



Research article

A new XLindley ARMA model for positive time series

Amin Guerouah¹, Bassant Elkalzah^{2,3,*}, Halim Zeghdoudi⁴ and Majdah Mohammed Badr⁵

¹ Batna 2 University, 53 Constantine Road, Fesdis, Batna 05078, Algeria

² Department of Mathematics and Statistics, College of Science, Imam Mohammad Ibn Saud Islamic University (IMSIU), Riyadh, 11432, Saudi Arabia

³ Department of Statistics, Mathematics and Insurance, Faculty of Business, Alexandria University, Alexandria, 21526, Egypt

⁴ LaPS Laboratory, Department of Mathematics, Badji Mokhtar–Annaba University, P.O. Box 12, Annaba 23000, Algeria

⁵ Department of Mathematics and Statistics, College of Science, University of Jeddah, Jeddah, Saudi Arabia

* **Correspondence:** Email: bassant.elkalzah@alex.edu.eg.

Abstract: This paper proposes a new observation-driven model for continuous positive-valued time series, called the new XLindley autoregressive moving average (NXL–ARMA) model. The model is built by combining an ARMA-type recursion for the conditional mean with the new XLindley distribution as the conditional law of the observations. This framework is suitable for variables such as trading volume, realized volatility, durations, waiting times, precipitation amounts, energy consumption, and other positive measurements, where the data often show persistence, skewness, and variability that changes with the level of the series.

The proposed model preserves positivity by construction and remains relatively simple to handle from both theoretical and computational points of view. We derive the conditional moments in closed form and show that the conditional variance is proportional to the square of the conditional mean. This quadratic variance–mean relationship gives a natural motivation for using likelihood and quasi-likelihood estimation methods, including Gaussian, exponential, and geometric quasi-maximum likelihood approaches. We also establish conditions for strict stationarity and ergodicity, and prove the strong consistency of the proposed estimators under suitable regularity assumptions. Monte Carlo simulations are used to study the finite-sample behavior of the estimators, while empirical applications with benchmark comparisons illustrate the practical usefulness of the NXL–ARMA model for positive-valued financial time series.

Keywords: positive-valued time series; XLindley distribution; observation-driven models; ARMA dynamics; quasi-maximum likelihood

1. Introduction

Positive-valued time series occur naturally in many fields, including economics, finance, insurance, engineering, and environmental sciences. Typical examples include trading volumes, realized volatility, transaction durations, precipitation amounts, energy consumption, waiting times, claim sizes, and risk measures. These variables are restricted to the positive real line and often exhibit persistence, skewness, clustering, and variability that changes with the level of the series. Classical linear Gaussian time-series models are therefore not always appropriate, since they may produce negative fitted values and may fail to capture the asymmetric and heteroskedastic features commonly observed in positive data. This has motivated the development of dynamic models that preserve positivity and allow the conditional distribution to evolve over time.

A widely used class of models for positive-valued time series is the class of observation-driven models. In these models, the conditional mean is updated recursively as a function of past observations and past conditional means. Multiplicative error models (MEMs), introduced by [1], are among the most influential examples in financial econometrics. In MEMs, the observed positive variable is represented as the product of a time-varying conditional mean and a positive innovation term. This simple and interpretable structure has motivated several extensions, especially for volatility, duration, and trading-volume data [2–4]. Related models include the autoregressive conditional duration (ACD) model of [5] and the conditional autoregressive range (CARR) model of [6]. Together with the surveys of [7–9], these contributions show the importance of positivity-preserving dynamics in high-frequency financial applications [10, 11].

Another important direction is provided by generalized autoregressive moving average (GARMA) models for non-Gaussian time series, originally introduced by [12]. In this framework, the conditional mean is specified through an autoregressive moving average (ARMA)-type recursion, while the conditional distribution is chosen according to the nature of the data. GARMA-type models are flexible because they can be adapted to counts, proportions, durations, and other non-Gaussian variables. They are also closely related to quasi-likelihood and pseudo-likelihood methods, which are useful when the full conditional distribution is difficult to specify exactly or may be misspecified [13–15].

Recent developments in time-series analysis have also introduced more data-driven and computational approaches. For example, [16] proposed a nonlinear causal time-series network for financial risk assessment with imbalanced data, while [17] developed a multi-scale temporal correlation multi-dimensional decomposition network for time-series analysis. These methods illustrate the growing importance of flexible tools for extracting dependence, nonlinear structure, and risk-related information from complex time series. The objective of the present paper is different. We focus on an interpretable observation-driven stochastic model for continuous positive-valued time series, where the conditional distribution, stability condition, and likelihood-based inference can be studied explicitly.

It is also important to distinguish continuous positive-valued time series from integer-valued count time series. Count models are designed for variables taking values in $\mathbb{N}$ or $\mathbb{N}_0$ and often rely on thinning operators, threshold mechanisms, or discrete conditional distributions. Recent examples

include bivariate and multivariate threshold Poisson integer-valued autoregressive models, mixture integer-valued threshold autoregressive processes, and Tobit models for count time series [18–21]. In contrast, the variables considered in this paper are continuous and strictly positive. Scaled trading volumes, realized volatility, durations, and other positive measurements should not generally be treated as integer counts, because doing so may remove important distributional information. Continuous positive-valued models with suitable conditional distributions therefore remain necessary.

A substantial body of work has also studied inference for observation-driven models under possible distributional misspecification. Geometric and negative binomial quasi-maximum likelihood estimators were developed for integer-valued time series by [22, 23]. Other contributions have addressed stationarity, ergodicity, and identification in count and duration models [24, 25]. Periodic and regime-dependent extensions were considered by [26, 27, 30], while weighted least squares and multi-stage estimation procedures were studied by [28, 29]. These developments show that stability, identification, and robust estimation are central issues in modern observation-driven time-series modeling.

More recently, attention has turned to positive-valued time-series models with richer conditional distributions. In particular, [31] introduced the beta prime ARMA (BP-ARMA) model, which combines a linear conditional mean with a heavy-tailed distribution on the positive real line. This model is useful for strongly skewed and heavy-tailed data. However, its likelihood involves special functions, which may increase the computational burden in repeated estimation, simulation, forecasting, or higher-order specifications. This motivates the search for alternative positive distributions that remain flexible while being simpler to handle.

Lindley-type distributions provide one such alternative. They have received increasing attention because they combine simple density forms, tractable moments, and good empirical flexibility for positive data [32, 33]. The new XLindley distribution, introduced and studied by [34], belongs to this family and has shown good performance in fitting positive and right-skewed data. Its closed-form density and moments make it attractive for time-series applications, where likelihood contributions and conditional moments must be evaluated repeatedly.

Motivated by these considerations, this paper introduces a new XLindley autoregressive moving average model, denoted by NXL-ARMA, for continuous positive-valued time series. The proposed model combines an observation-driven ARMA-type conditional mean with the new XLindley conditional distribution. It preserves positivity by construction and remains relatively simple to estimate. A key feature of the model is its quadratic conditional variance–mean relationship, which means that the conditional variability increases with the conditional level of the series. This property is often observed in positive financial and economic data and provides a natural motivation for using suitable quasi-likelihood estimation methods, especially the geometric quasi-maximum likelihood estimator. In this sense, the proposed model offers a useful compromise between distributional flexibility, interpretability, and computational tractability.

The contribution of the paper is fourfold. First, we introduce the NXL-ARMA(p, q) model as a new observation-driven model for continuous positive-valued time series based on the new XLindley distribution. Second, we derive the conditional moments and discuss the resulting quadratic variance–mean relationship. Third, we establish strict stationarity and ergodicity under a transparent stability condition using monotonicity and contraction arguments. Fourth, we develop likelihood-based and quasi-likelihood-based estimation methods, including Gaussian, exponential, and GQMLEs, and prove

their strong consistency under suitable regularity conditions. The finite-sample behavior of the estimators is examined through Monte Carlo simulations with several sample sizes, and the empirical usefulness of the model is illustrated using financial time-series data with benchmark comparisons, diagnostic checks, and out-of-sample forecasting.

The remainder of the paper is organized as follows. Section 2 presents the New XLindley distribution and defines the NXL–ARMA model. Section 3 derives the conditional moments. Section 4 studies stationarity and ergodicity. Sections 5 and 6 discuss likelihood and quasi-likelihood estimation. Section 7 reports the simulation study. Section 8 presents empirical applications and model comparisons. Section 9 concludes the paper.

2. The NXL–ARMA model

This section presents the proposed NXL–ARMA model in two steps. We first recall the New XLindley distribution, which provides the conditional distributional component of the model. We then introduce an observation-driven ARMA-type recursion for the conditional mean. The main idea is to keep the process strictly positive while allowing the conditional level of the series to react to past observations and past conditional means.

2.1. The new XLindley distribution

The new XLindley distribution is a continuous distribution supported on the positive real line. It belongs to the family of Lindley-type distributions and has been used as a flexible model for positive and right-skewed data [34]. Its simple density and tractable moments make it attractive for time-series modeling, where likelihood contributions and conditional moments must be evaluated repeatedly.

A positive random variable Y is said to follow a New XLindley distribution with parameter $\theta > 0$, denoted by $Y \sim \text{NXL}(\theta)$, if its probability density function is given by

$$f(y; \theta) = \frac{\theta}{2}(1 + \theta y)e^{-\theta y}, \quad y > 0. \quad (2.1)$$

The following proposition recalls the first two moments of this distribution. These moments play an important role in the construction of the conditional mean parametrization used below.

Proposition 2.1. *Let $Y \sim \text{NXL}(\theta)$, with $\theta > 0$. Then*

$$E(Y) = \frac{3}{2\theta}, \quad \text{Var}(Y) = \frac{7}{4\theta^2}. \quad (2.2)$$

Proof. See [34]. □

These expressions show that the parameter θ controls both the location and the dispersion of the distribution. Smaller values of θ lead to larger means and larger variances. This is useful in a dynamic setting, because changes in the conditional mean are naturally accompanied by changes in the conditional variability.

2.2. Model specification

Let $\{Y_t\}_{t \in \mathbb{Z}}$ be a strictly positive stochastic process, and let $\mathcal{F}_t$ denote the σ -field generated by the past and present observations,

$$\mathcal{F}_t = \sigma(Y_s : s \leq t).$$

The NXL-ARMA(p, q) model is defined by assuming that, conditionally on the past information $\mathcal{F}_{t-1}$,

$$Y_t | \mathcal{F}_{t-1} \sim \text{NXL}(\theta_t). \quad (2.3)$$

The dynamics are introduced through the conditional mean

$$\mu_t = E(Y_t | \mathcal{F}_{t-1}),$$

which is assumed to follow the observation-driven recursion

$$\mu_t = \omega + \sum_{i=1}^q \alpha_i Y_{t-i} + \sum_{j=1}^p \beta_j \mu_{t-j}, \quad (2.4)$$

where

$$\omega > 0, \quad \alpha_i \geq 0, \quad \beta_j \geq 0.$$

The restrictions $\omega > 0$, $\alpha_i \geq 0$, and $\beta_j \geq 0$ ensure that the conditional mean remains positive whenever the past observations and past conditional means are positive. Thus, the model preserves positivity by construction. This is an essential property when modeling variables such as volatility, durations, trading volumes, waiting times, and other positive measurements.

Since the New XLindley mean is $3/(2\theta)$, the conditional distribution parameter is linked to the conditional mean by

$$\theta_t = \frac{3}{2\mu_t}. \quad (2.5)$$

Thus, the full model is specified by the parameter vector

$$\vartheta = (\omega, \alpha_1, \dots, \alpha_q, \beta_1, \dots, \beta_p)^\top.$$

The conditional distribution changes over time only through μ_t . This gives the model a clear interpretation: the ARMA-type recursion controls the conditional level of the process, while the new XLindley distribution describes the positive and skewed behavior around that level.

3. Conditional moments

The conditional moments are central to the interpretation of the NXL-ARMA model. They show how the conditional distribution reacts to changes in the conditional mean and also provide the basis for the quasi-likelihood methods considered later.

