β_1	0.5544	0.8003	0.5687
Persistence $\alpha_1 + \beta_1$	0.999	0.976	0.999
LB std. resid. p (3 lags)	0.94 / 1.00 / 0.999	0.96 / 0.040 / 0.29	0.96 / 0.999 / 0.82
LB sq. resid. p (3 lags)	0.85 / 0.999 / 1.00	0.75 / 0.89 / 0.96	0.86 / 0.40 / 0.61
ARCH-LM p (3 lags)	0.90 / 0.996 / 1.00	0.54 / 0.83 / 0.92	0.41 / 0.68 / 0.84

Diagnostic	CBSK	QZSM	SQB
Nyblom joint (1% crit.)	8.24 (3.05)	4.08 (3.05)	2.98 (2.59)
Pearson GoF (sig. groupings)	g=40: 0.010 (others n.s.)	g=20: 0.047 (others n.s.)	g=20: 0.019; g=40: 0.002; g=50: 0.006
AIC / log-likelihood	5.075 / -2228.2	5.170 / -2269.8	4.563 / -2004.0

Source: authors' computations. LB = weighted Ljung–Box test; p-values reported at three representative lags. ARCH-LM at lags 3/5/7. Nyblom joint statistic compared with the 1% asymptotic critical value. Pearson goodness-of-fit reports the grouping(s) at which the null of correct distribution is rejected at 5%.

6.1 Residual adequacy

On the criteria that matter most for a well-specified GARCH model, the three final specifications perform creditably. The weighted Ljung–Box tests on the standardised residuals indicate that serial correlation in the conditional mean has been largely removed: p-values are high across lags for CBSK and SQB, and for QZSM the only marginal exception is a single intermediate lag ($p \approx 0.04$), with the longer-lag test comfortably non-significant. The Ljung–Box tests on the squared standardised residuals and the weighted ARCH-LM tests are non-significant at all reported lags for all three stocks, implying that the sGARCH(1,1) variance equation has successfully absorbed the conditional heteroscedasticity present in the raw returns. The sign-bias tests are non-significant throughout, providing no evidence of asymmetric (leverage) effects that a symmetric sGARCH specification would miss. To this extent, the volatility models are statistically adequate.

6.2 Volatility persistence and heavy tails

Two features of the estimates deserve emphasis. First, volatility is highly persistent in every case: the sum $\alpha_1 + \beta_1$ equals 0.999 for both CBSK and SQB and 0.976 for QZSM, indicating near-integrated conditional variance in which shocks to volatility decay very slowly. This persistence is economically meaningful: turbulent periods are expected to be long-lived but it

also implies that long-horizon variance forecasts are dominated by the persistence parameter and are correspondingly imprecise. Second, the estimated Student-t shape parameters are small ($\nu \approx 3.10$ for CBSK, 3.19 for QZSM and 2.80 for SQB), confirming pronounced fat tails. The SQB estimate is particularly notable: with roughly 2.8 degrees of freedom the fitted distribution lies close to the boundary ($\nu = 2$) below which the variance is undefined, signalling extremely heavy tails and counselling caution in interpreting moment-based quantities for that stock.

