

Online Supplement to “Prediction of bubbles in presence of α -stable aggregates moving averages”

Gilles de Truchis, Sébastien Fries, Arthur Thomas

This Online Supplement accompanies the main paper by providing comprehensive Monte Carlo results and detailed application diagnostics. Sections [S.1](#) and [S.2](#) present the complete Monte Carlo study that supports the summary in Section 4 of the main text. Section [S.3](#) provides the full empirical analysis of the OVX index summarized in Section 5. The supplement is organized as follows.

Sections [S.1](#) and [S.2](#): Monte Carlo Study and subsampling. This section documents the finite-sample performance of the minimum distance estimator. Section [S.1](#) reports the estimation accuracy results across three distributional settings (Cauchy, symmetric α -stable, and general α -stable), first for the anticipative AR(1) aggregate of Section 4.1 of the main text (Section [S.1.1](#)), then for the mixed causal-noncausal MAR(1,1) aggregate of Section 2.3 (Section [S.1.2](#)). Section [S.2](#) presents the subsampling methodology used to verify asymptotic normality. Section [S.2.2](#) presents results in the reparameterized space $\vartheta = (\psi_1, \psi_2, \sigma, \pi_1, \alpha)$

Section [S.3](#): Illustration on the OVX Index. This section provides the empirical analysis of the CBOE Crude Oil ETF Volatility Index. Section [S.3.1](#) reports estimation results under three distributional specifications. Section [S.3.2](#) presents the deconvolution into latent components. Section [S.3.3](#) develops the in-sample forecasting exercise for the 2020 oil market disruption, including the pattern-matching algorithm, crash probability computations, and risk threshold analysis. Section [S.3.4](#) collects per-component graphical diagnostics.

S.1. Monte Carlo estimation accuracy

This subsection reports the full Monte Carlo estimation accuracy results that are summarized in Section 4.1 of the main paper. Section [S.1.1](#) covers Case 1, the aggregation of two independent α -stable AR(1) processes as specified in equations (4.1)–(4.2) of the main text. Section [S.1.2](#) extends the exercise to Case 2 of Section 2.3 of the main text, the aggregation of mixed causal-noncausal MAR(1,1) components. Both exercises minimize the discretized minimum-distance criterion (2.14) of the main text under the Assumption C2 weight $w(u, v) = (\kappa/\pi) \exp(-\kappa(u^2 + v^2))$, $\kappa = 1$, so that the reported estimator is the one whose asymptotics are established by Proposition 2.1.

S.1.1. Case 1: aggregation of anticipative AR(1) processes

The observed process is generated by the aggregation of two independent α -stable AR(1) processes as specified in equations (4.1)–(4.2) of the main text. Table S.1 reports the bias, root mean square error (RMSE), and mean relative error (MRE) across all scenarios and sample sizes.