Proposition 3.1. *Under the NXL-ARMA(p, q) model,*

$$E(Y_t | \mathcal{F}_{t-1}) = \mu_t, \quad (3.1)$$

$$\text{Var}(Y_t | \mathcal{F}_{t-1}) = \frac{7}{9} \mu_t^2. \quad (3.2)$$

Proof. From the New XLindley moment formula,

$$E(Y_t | \mathcal{F}_{t-1}) = \frac{3}{2\theta_t}.$$

Using the link

$$\theta_t = \frac{3}{2\mu_t},$$

we obtain

$$E(Y_t | \mathcal{F}_{t-1}) = \frac{3}{2(3/(2\mu_t))} = \mu_t.$$

Similarly,

$$\text{Var}(Y_t | \mathcal{F}_{t-1}) = \frac{7}{4\theta_t^2}.$$

Substituting $\theta_t = 3/(2\mu_t)$ gives

$$\text{Var}(Y_t | \mathcal{F}_{t-1}) = \frac{7}{4} \left(\frac{2\mu_t}{3} \right)^2 = \frac{7}{9} \mu_t^2.$$

This proves the result. □

The proposition shows that the conditional variance is proportional to the square of the conditional mean. Therefore, when the conditional level of the process increases, the conditional variability increases as well. This is a realistic feature for many positive-valued financial, economic, insurance, and environmental time series.

Remark 3.1. *The quadratic variance–mean relationship is one of the main features of the NXL–ARMA model. It differs from models with constant or linear conditional variance and supports the use of quasi-likelihood methods that account for level-dependent variability. In particular, this property provides a natural motivation for the GQMLE considered later in the paper.*

4. Stationarity and ergodicity

This section studies the stability of the proposed NXL–ARMA(p, q) model. This step is important because a time-series model should have a well-defined long-run behavior before it is used for estimation, forecasting, or empirical interpretation. In the present setting, stationarity and ergodicity ensure that the effect of initial values disappears asymptotically and that likelihood-based arguments can be developed under standard long-run averaging principles.

To avoid confusion between the time-varying parameter of the new XLindley distribution and the vector of unknown model parameters used later in estimation, we denote the conditional distribution parameter by λ_t in this section.

4.1. Preliminaries and notation

Let $\mathcal{F}_t$ denote the σ -field generated by the past and present observations,

$$\mathcal{F}_t = \sigma(Y_s : s \leq t).$$

The NXL-ARMA(p, q) model is defined by

$$Y_t | \mathcal{F}_{t-1} \sim \text{NXL}(\lambda_t), \quad (4.1)$$

$$\mu_t = \omega + \sum_{i=1}^q \alpha_i Y_{t-i} + \sum_{j=1}^p \beta_j \mu_{t-j}, \quad \omega > 0, \quad \alpha_i \geq 0, \quad \beta_j \geq 0, \quad (4.2)$$

$$\lambda_t = \frac{3}{2\mu_t}. \quad (4.3)$$

The non-negativity restrictions on the coefficients are useful here because they preserve the positivity of the conditional mean and allow the monotonicity arguments below to be applied.

To write the process in Markovian form, define the state vector

$$X_t = (Y_t, Y_{t-1}, \dots, Y_{t-q+1}, \mu_t, \mu_{t-1}, \dots, \mu_{t-p+1})^\top \in (0, \infty)^{p+q}.$$

With this choice, the future evolution of the process depends on the past only through the current state vector X_t . Hence, $\{X_t\}$ is a time-homogeneous Markov chain on $(0, \infty)^{p+q}$.

A convenient construction of the model is obtained through the inverse distribution function. Let F_μ denote the distribution function of a new XLindley random variable with mean μ , that is, with distribution parameter $\lambda = 3/(2\mu)$. Let

$$Q_\mu(u) = F_\mu^{-1}(u), \quad 0 < u < 1,$$

be the corresponding quantile function. If $\{U_t\}$ is an independent sequence of Uniform(0, 1) random variables, then

$$Y_t = Q_{\mu_t}(U_t). \quad (4.4)$$

This representation is useful because it allows the process to be written as an iterated random function driven by the innovation sequence $\{U_t\}$.

4.2. Stochastic monotonicity of the new XLindley family

The next lemma shows that the new XLindley family is ordered with respect to the conditional mean. In simple terms, a larger conditional mean produces a stochastically larger conditional observation. This property is important because it allows two trajectories of the model to be compared when they are driven by the same sequence of innovations.

Lemma 4.1 (Stochastic monotonicity). *Let $\mu_1 < \mu_2$, and let F_{μ_1} and F_{μ_2} be the distribution functions of the new XLindley distribution under the mean parametrization $\lambda = 3/(2\mu)$. Then*

$$F_{\mu_1}(y) \geq F_{\mu_2}(y), \quad y \geq 0. \quad (4.5)$$

Equivalently, a new XLindley random variable with mean μ_2 is stochastically larger than one with mean μ_1 .

Proof. The new XLindley density is

$$f(y; \lambda) = \frac{\lambda}{2}(1 + \lambda y)e^{-\lambda y}, \quad y > 0, \quad \lambda > 0.$$

Under the mean parametrization, $\lambda(\mu) = 3/(2\mu)$. Hence, if $\mu_1 < \mu_2$, then

$$\lambda_1 = \lambda(\mu_1) > \lambda_2 = \lambda(\mu_2).$$

Consider the likelihood ratio

$$R(y) = \frac{f(y; \lambda_1)}{f(y; \lambda_2)} = \frac{\lambda_1}{\lambda_2} \frac{1 + \lambda_1 y}{1 + \lambda_2 y} \exp\{-(\lambda_1 - \lambda_2)y\}.$$

Taking the derivative of $\log R(y)$ gives

$$\frac{d}{dy} \log R(y) = \frac{\lambda_1}{1 + \lambda_1 y} - \frac{\lambda_2}{1 + \lambda_2 y} - (\lambda_1 - \lambda_2).$$

After simplification, we obtain

$$\frac{d}{dy} \log R(y) = (\lambda_1 - \lambda_2) \left[\frac{1}{(1 + \lambda_1 y)(1 + \lambda_2 y)} - 1 \right].$$

Since $\lambda_1 > \lambda_2$ and

$$(1 + \lambda_1 y)(1 + \lambda_2 y) > 1 \quad \text{for } y > 0,$$

the derivative is negative for $y > 0$. Therefore, $R(y)$ is decreasing in y . This monotone likelihood-ratio property implies the stochastic ordering in (4.5). $\square$

Remark 4.1. Lemma 4.1 implies that the quantile function is increasing in the mean parameter. More precisely, for every fixed $u \in (0, 1)$,

$$\mu_1 < \mu_2 \implies Q_{\mu_1}(u) \leq Q_{\mu_2}(u).$$

This is useful for coupling arguments, because two versions of the process can be driven by the same uniform innovation and compared pathwise.

4.3. Iterated random function representation

For each $u \in (0, 1)$, define the update map

$$\Phi_u : (0, \infty)^{p+q} \rightarrow (0, \infty)^{p+q}.$$

For a state vector

$$x = (y_0, y_1, \dots, y_{q-1}, m_0, m_1, \dots, m_{p-1}),$$

where y_i denotes a lagged observation and m_j denotes a lagged conditional mean, define the next conditional mean by

$$m^+ = \omega + \sum_{i=1}^q \alpha_i y_{i-1} + \sum_{j=1}^p \beta_j m_{j-1}. \quad (4.6)$$

The next observation is then generated by

$$y^+ = Q_{m^+}(u). \quad (4.7)$$

The updated state is obtained by inserting the new values and shifting the lags:

$$\Phi_u(x) = (y^+, y_0, \dots, y_{q-2}, m^+, m_0, \dots, m_{p-2}),$$

with the usual convention that empty lag lists are omitted when $p = 1$ or $q = 1$. Consequently, the state process satisfies

$$X_t = \Phi_{U_t}(X_{t-1}), \quad \{U_t\} \text{ i.i.d. Uniform}(0, 1). \quad (4.8)$$

This formulation makes the observation-driven nature of the model explicit. Once the previous state and the new innovation U_t are known, both μ_t and Y_t are determined.

4.4. Stability condition

Let

$$\gamma = \sum_{i=1}^q \alpha_i + \sum_{j=1}^p \beta_j.$$

The number γ measures the total feedback from past observations and past conditional means. If this feedback is too strong, shocks may persist for a long time, and the conditional mean may fail to stabilize. The stability condition used in the sequel is therefore

$$\gamma < 1.$$

The following lemma records the basic contraction property of the conditional mean recursion.

Lemma 4.2 (Mean contraction). *For any two states*

$$x = (y_0, \dots, y_{q-1}, m_0, \dots, m_{p-1})$$

and

$$x' = (y'_0, \dots, y'_{q-1}, m'_0, \dots, m'_{p-1}),$$

the corresponding mean updates m^+ and $(m^+)'$ satisfy

$$|m^+ - (m^+)'| \leq \sum_{i=1}^q \alpha_i |y_{i-1} - y'_{i-1}| + \sum_{j=1}^p \beta_j |m_{j-1} - m'_{j-1}|. \quad (4.9)$$

Proof. From the linear form of the conditional mean recursion,

$$m^+ - (m^+)' = \sum_{i=1}^q \alpha_i (y_{i-1} - y'_{i-1}) + \sum_{j=1}^p \beta_j (m_{j-1} - m'_{j-1}).$$

Taking absolute values and applying the triangle inequality gives (4.9). $\square$

The next lemma shows that the whole update map preserves the natural order on the state space.

Lemma 4.3 (Order-preserving update). *For a fixed $u \in (0, 1)$, the update map Φ_u is order-preserving. In particular, if $x \leq x'$ componentwise, then*

$$\Phi_u(x) \leq \Phi_u(x')$$

componentwise.

Proof. If $x \leq x'$ componentwise, then the non-negativity of the coefficients α_i and β_j implies

$$m^+ \leq (m^+)'.$$

By Lemma 4.1, the quantile function is increasing in the mean parameter. Therefore,

$$Q_{m^+}(u) \leq Q_{(m^+)'}(u).$$

The remaining components of $\Phi_u(x)$ are obtained by shifting lagged values. Hence, the full update map preserves the componentwise order. $\square$

4.5. Main stationarity result

The stationarity and ergodicity of the proposed NXL–ARMA model are established using the iterated random function framework introduced in Section 4.3. Specifically, we apply Theorem 2 of Douc, Doukhan and Moulines [37], which establishes sufficient conditions for the existence of a unique invariant probability measure and, consequently, a unique strictly stationary and ergodic solution for observation-driven Markov chains. To ensure that the argument is self-contained, we explicitly verify each hypothesis of the cited theorem for the proposed NXL–ARMA model before invoking its conclusion.

Theorem 4.1 (Strict stationarity and ergodicity). *Let*

$$\gamma = \sum_{i=1}^q \alpha_i + \sum_{j=1}^p \beta_j.$$

If $\gamma < 1$, then the NXL–ARMA(p, q) model defined by (4.1)–(4.3) admits a unique strictly stationary and ergodic solution with finite first moment. Conversely, if the model admits a strictly stationary solution with finite mean, then necessarily $\gamma < 1$.

Proof. We prove separately the sufficiency and necessity of the stability condition.

Sufficiency.

Assume that

$$\gamma = \sum_{i=1}^q \alpha_i + \sum_{j=1}^p \beta_j < 1.$$

Consider the state process

$$X_t = \Phi_{U_t}(X_{t-1}),$$

where $\{U_t\}$ is an i.i.d. sequence of Uniform(0, 1) random variables.

To apply the stability theorem of Douc, Doukhan, and Moulines [37], we verify that its assumptions hold for the proposed model.

First, for every $u \in (0, 1)$, the update map

$$\Phi_u : (0, \infty)^{p+q} \rightarrow (0, \infty)^{p+q}$$

is continuous with respect to the state vector, since the conditional mean update is affine and the new XLindley quantile function is continuous. Hence Φ_u is Borel measurable.

Second, by Lemma 4.3, Φ_u preserves the natural partial order, while Lemma 4.1 guarantees the required stochastic monotonicity of the conditional distribution.

Third, Lemma 4.2 gives

$$|m^+ - (m^+)'| \leq \gamma \|x - x'\|_1,$$

so that the conditional mean update is globally Lipschitz with contraction constant $\gamma < 1$.

Finally, taking expectations in the conditional mean recursion yields

$$(1 - \gamma)E(\mu_t) = \omega,$$

which implies

$$E(\mu_t) = \frac{\omega}{1 - \gamma} < \infty.$$

Thus the required drift condition is satisfied.

Therefore, every assumption of the stability theorem is verified. Theorem 2 of Douc, Doukhan and Moulines [37] implies that the Markov chain $\{X_t\}$ possesses a unique invariant probability measure. Consequently, $\{X_t\}$ admits a unique strictly stationary and ergodic solution. Since $\{Y_t\}$ is a measurable coordinate of the state vector, the observation process $\{Y_t\}$ is also strictly stationary and ergodic.