Figure 4. Estimated conditional volatility dynamics, ARIMA-GARCH finalists

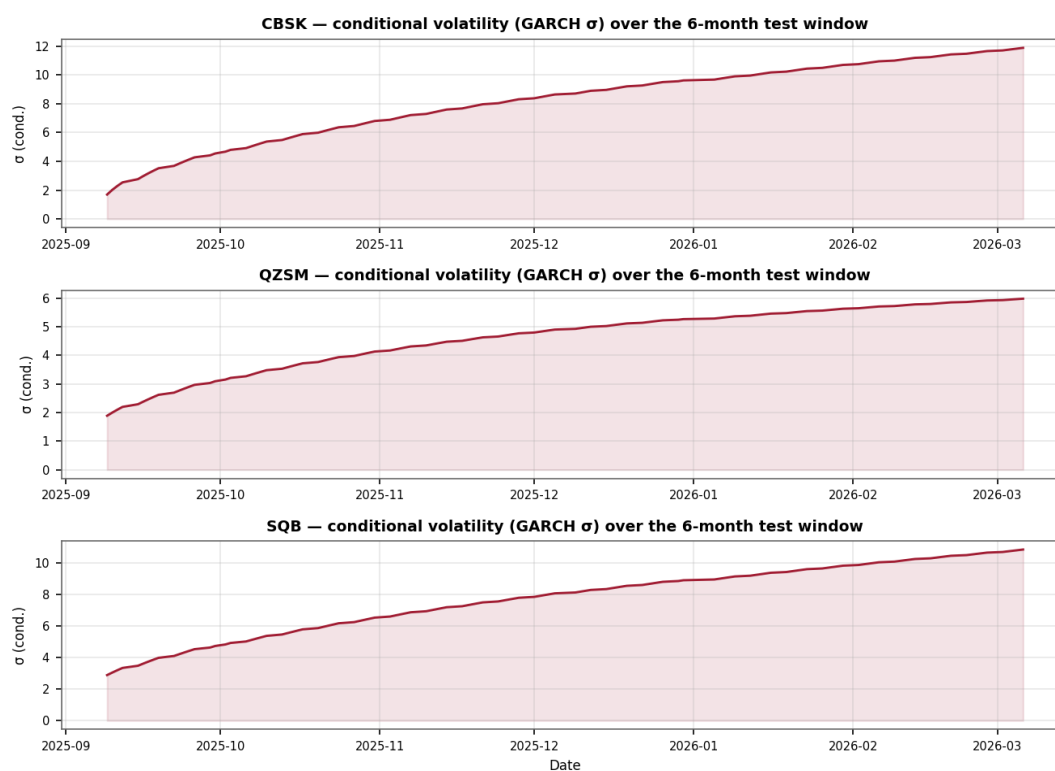


Figure 4. Estimated conditional volatility (GARCH σ) for the ARIMA–GARCH finalists

Conditional standard deviation over the six-month hold-out window for CBSK, QZSM and SQB. The series rise and cluster over time visible volatility clustering consistent with the high estimated persistence and the strong ARCH effects detected in Table 2. Source: authors' computations

6.3 Parameter stability and distributional fit

Against these strengths, two diagnostics flag genuine limitations that should not be glossed over. The Nyblom stability tests reject the constancy of the parameters for all three models: the joint statistic exceeds its 1% critical value in every case (8.24 for CBSK, 4.09 for QZSM and 2.98 for SQB). The instability is concentrated in the variance parameters for CBSK (the individual statistics for α_1 and β_1 greatly exceed their critical value), in the mean parameters for QZSM, and is joint-only for SQB (whose individual statistics all lie below the threshold). This points to possible structural change in the volatility process over the sampleplausible in a young market subject to evolving conditionsand suggests that rolling or regime-switching extensions merit investigation. In addition, the adjusted Pearson goodness-of-fit tests reject the assumed Student-t distribution at some groupings: mildly for CBSK and QZSM but more consistently for SQB (rejections at several groupings). The distributional fit is therefore imperfect, especially for SQB, reinforcing the earlier caution about that stock. In sum, the final ARIMA–GARCH models capture the mean and variance dynamics adequately and remove residual autocorrelation and ARCH, but they exhibit parameter instability and some distributional misspecification that temper confidence in their long-horizon forecasts.

7. Conclusion

This paper set out to evaluate and compare the forecasting performance of classical and volatility-adjusted time-series models for eight joint-stock companies listed in Uzbekistan, using daily prices from September 2022 to April 2026 within a three-year training and six-month testing design. Five modelsNaïve, Drift, ETS, ARIMA and ARIMA–GARCH with Student-t innovationswere estimated, compared on MAE, RMSE and MAPE, and the lowest-RMSE specification for each stock was re-estimated to forecast 2026Q3–2028Q4.