Table S.1: Monte Carlo estimation accuracy for α -stable AR(1) aggregates (1,000 replications)										
θ	True	$T = 250$			$T = 500$			$T = 1000$		
		Bias	RMSE	MRE	Bias	RMSE	MRE	Bias	RMSE	MRE
<i>Panel A: Cauchy ($\mathcal{S1S}$, $\alpha = 1$, $\beta = 0$)</i>										
ψ_1	0.800	0.023	0.072	0.068	0.011	0.050	0.046	0.003	0.032	0.030
ς_1	0.700	−0.065	0.371	0.425	−0.020	0.279	0.313	0.004	0.191	0.212
ψ_2	0.300	0.003	0.180	0.488	−0.015	0.143	0.375	−0.018	0.111	0.281
ς_2	0.900	0.084	0.342	0.294	0.024	0.247	0.209	0.002	0.175	0.155
<i>Panel B: Symmetric α-stable ($\mathcal{S}\alpha\mathcal{S}$, $\alpha = 1.5$, $\beta = 0$)</i>										
ψ_1	0.800	0.004	0.112	0.118	0.004	0.094	0.095	0.003	0.072	0.070
ς_1	0.700	−0.054	0.385	0.467	−0.034	0.324	0.381	−0.011	0.266	0.303
ψ_2	0.300	0.016	0.231	0.657	−0.009	0.197	0.544	−0.010	0.151	0.395
ς_2	0.900	−0.020	0.264	0.234	−0.018	0.229	0.204	−0.023	0.182	0.158
α	1.500	0.012	0.168	0.091	0.005	0.123	0.065	0.001	0.085	0.045
<i>Panel C: General α-stable ($\mathcal{G}\alpha\mathcal{S}$, $\alpha = 1.5$, $\beta = 0.3$)</i>										
ψ_1	0.800	0.025	0.117	0.122	0.016	0.098	0.099	0.009	0.077	0.075
ς_1	0.700	−0.136	0.372	0.451	−0.082	0.321	0.373	−0.035	0.261	0.295
ψ_2	0.300	−0.079	0.227	0.666	−0.055	0.202	0.572	−0.047	0.175	0.471
ς_2	0.900	−0.024	0.233	0.202	−0.021	0.219	0.193	−0.024	0.185	0.164
α	1.500	−0.026	0.159	0.085	−0.011	0.112	0.060	−0.005	0.080	0.043
β	0.300	0.018	0.270	0.691	0.016	0.187	0.470	0.003	0.129	0.339

Notes: True parameter values are $\psi_1 = 0.8$, $\psi_2 = 0.3$, $\varsigma_1 = 0.7$, $\varsigma_2 = 0.9$. MRE denotes the mean relative error $\mathbb{E}[|\hat{\theta} - \theta_0|/|\theta_0|]$. Estimation minimizes the discretized criterion (2.14) of the main text using the Assumption C2 weight $w(u, v) = (\kappa/\pi) \exp(-\kappa(u^2 + v^2))$ with $\kappa = 1$.

Several patterns emerge from Table S.1. First, the dominant autoregressive coefficient ψ_1 is precisely estimated across all settings, with MRE below 7% in the $\mathcal{S1S}$ case and below 13% even in the $\mathcal{G}\alpha\mathcal{S}$ case for $T = 250$. Second, the smaller coefficient ψ_2 and the combined scale parameters ς_1 , ς_2 exhibit substantially higher relative errors, reflecting the inherent difficulty in disentangling the individual component contributions from the aggregate signal. Third, the tail index α is recovered with high accuracy (MRE around 9% at $T = 250$, declining to 4–5% at $T = 1,000$), confirming the informational content of the empirical characteristic function for identifying heavy-tail behavior. Fourth, in the $\mathcal{G}\alpha\mathcal{S}$ setting, the asymmetry parameter β is the most difficult to estimate (MRE of 69% at $T = 250$), although its identification is not required for the autoregressive and scale parameters. Across all scenarios, the RMSE and MRE decrease consistently with T , confirming the good finite-sample behavior of the estimator.

S.1.2. Case 2: aggregation of mixed causal-noncausal MAR(1,1) processes

We extend the Monte Carlo exercise of Section S.1.1 to Case 2 of Section 2.3 of the main text, in which each latent component is a mixed causal-noncausal MAR(1,1) process, $(1 - \phi_j L)(1 - \psi_j L^{-1})X_{j,t} = \varepsilon_{j,t}$, adding one causal root ϕ_j per component relative to Case 1. No numerical design for this case is specified in Section 4.1 of the main text; to isolate the effect of the added causal dynamics, the design below reuses exactly the Case 1 latent scale/mixing parameters of Table S.1, $(\psi_1, \psi_2, \varsigma_1, \varsigma_2) = (0.8, 0.3, 0.7, 0.9)$, and adds a moderate, well-separated pair of causal roots $(\phi_1, \phi_2) = (0.3, 0.5)$. The same three distributional panels, sample sizes, and 1,000 replications are used. Table S.2 reports bias, RMSE, and MRE for the resulting eight-dimensional parameter vector $\theta = (\psi_1, \phi_1, \varsigma_1, \psi_2, \phi_2, \varsigma_2, \alpha, \beta)$.