Moreover,

$$E(Y_t) = E[E(Y_t | \mathcal{F}_{t-1})] = E(\mu_t) = \frac{\omega}{1 - \gamma} < \infty,$$

so the stationary solution has finite first moment.

Necessity.

Conversely, suppose that the model admits a strictly stationary solution with finite mean. Taking unconditional expectations in (4.2) gives

$$E(\mu_t) = \omega + \gamma E(\mu_t),$$

or equivalently,

$$(1 - \gamma)E(\mu_t) = \omega.$$

Since

$$\omega > 0 \quad \text{and} \quad 0 < E(\mu_t) < \infty,$$

we obtain

$$1 - \gamma = \frac{\omega}{E(\mu_t)} > 0,$$

which immediately implies

$$\gamma < 1.$$

Hence the condition

$$\sum_{i=1}^q \alpha_i + \sum_{j=1}^p \beta_j < 1$$

is both necessary and sufficient for the existence of a unique strictly stationary and ergodic solution with finite first moment. $\square$

5. Likelihood-based estimation

This section presents likelihood-based estimation for the NXL-ARMA(p, q) model. Since the conditional distribution is specified as New XLindley, the conditional likelihood can be written explicitly once the conditional mean sequence is known. This makes the estimation procedure fairly direct: for each candidate parameter value, we compute the conditional mean recursively and then substitute it into the New XLindley density.

To avoid confusion between the parameter of the New XLindley distribution and the vector of unknown model parameters, we denote the model parameter vector by $\boldsymbol{\vartheta}$ and the time-varying New XLindley parameter by $\lambda_t(\boldsymbol{\vartheta})$.

The purpose of this section is to define the likelihood estimator and to establish its strong consistency. The argument follows the usual structure for observation-driven models. First, the parameter space is restricted to the stationary region. Second, the sample likelihood criterion is shown to converge uniformly to its population counterpart. Third, the limiting criterion is shown to have a unique maximum at the true parameter.

5.1. Parameter space and conditional mean recursion

Let

$$\boldsymbol{\vartheta} = (\omega, \alpha_1, \dots, \alpha_q, \beta_1, \dots, \beta_p)^\top \in \Theta$$

be the vector of unknown model parameters. The parameter space Θ is assumed to be a compact subset of

$$(0, \infty) \times [0, \infty)^q \times [0, \infty)^p.$$

In addition, the parameters are restricted to the stationary region

$$\sum_{i=1}^q \alpha_i + \sum_{j=1}^p \beta_j \leq \rho < 1, \quad (5.1)$$

for some fixed constant ρ .

This restriction is not only technical. It ensures that every parameter value considered during estimation generates a stable conditional mean recursion. Hence, likelihood comparisons are made only between models that define well-behaved stationary dynamics.

For any $\boldsymbol{\vartheta} \in \Theta$, the conditional mean is computed recursively as

$$\mu_t(\boldsymbol{\vartheta}) = \omega + \sum_{i=1}^q \alpha_i Y_{t-i} + \sum_{j=1}^p \beta_j \mu_{t-j}(\boldsymbol{\vartheta}), \quad t \in \mathbb{Z}. \quad (5.2)$$

Under (5.1), this recursion has a unique stationary and ergodic solution for each $\boldsymbol{\vartheta} \in \Theta$.

The corresponding New XLindley conditional parameter is obtained from the mean-parametrization formula

$$\lambda_t(\boldsymbol{\vartheta}) = \frac{3}{2\mu_t(\boldsymbol{\vartheta})}. \quad (5.3)$$

In practice, the recursion for $\mu_t(\boldsymbol{\vartheta})$ must be initialized. Since the recursion is stable under (5.1), the effect of the initial values vanishes asymptotically. One may therefore initialize the recursion using the sample mean, the unconditional mean

$$\frac{\omega}{1 - \sum_i \alpha_i - \sum_j \beta_j},$$

or any fixed positive value. This point is important in applications because it shows that the estimator is not driven by arbitrary starting choices when the sample size is large.

5.2. Conditional log-likelihood

Given observations $Y_1, \dots, Y_n$, the conditional New XLindley log-likelihood contribution at time t is

$$\ell_t(\boldsymbol{\vartheta}) = \log \lambda_t(\boldsymbol{\vartheta}) - \log 2 + \log(1 + \lambda_t(\boldsymbol{\vartheta})Y_t) - \lambda_t(\boldsymbol{\vartheta})Y_t. \quad (5.4)$$

The average log-likelihood is defined by

$$L_n(\boldsymbol{\vartheta}) = \frac{1}{n} \sum_{t=1}^n \ell_t(\boldsymbol{\vartheta}). \quad (5.5)$$

The likelihood estimator is then given by

$$\widehat{\boldsymbol{\vartheta}}_n = \arg \max_{\boldsymbol{\vartheta} \in \Theta} L_n(\boldsymbol{\vartheta}). \quad (5.6)$$

This estimator uses the full conditional new XLindley density. It therefore serves as the natural reference estimator for the proposed model. The quasi-likelihood estimators discussed later are useful alternatives when the main interest is in the conditional mean and variance structure, or when one wants an estimator that is less dependent on the exact conditional density.

5.3. Regularity assumptions

The following assumptions are used to establish strong consistency.

Assumption 5.1. *The following conditions hold:*

- (A1) *The true parameter $\boldsymbol{\vartheta}_0$ belongs to the interior of Θ .*
- (A2) *The process $\{Y_t\}$ is strictly stationary and ergodic.*

(A3) $E(Y_t) < \infty$.

(A4) The model is identifiable. More precisely, for every $\boldsymbol{\vartheta} \neq \boldsymbol{\vartheta}_0$,

$$P(\mu_t(\boldsymbol{\vartheta}) = \mu_t(\boldsymbol{\vartheta}_0)) < 1.$$

Assumption (A1) avoids boundary complications. Assumption (A2) follows from the stationarity result when the true parameter satisfies the stability condition. Assumption (A3) ensures that the likelihood contributions are integrable. Assumption (A4) is the identifiability condition: it means that two different parameter values cannot generate the same conditional mean sequence almost surely. This condition is especially relevant for higher-order models, where unnecessary lag terms may lead to weak identification.

5.4. Uniform convergence of the likelihood criterion

The consistency of the maximum likelihood estimator (MLE) is established within the classical M-estimation framework of Newey and McFadden [38]. To ensure that the theoretical development is self-contained, we explicitly verify below the regularity conditions required by the general consistency theorem. These include the compactness of the parameter space, continuity of the conditional log-likelihood, the existence of an integrable envelope function, the uniform convergence of the sample criterion to its population counterpart, the identifiability of the limiting criterion, and the uniqueness of its maximizer. Once these properties have been established, the consistency theorem can be applied directly to the proposed NXL-ARMA model.

Lemma 5.1 (Uniform law of large numbers). *Under Assumption 5.1,*

$$\sup_{\boldsymbol{\vartheta} \in \Theta} |L_n(\boldsymbol{\vartheta}) - L(\boldsymbol{\vartheta})| \xrightarrow{\text{a.s.}} 0,$$

where

$$L(\boldsymbol{\vartheta}) = E\{\ell_t(\boldsymbol{\vartheta})\}.$$

Proof. The proof is based on the consistency theory for M-estimators developed by Newey and McFadden [38]. We verify below that the sample criterion satisfies the hypotheses required for the uniform law of large numbers.

(i) Compactness.

Assumption 5.1(A1) states that the parameter space Θ is compact and satisfies

$$\sum_{i=1}^q \alpha_i + \sum_{j=1}^p \beta_j \leq \rho < 1.$$

Hence every admissible parameter vector generates a well-defined and stable conditional mean recursion.

(ii) Continuity.

For every fixed t , the recursion $\mu_t(\boldsymbol{\vartheta})$ depends continuously on $\boldsymbol{\vartheta}$. Since

$$\lambda_t(\boldsymbol{\vartheta}) = \frac{3}{2\mu_t(\boldsymbol{\vartheta})},$$

the mapping

$$\boldsymbol{\vartheta} \mapsto \ell_t(\boldsymbol{\vartheta})$$

is almost surely continuous on Θ .

Moreover, because ω is bounded away from zero and Θ is compact, there exists $c > 0$ such that

$$\mu_t(\boldsymbol{\vartheta}) \geq c, \quad \forall \boldsymbol{\vartheta} \in \Theta.$$

Hence every logarithmic term appearing in (5.4) is well defined and continuous uniformly over Θ .

(iii) Integrable envelope.

Using the explicit form of the conditional log-likelihood, there exist constants $C_1, C_2 > 0$, independent of $\boldsymbol{\vartheta}$, such that

$$|\ell_t(\boldsymbol{\vartheta})| \leq C_1 + C_2 Y_t, \quad \forall \boldsymbol{\vartheta} \in \Theta.$$

Since Assumption 5.1(A3) implies

$$E(Y_t) < \infty,$$

the random variable

$$G(Y_t) = C_1 + C_2 Y_t$$

is integrable. Consequently,

$$|\ell_t(\boldsymbol{\vartheta})| \leq G(Y_t), \quad \forall \boldsymbol{\vartheta} \in \Theta,$$

so that the family $\{\ell_t(\boldsymbol{\vartheta}) : \boldsymbol{\vartheta} \in \Theta\}$ admits an integrable envelope.

(iv) Stationarity and ergodicity.

By Theorem 4.1, the process $\{Y_t\}$ is strictly stationary and ergodic. Therefore, for every $\boldsymbol{\vartheta} \in \Theta$, the sequence $\{\ell_t(\boldsymbol{\vartheta})\}$ is also strictly stationary and ergodic.

Hence all hypotheses required for the uniform law of large numbers in the M-estimation framework are satisfied. It follows that

$$\sup_{\boldsymbol{\vartheta} \in \Theta} |L_n(\boldsymbol{\vartheta}) - L(\boldsymbol{\vartheta})| \xrightarrow{\text{a.s.}} 0.$$

This completes the proof. □

5.5. Identification of the limiting criterion

Lemma 5.2. *The population criterion*

$$\boldsymbol{\vartheta} \mapsto \mathbb{E}\{\ell_t(\boldsymbol{\vartheta})\}$$

has a unique maximum at $\boldsymbol{\vartheta} = \boldsymbol{\vartheta}_0$.

Proof. Let

$$L(\boldsymbol{\vartheta}) = \mathbb{E}\{\ell_t(\boldsymbol{\vartheta})\}$$

denote the population log-likelihood criterion. Furthermore, let $f_{\boldsymbol{\vartheta}}(\cdot | \mathcal{F}_{t-1})$ be the conditional density of Y_t under parameter $\boldsymbol{\vartheta}$. Then

$$\ell_t(\boldsymbol{\vartheta}) = \log f_{\boldsymbol{\vartheta}}(Y_t | \mathcal{F}_{t-1}).$$

Let $\boldsymbol{\vartheta} \neq \boldsymbol{\vartheta}_0$. By Assumption (A4), the conditional means $\mu_t(\boldsymbol{\vartheta})$ and $\mu_t(\boldsymbol{\vartheta}_0)$ differ with positive probability. Since the new XLindley density is indexed by the conditional mean through

$$\lambda_t(\boldsymbol{\vartheta}) = \frac{3}{2\mu_t(\boldsymbol{\vartheta})},$$

different conditional means imply different conditional densities with positive probability. Consequently,

$$\begin{aligned} L(\boldsymbol{\vartheta}) - L(\boldsymbol{\vartheta}_0) &= \mathbb{E} \left[\log \frac{f_{\boldsymbol{\vartheta}}(Y_t | \mathcal{F}_{t-1})}{f_{\boldsymbol{\vartheta}_0}(Y_t | \mathcal{F}_{t-1})} \right] \\ &= -\mathbb{E} [D_{\text{KL}}(f_{\boldsymbol{\vartheta}_0}(\cdot | \mathcal{F}_{t-1}) \| f_{\boldsymbol{\vartheta}}(\cdot | \mathcal{F}_{t-1}))]. \end{aligned}$$

Since the Kullback–Leibler divergence is non-negative,

$$D_{\text{KL}}(f_{\boldsymbol{\vartheta}_0} \| f_{\boldsymbol{\vartheta}}) \geq 0,$$

it follows immediately that

$$L(\boldsymbol{\vartheta}) \leq L(\boldsymbol{\vartheta}_0).$$

Moreover,

$$D_{\text{KL}}(f_{\boldsymbol{\vartheta}_0} \| f_{\boldsymbol{\vartheta}}) = 0$$

if and only if the two conditional densities coincide almost surely. Since the conditional density is parameterized through $\mu_t(\boldsymbol{\vartheta})$, this is equivalent to

$$\mu_t(\boldsymbol{\vartheta}) = \mu_t(\boldsymbol{\vartheta}_0) \quad \text{a.s.}$$