The key empirical findings are fourfold. First, all eight price series are integrated of order one, and six of the eight return series exhibit significant ARCH effects. Second, no single model is universally best: under the RMSE criterion, ETS is preferred for ALKB, HMKB and IPKY; ARIMA–GARCH for

CBSK, QZSM and SQB; the Drift model for URTS; and ARIMA for UZTL. Third, the volatility-adjusted model is selected only where ARCH effects are present and is never selected indeed performs worst for the two stocks (IPKY, URTS) without them, yet simple benchmarks remain highly competitive even for volatile series, and the margins by which ARIMA–GARCH wins are slim. Fourth, the final forecasts imply predominantly upward trajectories, with one flat path (UZTL, whose driftless ARIMA collapses to a constant forecast) and one decline (QZSM); several upward forecasts are aggressive linear extrapolations that warrant cautious interpretation.

The paper contributes rare, methodologically careful model-comparison evidence from a frontier equity market that is largely absent from the international forecasting literature, and it links model choice explicitly to the statistical properties of each series. Its practical message match the model to the security, retain simple benchmarks, and treat forecasts as conditional estimates directly useful to investors, analysts and market-development policy in Uzbekistan.

Several limitations should be acknowledged. The forecasting models are univariate and therefore omit macroeconomic and firm-level fundamentals; the sample, while sufficient for daily estimation, is relatively short by the standards of mature markets and is drawn from a thinly traded environment; the diagnostic analysis reveals parameter instability and imperfect distributional fit in the volatility models; and the long-horizon forecasts rest on the persistence of recent dynamics. The price unit and source exchange should also be confirmed in the final version.

These limitations point to a rich agenda for future research. Promising extensions include augmenting the models with macroeconomic variables (such as interest rates, inflation and the exchange rate); applying machine-learning and hybrid methods capable of capturing nonlinearity; comparing forecasts across daily, weekly and monthly frequencies; incorporating liquidity indicators; adopting regime-switching or time-varying-parameter models to accommodate the instability detected here; expanding the cross-section of listed companies; and benchmarking Uzbekistan against other emerging and frontier markets. Such work would sharpen both the accuracy and the economic interpretation of equity-price forecasts in developing capital markets.

References

1. Andersen, T. G., Bollerslev, T., Diebold, F. X., & Labys, P. (2003). Modeling and forecasting realized volatility. *Econometrica*, 71(2), 579–625.
2. Atsalakis, G. S., & Valavanis, K. P. (2009). Surveying stock market forecasting techniques – Part II: Soft computing methods. *Expert Systems with Applications*, 36(3), 5932–5941.
3. Bekaert, G., & Harvey, C. R. (1997). Emerging equity market volatility. *Journal of Financial Economics*, 43(1), 29–77.
4. Bollerslev, T. (1986). Generalized autoregressive conditional heteroskedasticity. *Journal of Econometrics*, 31(3), 307–327.
5. Bollerslev, T. (1987). A conditionally heteroskedastic time series model for speculative prices and rates of return. *Review of Economics and Statistics*, 69(3), 542–547.
6. Box, G. E. P., & Jenkins, G. M. (1970). *Time Series Analysis: Forecasting and Control*. San Francisco: Holden-Day.
7. Brooks, C. (2019). *Introductory Econometrics for Finance* (4th ed.). Cambridge: Cambridge University Press.
8. Cont, R. (2001). Empirical properties of asset returns: stylized facts and statistical issues. *Quantitative Finance*, 1(2), 223–236.
9. Dickey, D. A., & Fuller, W. A. (1979). Distribution of the estimators for autoregressive time series with a unit root. *Journal of the American Statistical Association*, 74(366), 427–431.
10. Diebold, F. X., & Mariano, R. S. (1995). Comparing predictive accuracy. *Journal of Business & Economic Statistics*, 13(3), 253–263.
11. Engle, R. F. (1982). Autoregressive conditional heteroscedasticity with estimates of the variance of United Kingdom inflation. *Econometrica*, 50(4), 987–1007.
12. Engle, R. F., & Bollerslev, T. (1986). Modelling the persistence of conditional variances. *Econometric Reviews*, 5(1), 1–50.
13. Fama, E. F. (1965). The behavior of stock-market prices. *Journal of Business*, 38(1), 34–105.
14. Fama, E. F. (1970). Efficient capital markets: A review of theory and empirical work. *Journal of Finance*, 25(2), 383–417.