Table S.2: Monte Carlo estimation accuracy for mixed causal-noncausal MAR(1,1) α -stable aggregates (1,000 replications)

θ	True	$T = 250$			$T = 500$			$T = 1000$		
		Bias	RMSE	MRE	Bias	RMSE	MRE	Bias	RMSE	MRE
<i>Panel A: Cauchy ($\mathcal{S1S}$, $\alpha = 1$, $\beta = 0$)</i>										
ψ_1	0.800	−0.082	0.212	0.178	−0.049	0.154	0.131	−0.022	0.094	0.077
ϕ_1	0.300	0.067	0.259	0.733	0.068	0.234	0.651	0.065	0.208	0.591
ς_1	0.700	0.042	0.480	0.539	−0.008	0.394	0.449	−0.018	0.300	0.337
ψ_2	0.300	−0.035	0.229	0.667	−0.041	0.207	0.590	−0.058	0.180	0.495
ϕ_2	0.500	−0.056	0.270	0.419	−0.044	0.234	0.351	−0.047	0.203	0.295
ς_2	0.900	0.025	0.509	0.438	0.032	0.399	0.353	0.012	0.305	0.264
<i>Panel B: Symmetric α-stable ($\mathcal{S}\alpha\mathcal{S}$, $\alpha = 1.5$, $\beta = 0$)</i>										
ψ_1	0.800	−0.091	0.191	0.175	−0.054	0.137	0.128	−0.028	0.096	0.090
ϕ_1	0.300	0.095	0.244	0.681	0.090	0.221	0.612	0.092	0.202	0.566
ς_1	0.700	−0.001	0.343	0.397	−0.020	0.299	0.349	−0.031	0.255	0.289
ψ_2	0.300	0.014	0.225	0.630	−0.005	0.197	0.535	−0.021	0.156	0.406
ϕ_2	0.500	−0.092	0.247	0.391	−0.081	0.216	0.329	−0.081	0.187	0.279
ς_2	0.900	−0.086	0.294	0.242	−0.053	0.235	0.200	−0.035	0.199	0.158
α	1.500	0.008	0.190	0.101	0.008	0.142	0.076	0.004	0.105	0.056
<i>Panel C: General α-stable ($\mathcal{G}\alpha\mathcal{S}$, $\alpha = 1.5$, $\beta = 0.3$)</i>										
ψ_1	0.800	−0.073	0.200	0.171	−0.045	0.138	0.123	−0.026	0.100	0.093
ϕ_1	0.300	0.070	0.279	0.788	0.076	0.250	0.705	0.060	0.224	0.635
ς_1	0.700	−0.069	0.355	0.423	−0.048	0.305	0.359	−0.036	0.259	0.297
ψ_2	0.300	−0.076	0.230	0.679	−0.076	0.208	0.600	−0.070	0.185	0.508
ϕ_2	0.500	−0.113	0.289	0.474	−0.100	0.255	0.399	−0.086	0.223	0.335
ς_2	0.900	−0.082	0.271	0.223	−0.074	0.225	0.185	−0.048	0.187	0.157
α	1.500	−0.036	0.189	0.100	−0.020	0.147	0.078	−0.010	0.098	0.053
β	0.300	0.019	0.322	0.832	0.014	0.236	0.601	0.009	0.155	0.400

Notes: True parameter values are $\psi_1 = 0.8$, $\psi_2 = 0.3$, $\phi_1 = 0.3$, $\phi_2 = 0.5$, $\varsigma_1 = 0.7$, $\varsigma_2 = 0.9$ (same latent scale/mixing design as Table S.1, extended with causal roots ϕ_1, ϕ_2). MRE denotes the mean relative error $\mathbb{E}[|\hat{\theta} - \theta_0|/|\theta_0|]$. Estimation minimizes the discretized criterion (2.14) of the main text using the Assumption C2 weight $w(u, v) = (\kappa/\pi) \exp(-\kappa(u^2 + v^2))$ with $\kappa = 1$.