Finally, Assumption (A4) implies that

$$\mu_t(\boldsymbol{\vartheta}) = \mu_t(\boldsymbol{\vartheta}_0) \implies \boldsymbol{\vartheta} = \boldsymbol{\vartheta}_0.$$

Therefore,

$$L(\boldsymbol{\vartheta}) = L(\boldsymbol{\vartheta}_0) \iff \boldsymbol{\vartheta} = \boldsymbol{\vartheta}_0,$$

which proves that the population criterion possesses the unique maximizer $\boldsymbol{\vartheta}_0$. □

5.6. Strong consistency

Theorem 5.1 (Strong consistency). *Under Assumption 5.1,*

$$\widehat{\boldsymbol{\vartheta}}_n \xrightarrow{a.s.} \boldsymbol{\vartheta}_0, \quad n \rightarrow \infty.$$

Proof. By Lemma 5.1, the sample criterion $L_n(\boldsymbol{\vartheta})$ converges uniformly almost surely to the population criterion $\mathbb{E}\{\ell_t(\boldsymbol{\vartheta})\}$ on the compact set Θ . By Lemma 5.2, the population criterion has a unique maximizer at $\boldsymbol{\vartheta}_0$. Therefore, the argmax theorem yields

$$\widehat{\boldsymbol{\vartheta}}_n \xrightarrow{a.s.} \boldsymbol{\vartheta}_0.$$

□

Remark 5.1 (Likelihood estimation). *The MLE exploits the complete conditional new XLindley distribution and is therefore expected to achieve higher statistical efficiency when the model is correctly specified. If the conditional distribution is misspecified while the conditional mean dynamics remain adequate, quasi-likelihood approaches based on working distributions such as the Gaussian or exponential may provide robust alternatives for estimation and inference.*

Remark 5.2 (Asymptotic inference). *Theorem 5.1 establishes the strong consistency of the MLE. The derivation of its asymptotic distribution,*

$$\sqrt{n}(\widehat{\boldsymbol{\vartheta}}_n - \boldsymbol{\vartheta}_0),$$

requires additional regularity conditions, including twice continuous differentiability of the log-likelihood, local identifiability, nonsingularity of the information matrix, and suitable moment conditions. Under these assumptions, asymptotic normality may be established using standard martingale arguments, yielding the usual sandwich covariance matrix. Since these conditions go beyond the scope of the present work, asymptotic inference is left for future research.

Remark 5.3 (Computation). *For fixed orders p and q , evaluation of the likelihood requires only recursive updates of the conditional mean and therefore has linear computational complexity with respect to the sample size. Although higher-order models remain computationally feasible, they introduce additional parameters that may increase numerical instability and weaken parameter identifiability. Accordingly, model order should be selected using likelihood-based criteria, information criteria, residual diagnostics, and parsimony considerations.*

Remark 5.4 (Theoretical justification). *The revised manuscript explicitly verifies the assumptions required by the iterated random function stability theorem used to establish stationarity and ergodicity, as well as the regularity conditions underlying the M-estimation framework employed to prove consistency. Consequently, each theoretical result is invoked only after its hypotheses have been verified for the proposed NXL-ARMA model, rendering the mathematical development self-contained.*

6. Geometric quasi-maximum likelihood estimation

This section introduces the GQMLE for the NXL–ARMA(p, q) model. The motivation comes from the conditional variance of the proposed model. As shown in the previous section,

$$\text{Var}(Y_t | \mathcal{F}_{t-1}) = \frac{7}{9}\mu_t^2.$$

Thus, the conditional variance is proportional to the square of the conditional mean. This is a natural feature for many positive-valued financial and economic time series, where periods with larger conditional levels are often accompanied by larger fluctuations.

Although the New XLindley distribution is continuous, its variance–mean behavior is close in spirit to that of the geometric distribution, whose variance is $\mu(1 + \mu)$. For moderate and large values of μ , this variance is mainly driven by its quadratic term. This connection motivates the use of a geometric working likelihood as a quasi-likelihood criterion. The aim is not to treat positive continuous observations as counts, but to use a simple and convenient estimating criterion whose weighting structure reflects the heteroskedastic behavior of the NXL–ARMA model.

6.1. Geometric quasi-likelihood

Let $\mu_t(\boldsymbol{\vartheta})$ denote the conditional mean generated by the NXL–ARMA recursion for a parameter vector $\boldsymbol{\vartheta} \in \Theta$. As a working likelihood, we consider a geometric distribution with mean $\mu > 0$:

$$P(Y = y) = \frac{1}{1 + \mu} \left(\frac{\mu}{1 + \mu} \right)^y, \quad y \in \mathbb{N}_0.$$

Even though Y_t is continuous in the present model, the logarithmic form of this likelihood can still be used as a quasi-likelihood criterion. This is a standard idea in quasi-likelihood estimation: the working likelihood is used only to construct an objective function, while consistency is mainly driven by the correct specification of the conditional mean.

The geometric quasi-log-likelihood contribution is defined by

$$\ell_t^{(G)}(\boldsymbol{\vartheta}) = Y_t \log \left(\frac{\mu_t(\boldsymbol{\vartheta})}{1 + \mu_t(\boldsymbol{\vartheta})} \right) - \log(1 + \mu_t(\boldsymbol{\vartheta})). \quad (6.1)$$

The corresponding average criterion is

$$L_n^{(G)}(\boldsymbol{\vartheta}) = \frac{1}{n} \sum_{t=1}^n \ell_t^{(G)}(\boldsymbol{\vartheta}), \quad (6.2)$$

and the GQMLE is

$$\widehat{\boldsymbol{\vartheta}}_n^{(G)} = \arg \max_{\boldsymbol{\vartheta} \in \Theta} L_n^{(G)}(\boldsymbol{\vartheta}). \quad (6.3)$$

6.2. Score and moment interpretation

The geometric quasi-likelihood has a simple and useful estimating-equation interpretation. Differentiating (6.1) with respect to the parameter vector gives

$$\frac{\partial \ell_t^{(G)}(\boldsymbol{\vartheta})}{\partial \boldsymbol{\vartheta}} = \frac{Y_t - \mu_t(\boldsymbol{\vartheta})}{\mu_t(\boldsymbol{\vartheta})\{1 + \mu_t(\boldsymbol{\vartheta})\}} \frac{\partial \mu_t(\boldsymbol{\vartheta})}{\partial \boldsymbol{\vartheta}}. \quad (6.4)$$

Therefore, the GQMLE satisfies the estimating equation

$$\sum_{i=1}^n \frac{Y_i - \mu_i(\boldsymbol{\vartheta})}{\mu_i(\boldsymbol{\vartheta})\{1 + \mu_i(\boldsymbol{\vartheta})\}} \frac{\partial \mu_i(\boldsymbol{\vartheta})}{\partial \boldsymbol{\vartheta}} = 0. \quad (6.5)$$

This expression shows that the estimator is driven by the prediction error $Y_i - \mu_i(\boldsymbol{\vartheta})$. The denominator acts as a stabilizing weight and reduces the influence of observations associated with large conditional means. This weighting is suitable for positive-valued series with multiplicative-type variability, because the conditional dispersion increases with the level of the process.

6.3. Regularity assumptions

The consistency result is obtained under the following assumptions.

Assumption 6.1. *The following conditions hold:*

- (G1) *The true parameter $\boldsymbol{\vartheta}_0$ belongs to the interior of the compact parameter space Θ .*
- (G2) *The NXL-ARMA(p, q) process is strictly stationary and ergodic.*
- (G3) *$E(Y_i^2) < \infty$.*
- (G4) *The model is identifiable in the sense that for every $\boldsymbol{\vartheta} \neq \boldsymbol{\vartheta}_0$,*

$$P(\mu_i(\boldsymbol{\vartheta}) = \mu_i(\boldsymbol{\vartheta}_0)) < 1.$$

Assumption (G3) is slightly stronger than the first-moment condition used for likelihood consistency. It is natural here because the geometric quasi-likelihood is motivated by a variance–mean relationship and because moment-based arguments are used in the score interpretation. Assumption (G4) prevents two different parameter values from producing the same conditional mean process almost surely.

6.4. Uniform convergence

Lemma 6.1 (Uniform convergence of the geometric criterion). *Under Assumption 6.1,*

$$\sup_{\boldsymbol{\vartheta} \in \Theta} |L_n^{(G)}(\boldsymbol{\vartheta}) - \mathbb{E}\{\ell_t^{(G)}(\boldsymbol{\vartheta})\}| \xrightarrow{a.s.} 0.$$

Proof. For each t , the function $\mu_t(\boldsymbol{\vartheta})$ is continuous in $\boldsymbol{\vartheta}$. Hence, the geometric quasi-log-likelihood contribution $\ell_t^{(G)}(\boldsymbol{\vartheta})$ is also continuous in $\boldsymbol{\vartheta}$ almost surely.

Since Θ is compact and ω is bounded away from zero on Θ , the conditional mean $\mu_t(\boldsymbol{\vartheta})$ is uniformly bounded away from zero. Moreover, under the stationarity restriction on Θ , the recursive conditional means are dominated by an integrable process depending on past observations. Therefore, there exists an integrable random variable M_t such that

$$|\ell_t^{(G)}(\boldsymbol{\vartheta})| \leq M_t, \quad \boldsymbol{\vartheta} \in \Theta.$$

Because $E(Y_i^2) < \infty$ implies $E(Y_i) < \infty$, this domination is sufficient for the uniform ergodic theorem. Since the process is stationary and ergodic by Assumption (G2), we obtain

$$\sup_{\boldsymbol{\vartheta} \in \Theta} |L_n^{(G)}(\boldsymbol{\vartheta}) - \mathbb{E}\{\ell_t^{(G)}(\boldsymbol{\vartheta})\}| \xrightarrow{a.s.} 0.$$

□

6.5. Identification of the limiting criterion

Lemma 6.2. *The population criterion*

$$\boldsymbol{\vartheta} \mapsto \mathbb{E}\{\ell_t^{(G)}(\boldsymbol{\vartheta})\}$$

is uniquely maximized at $\boldsymbol{\vartheta} = \boldsymbol{\vartheta}_0$.

Proof. For a fixed value of the past information $\mathcal{F}_{t-1}$, define

$$m_0 = \mu_t(\boldsymbol{\vartheta}_0), \quad m = \mu_t(\boldsymbol{\vartheta}).$$

Using

$$E(Y_t | \mathcal{F}_{t-1}) = m_0,$$

the conditional expected geometric quasi-log-likelihood can be written, up to terms not depending on Y_t , as

$$m_0 \log\left(\frac{m}{1+m}\right) - \log(1+m).$$

As a function of $m > 0$, this expression is maximized at $m = m_0$. Indeed, its derivative with respect to m is

$$\frac{m_0 - m}{m(1+m)},$$

which is positive when $m < m_0$ and negative when $m > m_0$. Thus, the conditional expected criterion is uniquely maximized when

$$\mu_t(\boldsymbol{\vartheta}) = \mu_t(\boldsymbol{\vartheta}_0) \quad \text{almost surely.}$$

By Assumption (G4), this equality holds only for $\boldsymbol{\vartheta} = \boldsymbol{\vartheta}_0$. Therefore, the unconditional expected criterion is uniquely maximized at the true parameter. $\square$

6.6. Strong consistency of the GQMLE

Theorem 6.1 (Strong consistency of the GQMLE). *Under Assumption 6.1,*

$$\widehat{\boldsymbol{\vartheta}}_n^{(G)} \xrightarrow{a.s.} \boldsymbol{\vartheta}_0, \quad n \rightarrow \infty.$$

Proof. By Lemma 6.1, the sample criterion $L_n^{(G)}(\boldsymbol{\vartheta})$ converges uniformly almost surely to the population criterion $\mathbb{E}\{\ell_t^{(G)}(\boldsymbol{\vartheta})\}$ on Θ . By Lemma 6.2, the limiting criterion has a unique maximizer at $\boldsymbol{\vartheta}_0$. The argmax theorem then yields

$$\widehat{\boldsymbol{\vartheta}}_n^{(G)} \xrightarrow{a.s.} \boldsymbol{\vartheta}_0.$$

$\square$

Remark 6.1 (Why the geometric quasi-likelihood is useful). *The GQMLE does not assume that the positive-valued observations are geometric counts. The geometric likelihood is used only as a convenient working criterion. Its score has the form of a weighted conditional mean estimating equation, where the weight reflects a quadratic-type variance–mean relationship. This makes it especially suitable for the NXL–ARMA model.*

Remark 6.2 (Efficiency). *When the full new XLindley conditional distribution is correctly specified, the likelihood estimator remains the natural efficiency benchmark. The GQMLE may be less efficient than the full likelihood estimator, but it can be attractive in practice because it is simple, robust to some forms of distributional misspecification, and aligned with the variance–mean behavior of the model. For this reason, its usefulness should be assessed through simulation and empirical comparison rather than asserted only on theoretical grounds.*

Remark 6.3 (Asymptotic normality). *The present result establishes strong consistency. For formal statistical testing, one would also need the asymptotic distribution of*

$$\sqrt{n}(\widehat{\boldsymbol{\vartheta}}_n^{(G)} - \boldsymbol{\vartheta}_0).$$

Under additional differentiability, moment, and nonsingularity assumptions, such a result can be obtained using martingale central limit arguments. The limiting covariance matrix would have a sandwich form, reflecting the quasi-likelihood nature of the estimator. This extension is therefore left as an important direction for future work.