15. Ghalanos, A. (2024). rugarch: Univariate GARCH models. R package (current version). Comprehensive R Archive Network.
16. Granger, C. W. J., & Newbold, P. (1986). *Forecasting Economic Time Series* (2nd ed.). San Diego: Academic Press.
17. Harvey, C. R. (1995). Predictable risk and returns in emerging markets. *Review of Financial Studies*, 8(3), 773–816.
18. Hyndman, R. J., & Athanasopoulos, G. (2021). *Forecasting: Principles and Practice* (3rd ed.). Melbourne: OTexts.
19. Hyndman, R. J., & Khandakar, Y. (2008). Automatic time series forecasting: The forecast package for R. *Journal of Statistical Software*, 27(3), 1–22.
20. Hyndman, R. J., & Koehler, A. B. (2006). Another look at measures of forecast accuracy. *International Journal of Forecasting*, 22(4), 679–688.
21. Hyndman, R. J., Koehler, A. B., Ord, J. K., & Snyder, R. D. (2008). *Forecasting with Exponential Smoothing: The State Space Approach*. Berlin: Springer.
22. Kearney, C. (2012). Emerging markets research: Trends, issues and future directions. *Emerging Markets Review*, 13(2), 159–183.
23. Ljung, G. M., & Box, G. E. P. (1978). On a measure of lack of fit in time series models. *Biometrika*, 65(2), 297–303.
24. Makridakis, S., & Hibon, M. (2000). The M3-Competition: results, conclusions and implications. *International Journal of Forecasting*, 16(4), 451–476.
25. Makridakis, S., Spiliotis, E., & Assimakopoulos, V. (2018). Statistical and machine learning forecasting methods: Concerns and ways forward. *PLoS ONE*, 13(3), e0194889.
26. Mandelbrot, B. (1963). The variation of certain speculative prices. *Journal of Business*, 36(4), 394–419.
27. Nelson, D. B. (1991). Conditional heteroskedasticity in asset returns: A new approach. *Econometrica*, 59(2), 347–370.
28. Pesaran, M. H., & Timmermann, A. (1995). Predictability of stock returns: Robustness and economic significance. *Journal of Finance*, 50(4), 1201–1228.

29. Poon, S.-H., & Granger, C. W. J. (2003). Forecasting volatility in financial markets: A review. *Journal of Economic Literature*, 41(2), 478–539.
30. R Core Team. (2024). *R: A Language and Environment for Statistical Computing*. Vienna: R Foundation for Statistical Computing.
31. Tsay, R. S. (2010). *Analysis of Financial Time Series* (3rd ed.). Hoboken, NJ: Wiley.
32. Welch, I., & Goyal, A. (2008). A comprehensive look at the empirical performance of equity premium prediction. *Review of Financial Studies*, 21(4), 1455–1508.
33. Zhang, G. P. (2003). Time series forecasting using a hybrid ARIMA and neural network model. *Neurocomputing*, 50, 159–175.
34. Sindarov, F.Q. (2024) Analysis of factors affecting the yield of bonds as an investment instrument. *World Bulletin of Management and Law (WBML)* Available Online, 35, pp. 19 – 28.