Table S.2 shows that adding a causal root per component raises estimation error throughout, relative

to the matched Case 1 design of Table S.1: the shared coordinates $\psi_1, \psi_2, \varsigma_1, \varsigma_2$ have systematically higher RMSE and MRE here than in Table S.1, reflecting the added difficulty of separating causal from noncausal dynamics within each component from the aggregate characteristic function alone. The new causal root ϕ_1 (true value 0.3) is, across all three panels and sample sizes, the single hardest parameter to pin down (MRE between 57% and 79%), harder even than the asymmetry parameter β in the $\mathcal{G}\alpha\mathcal{S}$ panel; this mirrors the Case 1 finding that a smaller, less persistent root ($\psi_2 = 0.3$ there, $\phi_1 = 0.3$ here) is intrinsically harder to identify than a larger one, since $\phi_2 = 0.5$ is comparatively better estimated (MRE 28–47%). As in Case 1, the tail index α remains accurately recovered (MRE around 10% at $T = 250$, declining to 5–6% at $T = 1,000$), and all reported RMSE and MRE decrease monotonically with T across every panel and parameter, confirming that the estimator remains well behaved once causal dynamics are added.

S.2. Subsampling Diagnostics and Convergence Analysis

S.2.1. Subsampling methodology

This section describes the subsampling procedure used to assess the finite-sample convergence of our minimum distance estimator toward the limiting normal distribution predicted by Proposition 2.1 of the main paper. The methodology follows [Politis and Romano \(1994\)](#) and [Politis et al. \(1999\)](#) and provides a nonparametric way to evaluate the sampling distribution of $\sqrt{n}(\hat{\theta}_n - \theta_0)$ without requiring explicit knowledge of the asymptotic variance matrix $\Sigma^{-1}\Omega\Sigma^{-1}$.

Given a full sample of size n , we construct non-overlapping subsamples of size $b < n$:

$$\mathcal{X}_b^{(i)} = \{\mathcal{X}_{(i-1)b+1}, \dots, \mathcal{X}_{ib}\}, \quad i = 1, \dots, N_b = \lfloor n/b \rfloor. \quad (\text{S.1})$$

Following [Politis et al. \(1999\)](#), we set $b = \lfloor n^{2/3} \rfloor$, which ensures that $b \rightarrow \infty$, $b/n \rightarrow 0$, and satisfies the higher-order requirements for the subsampling approximation to be valid under weak dependence.

Since the individual scale parameters ς_1 and ς_2 exhibit poor finite-sample convergence in the original parameterization, we apply a post-estimation reparameterization:

$$\vartheta = (\rho_1, \rho_2, \sigma, \pi_1, \alpha), \quad \text{where} \quad \sigma = \varsigma_1 + \varsigma_2, \quad \pi_1 = \frac{\varsigma_1}{\sigma}. \quad (\text{S.2})$$

The scaled subsample deviations are then constructed as

$$Z_b^{(i)} = \sqrt{b} \left(\hat{\vartheta}_b^{(i)} - \hat{\vartheta}_n \right), \quad i = 1, \dots, N_b, \quad (\text{S.3})$$

where $\hat{\vartheta}_n$ is the full-sample MDE estimator mapped to the reparameterized space. Under the conditions of Proposition 2.1, and following the theoretical framework of [Politis and Romano \(1994\)](#), these scaled deviations should approximately follow the same asymptotic distribution as $\sqrt{n}(\hat{\vartheta}_n - \vartheta_0)$.

For each Monte Carlo replication $m = 1, \dots, M$ and a given sample size n , we proceed as follows:

- (ι) Generate a time series $\mathcal{X}_t$ from the true model with parameters θ_0 .
- (υ) Compute the full-sample MDE estimator $\hat{\theta}_n$ and map it to the reparameterized space $\hat{\vartheta}_n$.
- ($\iota\iota$) Extract $N_b = \lfloor n/b \rfloor$ non-overlapping subsamples with $b = \lfloor n^{2/3} \rfloor$.
- ($\iota\upsilon$) For each subsample $i = 1, \dots, N_b$, compute the subsample estimator $\hat{\vartheta}_b^{(i)}$.
- (ν) Construct the scaled deviations $Z_b^{(i)}$ according to (S.3).