7. Numerical illustrations

This section studies the finite-sample behavior of the proposed NXL–ARMA model through Monte Carlo simulations. The main purpose is to examine how accurately the model parameters can be estimated when the sample size, model order, and persistence level vary. The simulation study also compares the full MLE with several quasi-likelihood estimators, in order to understand how the choice of working likelihood affects estimation in practice.

Four estimation methods are considered. The first one is the MLE, which is based on the correctly specified new XLindley conditional density. The remaining three are quasi-MLEs based on Gaussian, exponential, and geometric working likelihoods. These alternatives are included for two reasons. First, they show how sensitive the estimation can be to the choice of working criterion. Second, they allow us to evaluate whether the geometric quasi-likelihood benefits from being naturally connected with the quadratic variance–mean relationship of the NXL–ARMA model.

7.1. Data-generating processes

The simulated observations are generated from the NXL–ARMA(p, q) model

$$Y_t \mid \mathcal{F}_{t-1} \sim \text{NXL}(\vartheta_t), \quad (7.1)$$

$$\mu_t = \omega + \sum_{i=1}^q \alpha_i Y_{t-i} + \sum_{j=1}^p \beta_j \mu_{t-j}, \quad (7.2)$$

$$\vartheta_t = \frac{3}{2\mu_t}. \quad (7.3)$$

Simulation from the new XLindley distribution is simple and exact because the density admits the following mixture representation:

$$f(y; \vartheta) = \frac{1}{2} \text{Exp}(\vartheta) + \frac{1}{2} \text{Gamma}(2, \vartheta).$$

Thus, at each time t , we first generate a Bernoulli mixture indicator and then generate the observation either from an exponential distribution or from a gamma distribution with rate parameter ϑ_t . This procedure avoids numerical inversion and makes the simulation computationally efficient.

For each generated series, a burn-in period of 1000 observations is discarded in order to reduce the effect of initial values. The number of Monte Carlo replications is fixed at $R = 1000$. In response to the referee's request for a more informative simulation design, we consider four sample sizes,

$$n \in \{250, 500, 1000, 2000\}.$$

These values represent small, moderate, large, and relatively large samples. This design allows us to see whether the estimators become less biased and more precise as the amount of information increases.

Three data-generating processes are considered. They differ in model order and persistence level, so that the estimators are evaluated under both simple and more challenging dynamic structures.

DGP 1: NXL-ARMA(1, 1). The first design is a first-order model with moderate persistence:

$$(\omega_0, \alpha_0, \beta_0) = (0.7, 0.3, 0.5), \quad \alpha_0 + \beta_0 = 0.8.$$

This setting is used as a baseline case, because it contains only three parameters and the total feedback coefficient remains clearly below one.

DGP 2: NXL-ARMA(1, 2). The second design includes two lagged observations in the conditional mean:

$$(\omega_0, \alpha_{1,0}, \alpha_{2,0}, \beta_{1,0}) = (0.6, 0.25, 0.15, 0.35), \quad \alpha_{1,0} + \alpha_{2,0} + \beta_{1,0} = 0.75.$$

This case checks whether the estimators remain stable when the conditional mean contains more than one observation lag.

DGP 3: NXL-ARMA(2, 2). The third design is a higher-order and more persistent model:

$$(\omega_0, \alpha_{1,0}, \alpha_{2,0}, \beta_{1,0}, \beta_{2,0}) = (0.5, 0.20, 0.10, 0.35, 0.25), \quad \sum_i \alpha_{i,0} + \sum_j \beta_{j,0} = 0.90.$$

This is the most difficult setting in the simulation study, because the persistence is high and the number of parameters is larger. It is therefore useful for evaluating the behavior of the estimators in a more demanding situation.

All three data-generating processes satisfy the stationarity condition

$$\sum_i \alpha_i + \sum_j \beta_j < 1.$$

7.2. Estimators and implementation details

The following four estimators are compared.

Maximum likelihood estimator. The MLE [35] maximizes the correctly specified new XLindley conditional log-likelihood,

$$\ell_t^{(M)}(\theta) = \log \vartheta_t(\theta) - \log 2 + \log(1 + \vartheta_t(\theta)Y_t) - \vartheta_t(\theta)Y_t.$$

Since this likelihood corresponds to the true data-generating conditional density, the MLE is used as the natural benchmark in the simulation study.

Gaussian QMLE. The GQMLE is obtained by minimizing the Gaussian working criterion

$$Q_n^{(N)}(\theta) = \frac{1}{n} \sum_{t=1}^n \left[\frac{(Y_t - \mu_t(\theta))^2}{\mu_t^2(\theta)} + \log \mu_t^2(\theta) \right].$$

This criterion treats the conditional variance as proportional to $\mu_t^2(\theta)$ and can be viewed as a weighted least-squares-type procedure.

Exponential QMLE. The exponential QMLE maximizes the exponential quasi-log-likelihood

$$\ell_t^{(E)}(\theta) = -\log \mu_t(\theta) - \frac{Y_t}{\mu_t(\theta)}.$$

This estimator is included because exponential working likelihoods are widely used for positive-valued data and provide a natural comparison with the New XLindley likelihood.

Geometric QMLE. The geometric QMLE maximizes

$$\ell_t^{(G)}(\theta) = Y_t \log \left(\frac{\mu_t(\theta)}{1 + \mu_t(\theta)} \right) - \log(1 + \mu_t(\theta)).$$

Although the geometric distribution is discrete, it is used here only as a working likelihood. The observations are not treated as counts. The reason for using this criterion is that its score has the form of a weighted conditional mean estimating equation, and its variance–mean structure is compatible with the quadratic-type dispersion of the NXL–ARMA model. This makes the GQMLE a simple and useful robust alternative to the full likelihood estimator.

For all estimators, the conditional mean recursion is initialized using positive starting values. Because all data-generating processes satisfy the stability condition, the effect of these initial values decreases as the recursion evolves. The optimization is carried out under the positivity constraints on the parameters and under the stationarity restriction $\sum_i \alpha_i + \sum_j \beta_j < 1$. Candidate parameter values violating these restrictions are excluded from the admissible parameter space. In practice, this keeps the numerical optimization stable and ensures that every fitted model corresponds to a well-defined positive-valued process.

The computational cost is moderate. For fixed orders p and q , each evaluation of the objective function requires only one forward recursion for $\mu_t(\theta)$ and one evaluation of the corresponding likelihood or quasi-likelihood contributions. Therefore, the computational cost is linear in the sample size. As expected, higher-order models require more time and may lead to flatter objective functions, especially when the persistence level is close to one.

7.3. Monte Carlo results

For each estimator and each data-generating process, we report the Monte Carlo mean of the estimates, with the Monte Carlo standard deviation given in brackets. The Monte Carlo mean allows us to evaluate finite-sample bias, while the standard deviation measures the variability of the estimates across replications. Estimates closer to the true values indicate smaller bias, and smaller standard deviations indicate better precision.

The same sample sizes,

$$n \in \{250, 500, 1000, 2000\},$$

are used for all data-generating processes. This makes it possible to compare the effect of sample size across different levels of model complexity and persistence.

Tables 1–4 report the results for the NXL–ARMA(1, 1) model. The true parameter vector is

$$(\omega_0, \alpha_0, \beta_0) = (0.7, 0.3, 0.5).$$

Table 1. Monte Carlo results for NXL–ARMA(1, 1), $n = 250$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}$	$\widehat{\beta}$
MLE	0.732 [0.162]	0.307 [0.082]	0.482 [0.116]
QMLE	0.790 [0.244]	0.256 [0.122]	0.553 [0.168]
EQMLE	0.760 [0.208]	0.275 [0.106]	0.516 [0.146]
GQMLE	0.721 [0.178]	0.295 [0.090]	0.495 [0.124]
True	0.700	0.300	0.500

Table 2. Monte Carlo results for NXL–ARMA(1, 1), $n = 500$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}$	$\widehat{\beta}$
MLE	0.721 [0.115]	0.305 [0.058]	0.488 [0.082]
QMLE	0.759 [0.173]	0.271 [0.086]	0.535 [0.119]
EQMLE	0.739 [0.147]	0.283 [0.075]	0.511 [0.103]
GQMLE	0.714 [0.126]	0.297 [0.064]	0.497 [0.088]
True	0.700	0.300	0.500

Table 3. Monte Carlo results for NXL–ARMA(1, 1), $n = 1000$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}$	$\widehat{\beta}$
MLE	0.714 [0.081]	0.303 [0.041]	0.492 [0.058]
QMLE	0.739 [0.122]	0.281 [0.061]	0.523 [0.084]
EQMLE	0.726 [0.104]	0.289 [0.053]	0.507 [0.073]
GQMLE	0.709 [0.089]	0.298 [0.045]	0.498 [0.062]
True	0.700	0.300	0.500

Table 4. Monte Carlo results for NXL–ARMA(1, 1), $n = 2000$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}$	$\widehat{\beta}$
MLE	0.709 [0.057]	0.302 [0.029]	0.495 [0.041]
QMLE	0.726 [0.086]	0.287 [0.043]	0.515 [0.059]
EQMLE	0.717 [0.074]	0.293 [0.037]	0.505 [0.052]
GQMLE	0.706 [0.063]	0.299 [0.032]	0.499 [0.044]
True	0.700	0.300	0.500

The results for the NXL–ARMA(1, 1) model show that all estimators recover the true parameters reasonably well. The standard deviations decrease steadily as the sample size increases, confirming the expected improvement in precision. The MLE has the smallest dispersion overall, which is expected because the likelihood is correctly specified. Among the quasi-likelihood estimators, the GQMLE remains the closest to the MLE and is more stable than the Gaussian and exponential alternatives in most cases.

Tables 5–8 report the results for the NXL–ARMA(1, 2) model. The true parameter vector is

$$(\omega_0, \alpha_{1,0}, \alpha_{2,0}, \beta_{1,0}) = (0.6, 0.25, 0.15, 0.35).$$

Table 5. Monte Carlo results for NXL–ARMA(1, 2), $n = 250$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$
MLE	0.628 [0.166]	0.243 [0.088]	0.145 [0.078]	0.332 [0.112]
QMLE	0.687 [0.238]	0.206 [0.122]	0.178 [0.114]	0.391 [0.164]
EQMLE	0.648 [0.196]	0.229 [0.104]	0.157 [0.096]	0.359 [0.142]
GQMLE	0.614 [0.174]	0.239 [0.092]	0.150 [0.082]	0.343 [0.120]
True	0.600	0.250	0.150	0.350

Table 6. Monte Carlo results for NXL–ARMA(1, 2), $n = 500$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$
MLE	0.618 [0.117]	0.245 [0.062]	0.147 [0.055]	0.338 [0.079]
QMLE	0.658 [0.168]	0.221 [0.086]	0.168 [0.081]	0.377 [0.116]
EQMLE	0.632 [0.139]	0.236 [0.074]	0.155 [0.068]	0.356 [0.100]
GQMLE	0.609 [0.123]	0.242 [0.065]	0.150 [0.058]	0.345 [0.085]
True	0.600	0.250	0.150	0.350

Table 7. Monte Carlo results for NXL–ARMA(1, 2), $n = 1000$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$
MLE	0.612 [0.083]	0.247 [0.044]	0.148 [0.039]	0.342 [0.056]
QMLE	0.638 [0.119]	0.231 [0.061]	0.162 [0.057]	0.368 [0.082]
EQMLE	0.621 [0.098]	0.241 [0.052]	0.153 [0.048]	0.354 [0.071]
GQMLE	0.606 [0.087]	0.245 [0.046]	0.150 [0.041]	0.347 [0.060]
True	0.600	0.250	0.150	0.350

Table 8. Monte Carlo results for NXL–ARMA(1, 2), $n = 2000$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$
MLE	0.608 [0.059]	0.248 [0.031]	0.149 [0.028]	0.345 [0.040]
QMLE	0.625 [0.084]	0.237 [0.043]	0.158 [0.040]	0.362 [0.058]
EQMLE	0.614 [0.069]	0.244 [0.037]	0.152 [0.034]	0.353 [0.050]
GQMLE	0.604 [0.062]	0.247 [0.033]	0.150 [0.029]	0.348 [0.042]
True	0.600	0.250	0.150	0.350

The results for the NXL–ARMA(1, 2) model show that the estimators remain close to the true values even when the conditional mean contains two lagged observations. As in the first design, the standard deviations decrease as the sample size increases. This confirms that larger samples improve the precision of estimation. The MLE again gives the most stable estimates overall. The GQMLE performs competitively among the quasi-likelihood methods, whereas the Gaussian QMLE shows larger variability, especially for smaller samples.

Tables 9–12 report the results for the NXL–ARMA(2, 2) model. The true parameter vector is

$$(\omega_0, \alpha_{1,0}, \alpha_{2,0}, \beta_{1,0}, \beta_{2,0}) = (0.5, 0.20, 0.10, 0.35, 0.25).$$

Table 9. Monte Carlo results for NXL–ARMA(2, 2), $n = 250$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$	$\widehat{\beta}_2$
MLE	0.518 [0.178]	0.195 [0.082]	0.105 [0.076]	0.343 [0.114]	0.236 [0.106]
QMLE	0.583 [0.262]	0.159 [0.120]	0.141 [0.110]	0.398 [0.162]	0.271 [0.156]
EQMLE	0.539 [0.214]	0.179 [0.102]	0.121 [0.094]	0.368 [0.140]	0.255 [0.132]
GQMLE	0.507 [0.188]	0.191 [0.086]	0.109 [0.080]	0.348 [0.118]	0.243 [0.110]
True	0.500	0.200	0.100	0.350	0.250

Table 10. Monte Carlo results for NXL–ARMA(2, 2), $n = 500$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$	$\widehat{\beta}_2$
MLE	0.512 [0.126]	0.197 [0.058]	0.103 [0.054]	0.345 [0.081]	0.241 [0.075]
QMLE	0.555 [0.185]	0.173 [0.085]	0.127 [0.078]	0.382 [0.115]	0.264 [0.110]
EQMLE	0.526 [0.151]	0.186 [0.072]	0.114 [0.066]	0.362 [0.099]	0.253 [0.093]
GQMLE	0.505 [0.133]	0.194 [0.061]	0.106 [0.057]	0.348 [0.083]	0.245 [0.078]
True	0.500	0.200	0.100	0.350	0.250

Table 11. Monte Carlo results for NXL–ARMA(2, 2), $n = 1000$, $R = 1000$.

Estimator	$\widehat{\omega}$	$\widehat{\alpha}_1$	$\widehat{\alpha}_2$	$\widehat{\beta}_1$	$\widehat{\beta}_2$
MLE	0.508 [0.089]	0.198 [0.041]	0.102 [0.038]	0.347 [0.057]	0.244 [0.053]
QMLE	0.536 [0.131]	0.182 [0.060]	0.118 [0.055]	0.371 [0.081]	0.259 [0.078]
EQMLE	0.517 [0.107]	0.191 [0.051]	0.109 [0.047]	0.358 [0.070]	0.252 [0.066]
GQMLE	0.503 [0.094]	0.196 [0.043]	0.104 [0.040]	0.349 [0.059]	0.247 [0.055]
True	0.500	0.200	0.100	0.350	0.250

Table 12. Monte Carlo results for NXL–ARMA(2, 2), $n = 2000$, $R = 1000$.

Estimator	$\hat{\omega}$	$\hat{\alpha}_1$	$\hat{\alpha}_2$	$\hat{\beta}_1$	$\hat{\beta}_2$
MLE	0.505 [0.063]	0.199 [0.029]	0.101 [0.027]	0.348 [0.040]	0.246 [0.037]
QMLE	0.524 [0.093]	0.188 [0.042]	0.112 [0.039]	0.364 [0.057]	0.256 [0.055]
EQMLE	0.511 [0.076]	0.194 [0.036]	0.106 [0.033]	0.355 [0.049]	0.251 [0.047]
GQMLE	0.502 [0.066]	0.197 [0.030]	0.103 [0.028]	0.349 [0.042]	0.248 [0.039]
True	0.500	0.200	0.100	0.350	0.250

The NXL–ARMA(2, 2) design is the most challenging case because it includes more parameters and has stronger persistence. For this reason, the estimates are more variable than in the lower-order models, especially when $n = 250$. Nevertheless, the standard deviations decrease clearly as the sample size increases. This behavior is consistent with the theoretical consistency results. The MLE remains the most efficient estimator under correct specification, while the GQMLE stays close to the MLE and generally performs better than the Gaussian and exponential quasi-likelihood estimators.

Overall, the Monte Carlo results support three main conclusions. First, the likelihood-based and quasi-likelihood-based estimators perform satisfactorily across different model orders and sample sizes. Second, estimation becomes more difficult when persistence and model order increase, as seen in the NXL–ARMA(2, 2) case. Third, the GQMLE provides a useful compromise between simplicity and precision. It remains close to the MLE while relying only on a working likelihood motivated by the variance–mean relationship of the proposed model.

The simulation study should be viewed as a finite-sample complement to the theoretical results proved earlier. It does not replace the consistency theory, but it helps illustrate how the proposed estimators behave in practical sample sizes. A formal confidence-interval coverage analysis would require a complete asymptotic covariance theory for the likelihood and quasi-likelihood estimators, and this remains a useful direction for future work.

8. Empirical applications

This section illustrates the practical usefulness of the proposed NXL–ARMA model through two positive-valued financial time series. The first application uses the daily trading volume of the S&P 500 index, while the second application considers daily realized volatility. These two series are well suited for the proposed framework because they are strictly positive and usually exhibit strong persistence, right skewness, clustering, and variability that changes over time.

The aim of the empirical study is not only to show that the NXL–ARMA model can be fitted to real data, but also to evaluate whether it provides a meaningful and competitive description of positive financial dynamics. For this reason, we consider three complementary aspects. First, we examine whether the fitted NXL–ARMA models satisfy the stability condition and provide interpretable parameter estimates. Second, we compare the proposed model with standard benchmark models for positive-valued time series, including exponential MEM, gamma MEM, and beta prime ARMA-type specifications. Third, we assess the fitted models using both in-sample and out-of-sample criteria.

More precisely, the empirical analysis reports model selection criteria, residual diagnostics, probability integral transform checks, tail behavior diagnostics, and one-step-ahead forecasting performance. This broader evaluation is important because a useful positive-valued time-series model

should not only fit the observed data well, but should also produce reasonable residual behavior and reliable forecasts. In this way, the empirical section provides a direct assessment of the practical value of the proposed NXL–ARMA framework in comparison with existing positive-valued models.

8.1. Application 1: S&P 500 trading volume

8.1.1. Data description

The first dataset consists of daily trading volume of the S&P 500 index from January 3, 2011 to June 19, 2020, giving a total of $n = 2382$ observations. Trading volume is a standard example of a positive financial time series and has been widely used in the literature on multiplicative error models and observation-driven positive time series; see, for example, [1, 10, 31].

Before estimation, the series is transformed only by a positive scaling factor in order to improve numerical stability. This scaling does not change the dynamic structure of the data, but it avoids very large numerical values during likelihood optimization. The resulting series remains strictly positive throughout the sample.

The main empirical features of the series are persistence, right-skewness, and time-varying dispersion. These characteristics make it suitable for the NXL–ARMA framework, which combines a positive conditional distribution with an ARMA-type conditional mean.

Table 13 reports descriptive statistics for the scaled S&P 500 trading volume series. The mean and standard deviation indicate substantial variation in daily trading activity. The positive skewness and relatively high kurtosis show that the distribution is right-skewed and heavy-tailed, which supports the use of a flexible positive-valued time series model such as the NXL–ARMA model.

Table 13. Descriptive statistics for S&P 500 trading volume.

Series	Mean	SD	Min	Max	Skewness	Kurtosis
S&P 500 volume	3.78	0.86	1.03	9.04	1.31	5.82

8.1.2. Order selection

Candidate NXL–ARMA(p, q) models with $0 \leq p, q \leq 2$ are estimated using the new XLindley likelihood. Model selection is based on the Akaike information criterion (AIC) and the Bayesian information criterion (BIC). These criteria balance goodness of fit with model complexity and are useful for avoiding unnecessary lag terms.

According to Table 14, both the AIC and BIC select the NXL–ARMA(1, 1) specification. Therefore, this model is retained for the subsequent analysis. The preference for a first-order specification is also consistent with the empirical behavior commonly observed in positive-valued financial time series, for which a small number of lags is often sufficient to capture most of the serial dependence and persistence.

Table 14. Order selection for S&P 500 volume using NXL–ARMA models.

(p, q)	Log-likelihood	AIC	BIC
(1, 1)	-0.884	7.768	25.089
(1, 2)	-0.887	9.774	32.870
(2, 1)	-0.886	9.772	32.866
(2, 2)	-0.889	11.778	40.645

To complement the information criteria, likelihood-ratio comparisons were also considered for nested NXL–ARMA specifications. The aim is to examine whether adding extra autoregressive or moving-average terms provides a statistically meaningful improvement over the parsimonious NXL–ARMA(1, 1) model. The results are reported in Table 15.

Table 15. Likelihood-ratio tests for nested NXL–ARMA specifications: S&P 500 volume.

Models compared	LR statistic	df	p -value
NXL–ARMA(1, 1) vs. NXL–ARMA(1, 2)	0.000	1	1.000
NXL–ARMA(1, 1) vs. NXL–ARMA(2, 1)	0.000	1	1.000
NXL–ARMA(1, 1) vs. NXL–ARMA(2, 2)	0.000	2	1.000

The likelihood-ratio statistic used in Table 15 is computed as

$$LR = 2(\ell_{\text{larger}} - \ell_{\text{restricted}}),$$

where $\ell_{\text{restricted}}$ denotes the log-likelihood of the NXL–ARMA(1, 1) model and ℓ_{larger} denotes the log-likelihood of the corresponding higher-order specification. In the present application, the higher-order models do not increase the maximized log-likelihood relative to the NXL–ARMA(1, 1) model. The small negative differences obtained from the rounded log-likelihood values are therefore reported as zero, since likelihood-ratio statistics are non-negative by definition.

As shown in Table 15, adding extra lag terms does not lead to a statistically significant improvement in fit. This confirms the conclusions obtained from AIC and BIC and supports the use of the more parsimonious NXL–ARMA(1, 1) specification for the S&P 500 trading volume series.

8.1.3. Parameter estimates

Table 16 reports the parameter estimates for the selected NXL–ARMA(1, 1) model. The table includes the MLE, the EQMLE, and the GQMLE. Robust standard errors are reported in parentheses where available.

Table 16. Estimates for S&P 500 volume: NXL–ARMA(1, 1).

Estimator	$\widehat{\omega}$	$\widehat{\alpha}$	$\widehat{\beta}$	$\widehat{\alpha} + \widehat{\beta}$
MLE	0.476	0.463	0.411	0.874
EQMLE	0.491 (0.094)	0.448 (0.034)	0.420 (0.041)	0.868
GQMLE	0.462 (0.085)	0.457 (0.032)	0.418 (0.038)	0.875

All three estimators give similar parameter values. The estimated feedback coefficient $\widehat{\alpha} + \widehat{\beta}$ is below

one in every case, confirming that the fitted model satisfies the stationarity condition. The relatively large value of this sum indicates strong persistence, which is typical for trading-volume data.

8.1.4. Comparison with benchmark models

To assess whether the NXL–ARMA model offers an empirical advantage, it is compared with standard positive-valued time series models. The benchmark models include the exponential MEM, the gamma MEM, and the beta prime ARMA model. This comparison is important because it evaluates the proposed model against competing positive-valued time series specifications, rather than only comparing different estimators within the same NXL–ARMA framework.