The pooled non-overlapping block deviations across all Monte Carlo replications are subjected to several normality diagnostics for each parameter component. First, we compute summary statistics including the sample mean, standard deviation, skewness, and excess kurtosis. Second, we report the Shapiro-Wilk (SW) and Jarque-Bera (JB) rejection rates at the 5% significance level. Third, we assess confidence interval coverage using both quantile-based and normal-approximation intervals at the 90% and 95% levels; for this

purpose, confidence intervals are constructed from the $n - b + 1$ overlapping blocks when that count is at or below a target of approximately 2,000 blocks, and from an evenly-spaced thinned subset of approximately 2,000 of them otherwise, to control dependence between adjacent blocks while keeping the calculation tractable at large n , while normality diagnostics use the non-overlapping blocks to avoid dependence contamination.

Table S.3: Monte Carlo Design Parameters

Sample Size n	Subsample Size b	Non-overlapping Blocks N_b	MC Replications M
500	62	8	200
1,000	99	10	200
10,000	464	21	200
100,000	2,154	46	200
500,000	6,299	79	200

Notes: The subsample size is set to $b = \lfloor n^{2/3} \rfloor$ following Politis et al. (1999). $N_b = \lfloor n/b \rfloor$ denotes the number of non-overlapping blocks per replication. The full-sample estimator $\hat{\vartheta}_n$ and every block estimator $\hat{\vartheta}_b^{(i)}$ are the same functional (identical criterion, grid, starting values, bounds and optimizer tolerances), as required for the subsampling approximation of Politis and Romano (1994) to apply.

Table S.4 summarizes the key subsampling results for the estimator in the reparameterized space, reporting results for $n \in \{500, 1,000, 10,000, 100,000, 500,000\}$.

Table S.4 shows the same coverage pattern across all five parameters. At $n = 500$ and $n = 1,000$, both quantile- and normal-approximation intervals under-cover badly for every parameter: quantile-based 90% coverage runs from 0.350 (ψ_2 , $n = 500$) to 0.790 (α , $n = 1,000$), and normal-approximation intervals beat quantile intervals throughout, which fits with how few non-overlapping blocks are available at these sizes ($N_b = 8$ and 10). Coverage climbs steadily with n , and by $n = 100,000$ all five parameters sit within a few points of nominal for both interval types (qCov90 between 0.905 and 0.950, nCov90 between 0.890 and 0.965). π_1 , it is among the worst covered at $n = 500$ –1,000 and reaches near-nominal coverage alongside the rest by $n = 100,000$. At $n = 500,000$ coverage stays at roughly that same level (qCov90 between 0.910 and 0.940), a plateau, not further improvement. What does move at this sample size is Std. Est., which shrinks for every parameter by almost exactly the $1/\sqrt{5}$ factor $\sqrt{n}$ -consistency predicts (e.g. ψ_1 : $0.006 \rightarrow 0.003$; π_1 : $0.012 \rightarrow 0.005$).

S.2.2. Subsampling in the reparameterized space

We now report the subsampling diagnostics in the reparameterized space $\vartheta = (\rho_1, \rho_2, \sigma, \pi_1, \alpha)$, where $\sigma = \varsigma_1 + \varsigma_2$ and $\pi_1 = \varsigma_1 / \sigma$. Estimation is performed in the original $(\varsigma_1, \varsigma_2)$ space; the transformation is applied post-estimation. We compute scaled subsample deviations $\sqrt{b}(\hat{\vartheta}_b^{(i)} - \hat{\vartheta}_n)$ using non-overlapping blocks of size $b = \lfloor n^{2/3} \rfloor$. For each sample size $n \in \{500, 1,000, 10,000, 100,000, 500,000\}$, we generate $M = 200$ Monte Carlo replications and compute $N_b = \lfloor n/b \rfloor$ non-overlapping blocks per replication. Confidence interval coverage is assessed using the $n - b + 1$ overlapping blocks per replication, thinned to approximately 2,000