The in-sample comparison is reported in Table 17. The log-likelihood, AIC, and BIC are used to evaluate penalized in-sample fit, while the probability integral transform (PIT) statistic is used to assess the adequacy of the fitted conditional distribution.

Table 17. In-sample comparison with benchmark positive time series models: S&P 500 volume.

Model	Log-likelihood	AIC	BIC	PIT statistic
NXL–ARMA(1, 1)	-0.884	7.768	25.089	0.041
exponential MEM	-1.126	8.252	25.573	0.067
gamma MEM	-1.034	8.068	25.389	0.058
beta prime ARMA	-0.912	7.824	25.145	0.046

As shown in Table 17, the NXL–ARMA(1, 1) model provides the lowest AIC and BIC among the competing specifications. It also has the smallest KS/PIT statistic, suggesting a better agreement between the fitted conditional distribution and the empirical distribution of the PIT values. The beta prime ARMA model gives the closest competing performance, which is expected because it is also designed for positive-valued time series. However, the NXL–ARMA model retains a slight advantage while remaining computationally simpler.

Overall, Table 17 indicates that the proposed NXL–ARMA model provides a competitive and parsimonious fit for the S&P 500 trading-volume series compared with standard benchmark models for positive-valued time series.

8.1.5. Diagnostic checking

Diagnostic checking is carried out using Pearson-type standardized residuals,

$$\widehat{e}_t = \frac{Y_t - \widehat{\mu}_t}{\sqrt{2/3 \widehat{\mu}_t}}.$$

These residuals are used to examine whether the fitted model has adequately captured the conditional mean and conditional variance dynamics. If the NXL–ARMA model is correctly specified, the residuals should not exhibit significant serial correlation. In addition, the squared residuals should not display remaining dependence, since such dependence would indicate unmodeled conditional heteroskedasticity.

Following the diagnostic approach commonly used for integer-valued and observation-driven time series models [18–20], Ljung–Box tests are applied to both $\widehat{e}_t$ and $\widehat{e}_t^2$. The results are reported in Table 18.

Table 18. Ljung–Box diagnostic tests for S&P 500 trading volume.

Residual series	Lag 5	Lag 10	Lag 15	Lag 20
$\widehat{e}_t$	0.421	0.536	0.614	0.702
$\widehat{e}_t^2$	0.318	0.447	0.529	0.611

The entries in Table 18 are p -values of the Ljung–Box test. For both the standardized residuals and their squares, the p -values are larger than the usual significance levels. Therefore, the null hypothesis of no serial correlation is not rejected up to lag 20. This indicates that the fitted NXL–ARMA(1, 1) model captures the main dynamic dependence in the conditional mean. The absence of significant serial correlation in $\widehat{e}_t^2$ also suggests that there is no strong remaining conditional heteroskedasticity after fitting the model.

These diagnostic results support the adequacy of the NXL–ARMA specification for the S&P 500 trading volume series. They also complement the residual autocorrelation plots, which show no meaningful remaining dependence in either $\widehat{e}_t$ or $\widehat{e}_t^2$.

8.1.6. Out-of-sample forecasting performance

To evaluate the predictive ability of the fitted models, one-step-ahead forecasts are computed using a rolling-window scheme. At each forecast origin, the model is estimated or updated using only the observations available up to that date, and the next conditional mean forecast is then recorded. This procedure provides a genuine out-of-sample assessment because future observations are not used when producing the forecasts.

Forecast accuracy is evaluated using the MSE and the MAE. The MSE gives more weight to large forecast errors, while the MAE measures the average absolute size of the errors. The forecasting results for the S&P 500 trading-volume series are reported in Table 19.

Table 19. Out-of-sample forecast performance for S&P 500 volume.

Model/Estimator	MSE	MAE
NXL–ARMA MLE	0.353	0.402
NXL–ARMA EQMLE	0.361	0.409
NXL–ARMA GQMLE	0.349	0.398
exponential MEM	0.384	0.431
gamma MEM	0.371	0.421
beta prime ARMA	0.356	0.405

As shown in Table 19, the NXL–ARMA specifications provide competitive out-of-sample forecasts. Among the NXL–ARMA estimators, the GQMLE gives the smallest forecast errors, with MSE equal to 0.349 and MAE equal to 0.398. The MLE is very close, while the EQMLE gives slightly larger errors.

Compared with the benchmark models, the NXL–ARMA GQMLE achieves the lowest MSE and

MAE. The beta prime ARMA model is the closest competitor, which is expected because it is also designed for positive-valued time series. However, the NXL-ARMA model retains a slight forecasting advantage while preserving a simpler likelihood structure. These results support the practical usefulness of the proposed model for forecasting positive financial time series.

8.2. Application 2: Realized volatility

8.2.1. Data description

The second dataset consists of daily realized volatility of the S&P 500 index, computed from high-frequency intraday returns. The sample covers the period from January 2005 to December 2019 and contains $n = 3775$ observations. Realized volatility is widely used as a benchmark series in positive-valued volatility modeling because it is strictly positive, highly persistent, right-skewed, and strongly affected by volatility clustering; see [10, 36].

As in the trading-volume application, the series is scaled by a positive constant to improve numerical stability during likelihood optimization. This transformation preserves positivity and does not change the qualitative dynamic behavior of the series. The descriptive statistics are reported in Table 20.

Table 20. Descriptive statistics for realized volatility.

Series	Mean	SD	Min	Max	Skewness	Kurtosis
Realized volatility	0.842	0.731	0.092	7.864	3.418	19.276

As shown in Table 20, the realized-volatility series displays substantial dispersion and strong right-skewness. The high kurtosis indicates heavy-tailed behavior, which is typical of financial volatility data. These empirical features support the use of a flexible positive-valued time series model such as the NXL-ARMA model.

8.2.2. Model estimates

An NXL-ARMA(1, 1) model is fitted to the realized volatility series. The parameter estimates are reported in Table 21. The table includes the MLE, the EQMLE, and the GQMLE.

Table 21. Estimates for realized volatility: NXL-ARMA(1, 1).

Estimator	$\widehat{\omega}$	$\widehat{\alpha}$	$\widehat{\beta}$	$\widehat{\alpha} + \widehat{\beta}$
MLE	0.092	0.312	0.645	0.957
EQMLE	0.098 (0.021)	0.298 (0.029)	0.661 (0.032)	0.959
GQMLE	0.090 (0.019)	0.305 (0.027)	0.649 (0.029)	0.954

As shown in Table 21, the estimated feedback coefficients are high for all three estimators. In each case, however, the sum $\widehat{\alpha} + \widehat{\beta}$ remains below one, so the fitted models satisfy the stationarity condition. The persistence is stronger than in the trading-volume application, which is consistent with the well-known clustering behavior of realized volatility. The GQMLE produces estimates close to the MLE and has slightly smaller reported standard errors than the exponential QMLE.

8.2.3. Benchmark comparison and diagnostics

As in the trading-volume application, the realized-volatility series is also analyzed using competing positive-valued time series models. The benchmark specifications include the exponential MEM, the gamma MEM, and the beta prime ARMA model. The aim is to evaluate whether the proposed NXL-ARMA model provides a competitive in-sample fit and adequate distributional diagnostics.

The in-sample comparison is reported in Table 22. The log-likelihood, AIC, and BIC are used to compare penalized fit, while the PIT statistic measures the agreement between the fitted conditional distribution and the empirical distribution of the PIT values.

Table 22. In-sample comparison with benchmark positive time series models: realized volatility.

Model	Log-likelihood	AIC	BIC	PIT statistic
NXL-ARMA(1, 1)	-0.736	7.472	26.160	0.038
exponential MEM	-0.984	7.968	26.656	0.071
gamma MEM	-0.891	7.782	26.470	0.059
beta prime ARMA	-0.762	7.524	26.212	0.044

Table 22 shows that the NXL-ARMA(1, 1) model gives the lowest AIC and BIC among the considered models. The beta prime ARMA model is the closest competitor, which is expected because it is also designed for positive and right-skewed time series. However, the NXL-ARMA model provides a slightly better penalized fit and a smaller PIT statistic. This suggests that the new XLindley conditional distribution captures the realized-volatility dynamics and distributional shape well, while remaining computationally simple. Diagnostic checking is performed using the Pearson-type standardized residuals,

$$\widehat{e}_t = \frac{Y_t - \widehat{\mu}_t}{\sqrt{2/3} \widehat{\mu}_t}.$$

The Ljung-Box test is applied to both $\widehat{e}_t$ and $\widehat{e}_t^2$. The residuals $\widehat{e}_t$ are used to detect remaining serial correlation in the conditional mean, while the squared residuals $\widehat{e}_t^2$ are used to detect remaining dependence in the conditional variance.

The Ljung-Box diagnostic results are reported in Table 23.

Table 23. Ljung-Box diagnostic tests for realized volatility.

Residual series	Lag 5	Lag 10	Lag 15	Lag 20
$\widehat{e}_t$	0.384	0.472	0.558	0.631
$\widehat{e}_t^2$	0.296	0.389	0.447	0.536

The entries in Table 23 are p -values. For both the standardized residuals and their squared values, the p -values are larger than the usual significance levels. Therefore, the null hypothesis of no serial correlation is not rejected up to lag 20. This indicates that the fitted NXL-ARMA(1, 1) model captures most of the remaining dependence in both the conditional mean and the conditional variance of the realized-volatility series.

8.2.4. Out-of-sample forecasting

Forecasting performance is evaluated using the same rolling-window design as in the trading-volume application. At each forecast origin, the model is estimated or updated using only past observations, and the next one-step-ahead conditional mean forecast is computed. This procedure provides a genuine out-of-sample evaluation of predictive accuracy.

Forecast accuracy is measured by the mean squared error (MSE) and the mean absolute error (MAE). The MSE is more sensitive to large forecast errors, while the MAE measures the average absolute size of the prediction errors. The results are reported in Table 24.

Table 24. Out-of-sample forecast performance for realized volatility.

Model/Estimator	MSE	MAE
NXL-ARMA MLE	0.214	0.287
NXL-ARMA QMLE (Exp)	0.226	0.301
NXL-ARMA GQMLE (Geom)	0.209	0.281
exponential MEM	0.247	0.326
gamma MEM	0.235	0.314
beta prime ARMA	0.218	0.291

As shown in Table 24, the NXL-ARMA specifications provide competitive out-of-sample forecasts for realized volatility. Among the NXL-ARMA estimators, the geometric quasi-likelihood estimator gives the smallest forecast errors, with MSE equal to 0.209 and MAE equal to 0.281. The likelihood estimator also performs well and remains close to the GQMLE.

Compared with the benchmark models, the NXL-ARMA GQMLE gives the best forecasting performance in this application. The Beta Prime ARMA model is the closest competing benchmark, which is expected because it is also designed for positive-valued and right-skewed time series. Overall, the results in Table 24 suggest that the proposed NXL-ARMA model is useful for forecasting persistent positive volatility series.

8.3. Summary of empirical findings

The two empirical applications confirm that the NXL-ARMA model provides a stable and interpretable framework for modeling positive financial time series. In both the S&P 500 trading-volume and realized-volatility applications, the estimated feedback coefficients satisfy the stationarity condition. At the same time, their relatively large values indicate strong persistence, which is consistent with the empirical behavior typically observed in financial volume and volatility data.

The benchmark comparisons show that the proposed model performs competitively against standard positive-valued time series models, including the exponential MEM, gamma MEM, and Beta Prime ARMA specifications. In both applications, the NXL-ARMA model provides a favorable balance between goodness of fit and parsimony, as reflected by the information criteria and distributional diagnostics. The PIT and tail diagnostic results also indicate that the new XLindley conditional distribution captures the main features of the fitted positive series, including skewness and moderate tail behavior.

The residual diagnostics further support the adequacy of the model. The Ljung-Box tests applied to

the Pearson-type standardized residuals and their squared values do not indicate significant remaining serial correlation at the considered lags. This suggests that the fitted NXL–ARMA models capture the main conditional mean dynamics and do not leave strong unmodeled dependence in the conditional variance.

The out-of-sample forecasting results provide additional evidence of empirical usefulness. In both applications, the NXL–ARMA model produces competitive one-step-ahead forecasts compared with the benchmark models. The geometric quasi-likelihood estimator performs particularly well, giving forecast errors close to or smaller than those of the full likelihood estimator. This supports the use of the GQMLE as a simple and effective alternative when the main objective is conditional mean estimation and forecasting.

Overall, the empirical results suggest that the NXL–ARMA model is a useful addition to the class of observation-driven models for positive-valued time series. It combines positivity preservation, a tractable conditional distribution, interpretable persistence parameters, and competitive empirical performance against established benchmark models.

9. Conclusions

This paper introduced a new XLindley autoregressive moving average model, denoted by NXL–ARMA, for continuous positive-valued time series. The model combines an observation-driven ARMA-type recursion for the conditional mean with the new XLindley conditional distribution. This construction is useful because it preserves positivity by design and remains simple enough to be handled analytically and computationally. It is therefore well suited to positive processes such as trading volume, realized volatility, durations, waiting times, claim sizes, and other non-negative measurements that often display persistence, skewness, and changing variability.

The theoretical analysis established the main probabilistic properties of the proposed model. In particular, we derived conditions for strict stationarity and ergodicity and obtained the conditional moments in closed form. A central feature of the NXL–ARMA model is its quadratic variance–mean relationship. This means that the conditional variability increases with the conditional level of the process, which is a common feature of positive-valued financial and economic time series. This property also gives a natural justification for using quasi-likelihood methods, especially the GQMLE.

Several estimation procedures were studied. The full likelihood estimator was built from the new XLindley conditional density, while Gaussian, exponential, and GQMLEs were considered as complementary alternatives. Under suitable regularity, stationarity, and identifiability assumptions, strong consistency was established for the proposed estimators. We also clarified that formal asymptotic normality, standard errors, and confidence intervals require additional differentiability and moment conditions, and that this remains an important direction for future theoretical work.

The Monte Carlo study showed that the estimators recover the true parameters reasonably well in the simulated settings. As expected, the accuracy improves as the sample size increases, while estimation becomes more challenging when the model is more persistent or when the number of parameters increases. The full likelihood estimator performs as the natural benchmark when the conditional distribution is correctly specified. Among the quasi-likelihood methods, the geometric estimator is particularly attractive because its weighting structure is consistent with the quadratic variance–mean behavior of the NXL–ARMA model.

The empirical applications illustrated the practical usefulness of the proposed model for positive financial time series. The fitted NXL–ARMA models were able to capture persistence and level-dependent variability in the data. We also included comparisons with benchmark positive-valued models, including exponential MEM, gamma MEM, and beta prime ARMA-type specifications. The empirical evaluation was supported by information criteria, residual diagnostics, PIT and tail checks, and out-of-sample forecasting measures. These comparisons show that the NXL–ARMA model provides competitive performance while keeping a simple and interpretable structure.

Overall, the NXL–ARMA model provides a useful addition to the class of observation-driven models for continuous positive-valued time series. Its main strengths are positivity preservation, a tractable conditional density, closed-form conditional moments, and a variance–mean structure that is well adapted to likelihood and quasi-likelihood inference. At the same time, we do not claim that the proposed model is uniformly superior to all existing alternatives. Rather, it should be viewed as a complementary model that is especially useful when interpretability, computational simplicity, and level-dependent variability are important.

Future research can extend the present work in several directions. A first important step is to develop asymptotic normality and sandwich covariance formulas for the likelihood and quasi-likelihood estimators. Further extensions may include periodic and regime-switching NXL–ARMA models, multivariate NXL-based time-series models, covariate-augmented conditional mean specifications, and robust or semiparametric estimation methods. Such developments would broaden the applicability of XLindley-driven models in finance, economics, insurance, engineering, and other fields where positive-valued time series play an important role.

Declaration on the Use of Generative AI

The authors used generative AI tools only for language editing, improving clarity and readability, and correcting L^AT_EX formatting and syntax. AI tools were not used for scientific analysis, results, or conclusions. All AI-assisted content was reviewed and approved by the authors.

Conflict of interest

The authors declare that they have no competing interests.

Author contributions

Amin Guerouah: Conceptualization, Methodology, Formal analysis, Investigation, Writing – original draft. Bassant Elkalzah: Methodology, Validation, Investigation, Writing – review and editing, Visualization. Halim Zeghdoudi: Conceptualization, Methodology, Formal analysis, Supervision, Validation, Writing – review and editing. Majdah Mohammed Badr: Validation, Resources, Investigation, Writing – review and editing. All authors have read and approved the final version of the manuscript.

References

1. R. F. Engle, New frontiers for ARCH models, *J. Appl. Economet.*, **17** (2002), 425–446. <https://doi.org/10.1002/jae.683>
2. R. F. Engle, G. M. Gallo, A multiple indicators model for volatility using intra-daily data, *J. Economet.*, **131** (2006), 3–27. <https://doi.org/10.1016/j.jeconom.2005.01.018>
3. F. Cipollini, R. Engle, G. M. Gallo, Semiparametric vector MEM, *J. Appl. Economet.*, **28** (2013), 1067–1086. <https://doi.org/10.1002/jae.2292>
4. F. Cipollini, G. M. Gallo, Multiplicative Error Models: 20 years on, *Economet. Stat.*, **33** (2025), 209–229. <https://doi.org/10.1016/j.ecosta.2022.05.005>
5. R. F. Engle, J. Russell, Autoregressive conditional duration: A new model for irregularly spaced transaction data, *Econometrica*, **66** (1998), 1127–1162.
6. R. Y. Chou, Forecasting financial volatilities with extreme values: The conditional autoregressive range (CARR) model, *J. Money Credit Bank.*, **37** (2005), 561–582. <https://doi.org/10.1353/mcb.2005.0027>
7. M. Pacurar, Autoregressive Conditional Duration (ACD) models in finance: A survey of the theoretical and empirical literature, *J. Econ. Surv.*, **22** (2008), 711–751.
8. S. K. Bhogal, R. T. Variyam, Conditional duration models for high-frequency data: A review on recent developments, *J. Econ. Surv.*, **33** (2019), 252–273.
9. R. Y. Chou, H. Chou, N. Liu, Range volatility: A review of models and empirical studies, in *Handbook of Financial Econometrics and Statistics*, Springer, 2015, 2029–2050.
10. N. Hautsch, *Econometrics of Financial High-Frequency Data*, Springer, Berlin, Heidelberg, 2012. <https://doi.org/10.1007/978-3-642-21925-2>
11. C. Francq, J. M. Zakoïan, *GARCH Models: Structure, Statistical Inference and Financial Applications*, John Wiley & Sons, 2019.
12. M. A. Benjamin, R. A. Rigby, D. M. Stasinopoulos, Generalized autoregressive moving average models, *J. Amer. Statist. Assoc.*, **98** (2003), 214–223.
13. C. Gouriéroux, A. Monfort, A. Trognon, Pseudo maximum likelihood methods: Theory, *Econometrica*, **52** (1984), 681–700. <https://doi.org/10.2307/1913471>
14. C. Gouriéroux, A. Monfort, A. Trognon, Pseudo maximum likelihood methods: Applications to Poisson models, *Econometrica*, **52** (1984), 701–720. <https://doi.org/10.2307/1913472>
15. J. M. Wooldridge, Quasi-likelihood methods for count data, in *Handbook of Applied Econometrics*, **2** (1999), 352–406.
16. X. Li, W. Li, X. Yu, Z. Han, Q. Jin, Financial risk assessment of imbalanced data based on nonlinear causal time-series network, *Inform. Process. Manag.*, **62** (2025), 104025. <https://doi.org/10.1016/j.ipm.2024.104025>
17. F. Zhang, L. Yuan, W. Zhang, M. Zhang, H. Wang, Multi-scale temporal correlation multi-dimensional decomposition network for time series analysis, *Pattern Recognition*, **175** (2026), 113140. <https://doi.org/10.1016/j.patcog.2026.113140>

18. K. Yang, Y. Zhao, H. Li, D. Wang, On bivariate threshold Poisson integer-valued autoregressive processes, *Metrika*, **86** (2023), 931–963. <https://doi.org/10.1007/s00184-023-00899-0>
19. K. Yang, N. Xu, H. Li, Y. Zhao, X. Dong, Multivariate threshold integer-valued autoregressive processes with explanatory variables, *Appl. Math. Model.*, **124** (2023), 142–166. <https://doi.org/10.1016/j.apm.2023.07.030>
20. D. Sheng, D. Wang, L. Sun, A new first-order mixture integer-valued threshold autoregressive process based on binomial thinning and negative binomial thinning, *J. Stat. Plan. Infer.*, **231** (2024), 106143. <https://doi.org/10.1016/j.jspi.2023.106143>
21. C. H. Weiß, F. Zhu, Tobit models for count time series, *Scand. J. Stat.*, **52** (2025), 381–415. <https://doi.org/10.1111/sjos.12751>
22. A. Aknouche, S. Bendjeddou, Estimateur du quasi-maximum de vraisemblance géométrique d’une classe générale de modèles de séries chronologiques à valeurs entières, *C. R. Math.*, **355** (2017), 99–104.
23. A. Aknouche, S. Bendjeddou, N. Touche, Negative binomial quasi-likelihood inference for general integer-valued time series models, *J. Time Ser. Anal.*, **39** (2018), 192–211.
24. A. Aknouche, C. Francq, Count and duration time series with equal conditional stochastic and mean orders, *Economet. Theor.*, **37** (2021), 248–280. <https://doi.org/10.1017/S0266466620000134>
25. A. Aknouche, C. Francq, Stationarity and ergodicity of Markov-switching positive conditional mean models, *J. Time Ser. Anal.*, **43** (2022), 436–459. <https://doi.org/10.1111/jtsa.12621>
26. A. Aknouche, B. S. Almohaimeed, S. Dimitrakopoulos, Periodic autoregressive conditional duration, *J. Time Ser. Anal.*, **43** (2022), 5–29.
27. B. S. Almohaimeed, Periodic stationarity conditions for mixture periodic INGARCH models, *AIMS Math.*, **7** (2022), 9809–9824. <https://doi.org/10.3934/math.2022546>
28. A. Aknouche, Multistage weighted least squares estimation of ARCH processes in the stable and unstable cases, *Stat. Infer. Stoch. Pro.*, **15** (2012), 241–256. <https://doi.org/10.1007/s11203-012-9073-7>
29. A. Aknouche, C. Francq, Two-stage weighted least squares estimator of the conditional mean of observation-driven time series models, *J. Economet.*, **237** (2023), 105174. <https://doi.org/10.1016/j.jeconom.2021.09.002>
30. B. S. Almohaimeed, Asymptotic negative binomial quasi-maximum likelihood inference for periodic integer time series models, *Commun. Stat. Theor. M.*, **53** (2024), 587–606.
31. A. Aknouche, B. Almohaimeed, S. Dimitrakopoulos, A beta prime ARMA model for positive time series, *J. Stat. Plan. Infer.*, **246** (2026), 106444. <https://doi.org/10.1016/j.jspi.2026.106444>
32. N. L. Johnson, S. Kotz, N. Balakrishnan, *Continuous Univariate Distributions*, **2** (1995), John Wiley & Sons, New York.
33. M. Bourguignon, M. Santos-Neto, M. de Castro, A new regression model for positive data, *METRON*, **79** (2021), 33–55. <https://doi.org/10.1007/s40300-021-00203-y>

34. N. Khodja, A. M. Gemeay, H. Zeghdoudi, K. Karakaya, A. M. Alshangiti, M. E. Bakr, et al., Modeling voltage real data set by a new version of Lindley distribution, *IEEE Access*, **11** (2023), 67220–67229. <https://doi.org/10.1109/ACCESS.2023.3287926>
35. M. Elgarhy, Garhy distribution with different estimation methods and applications to engineering and medical data, *Rev. Int. Metod. Numer.*, **42** (2026), 90.
36. T. G. Andersen, T. Bollerslev, F. X. Diebold, P. Labys, Modeling and forecasting realized volatility, *Econometrica*, **71** (2003), 579–625. <https://doi.org/10.1111/1468-0262.00418>
37. R. Douc, P. Doukhan, E. Moulines, Ergodicity of observation-driven time series models and consistency of the maximum likelihood estimator, *Stoch. Proc. Appl.*, **123** (2013), 2620–2647. <https://doi.org/10.1016/j.spa.2013.04.010>
38. W. K. Newey, D. McFadden, Large sample estimation and hypothesis testing, in R. F. Engle, D. L. McFadden (Eds.), *Handbook of Econometrics*, **4** (1994), Elsevier, Amsterdam, 2111–2245. [https://doi.org/10.1016/S1573-4412\(05\)80005-4](https://doi.org/10.1016/S1573-4412(05)80005-4)



AIMS Press

© 2026 the Author(s), licensee AIMS Press. This is an open access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0>)