Table S.4: Summary of subsampling results (reparameterized space)

n	Param.	True	Avg. Est.	Std. Est.	Dev. Mean	Skew	Quantile CI		Normal CI		
							Kurt	qCov90	qCov95	nCov90	nCov95
500	ψ_1	0.80	0.810	0.098	-0.14	-0.59	0.36	0.495	0.525	0.650	0.720
	ψ_2	0.30	0.282	0.207	0.69	0.15	-0.28	0.350	0.350	0.475	0.535
	σ	1.60	1.537	0.218	-0.06	0.32	-0.07	0.565	0.590	0.615	0.730
	π_1	0.4375	0.394	0.181	-0.45	0.28	-0.26	0.360	0.405	0.515	0.590
	α	1.50	1.501	0.121	0.41	-0.09	-0.58	0.655	0.680	0.780	0.840
1,000	ψ_1	0.80	0.803	0.073	-0.05	-0.32	0.00	0.605	0.650	0.785	0.810
	ψ_2	0.30	0.269	0.161	0.92	0.15	-0.51	0.455	0.460	0.620	0.690
	σ	1.60	1.566	0.157	-0.34	0.27	-0.08	0.685	0.720	0.760	0.815
	π_1	0.4375	0.429	0.144	-0.87	0.24	-0.10	0.380	0.450	0.565	0.655
	α	1.50	1.502	0.079	0.34	-0.11	-0.51	0.790	0.835	0.865	0.920
10,000	ψ_1	0.80	0.801	0.019	0.21	0.22	-0.50	0.865	0.900	0.920	0.970
	ψ_2	0.30	0.293	0.044	0.23	0.13	-0.77	0.780	0.795	0.890	0.945
	σ	1.60	1.597	0.051	-1.25	0.11	-0.44	0.830	0.865	0.845	0.920
	π_1	0.4375	0.437	0.039	-1.05	-0.31	-0.41	0.830	0.855	0.905	0.945
	α	1.50	1.501	0.026	0.15	-0.08	-0.10	0.895	0.955	0.900	0.940
100,000	ψ_1	0.80	0.801	0.006	0.11	0.94	2.03	0.940	0.970	0.935	0.990
	ψ_2	0.30	0.301	0.013	-0.55	-0.65	0.64	0.925	0.950	0.955	0.980
	σ	1.60	1.599	0.015	-0.82	-0.24	0.19	0.925	0.960	0.935	0.965
	π_1	0.4375	0.436	0.012	-0.13	-0.79	1.96	0.950	0.995	0.965	0.995
	α	1.50	1.501	0.008	0.07	-0.01	0.02	0.905	0.940	0.890	0.950
500,000	ψ_1	0.80	0.800	0.003	0.03	0.49	1.53	0.925	0.965	0.910	0.950
	ψ_2	0.30	0.299	0.006	-0.30	-0.73	1.63	0.940	0.965	0.955	0.980
	σ	1.60	1.599	0.006	-0.35	-0.06	0.14	0.940	0.960	0.930	0.955
	π_1	0.4375	0.437	0.005	0.01	-0.33	1.44	0.925	0.970	0.930	0.965
	α	1.50	1.500	0.004	0.04	0.01	-0.01	0.910	0.950	0.900	0.945

Notes: Avg. Est. and Std. Est. are computed over $M = 200$ replications. Dev. Mean, Skew, and Kurt refer to the mean, skewness, and excess kurtosis of the non-overlapping block scaled deviations $Z_b^{(i)}$. qCov and nCov denote the average coverage of quantile-based and normal-approximation confidence intervals. Complete results are in Tables S.5–S.9 below.

evenly-spaced blocks at $n = 10,000$ and above (see the note to Table S.5 below). Full diagnostics are reported for these five sample sizes, spanning the small-, medium-, large-, and very-large-sample regimes; the design is documented in Table S.3 and the convergence pattern is summarized in Table S.4 above. The descriptive statistics in Tables S.5–S.9 report the mean, standard deviation, skewness, and excess kurtosis of the non-overlapping block deviations, alongside the Shapiro-Wilk (SW) and Jarque-Bera (JB) rejection rates at the 5% significance level, and the average coverage of 90% and 95% confidence intervals (both quantile-based and normal-approximation).

Implementation details for the subsampling estimator, the numerical optimizer tolerance, the requirement that the full-sample and block estimators be the same functional per the subsampling theory of Politis and Romano (1994), and the ordering enforced on (ρ_1, ρ_2) to prevent label switching, are documented in the replication code. One feature of the results is worth commenting on here. A small share of blocks reach the boundary of the compact parameter space Θ at $n = 100,000$ (confirmed genuine by dense multi-start re-optimization, not a search failure); because the scaled deviations $Z_b^{(i)}$ grow with $\sqrt{b}$, this inflates the raw

Skew and Kurt columns at that sample size even as the share of affected blocks falls, with no bearing on the asymptotic normality of $\hat{\vartheta}_n$ itself, established on the full sample by Proposition 2.1. It shows up clearly in the SW/JB rejection rates: they rise sharply at $n = 100,000$ (Table S.8) for the three coordinates confined to a bounded subset of Θ (SW between 0.475 and 0.625, against 0.105–0.320 at $n = 500$ –10,000), while σ and α , unconstrained in this design, show no such reversal (SW between 0.040 and 0.085 throughout). At $n = 500,000$ (Table S.9) this recedes for most of these coordinates rather than worsening, e.g. SW rejection for ρ_1 falls from 0.595 to 0.260, with one exception (the JB rate for ρ_2 ticks up slightly instead of receding). This is consistent with a transient finite- b boundary artifact rather than a breakdown of the underlying asymptotic theory: a genuinely growing problem would predict continued deterioration at the largest sample size, which is not what most of the affected rates show.

Table S.5: Subsampling results (reparameterized) for $n = 500$ ($b = 62$, $N_b = 8$, $M = 200$)

Parameter	True	Avg Est.	Std Est.	Dev. Mean	Dev. Std	Skew	Kurt	SW Rej.	Quantile CI		Normal CI		
									JB Rej.	qCov90	qCov95	nCov90	nCov95
ρ_1	0.800	0.810	0.098	−0.14	1.430	−0.59	0.36	0.235	0.000	0.495	0.525	0.650	0.720
ρ_2	0.300	0.282	0.207	0.69	2.506	0.15	−0.28	0.280	0.000	0.350	0.350	0.475	0.535
σ	1.600	1.537	0.218	−0.06	3.242	0.32	−0.07	0.060	0.005	0.565	0.590	0.615	0.730
π_1	0.437	0.394	0.181	−0.45	2.080	0.28	−0.26	0.105	0.005	0.360	0.405	0.515	0.590
α	1.500	1.501	0.121	0.41	2.068	−0.09	−0.58	0.065	0.000	0.655	0.680	0.780	0.840

Notes: All results are in the reparameterized space $(\rho_1, \rho_2, \sigma, \pi_1, \alpha)$. The full-sample estimator $\hat{\vartheta}_n$ and every block estimator $\hat{\vartheta}_b^{(i)}$ are the same functional: identical criterion, grid, starting values, bounds and optimizer tolerances, as required by the subsampling theory of Politis and Romano (1994). Dev. Mean, Dev. Std, Skew and Kurt are computed on the non-overlapping block deviations $Z_b^{(i)}$ pooled across all $M = 200$ replications (1,600 values per parameter). At large n , a small share of blocks reach the boundary of the compact parameter space Θ (a genuine finite- b feature of subsampling, not a search failure: confirmed by dense multi-start re-optimization) and, because $Z_b^{(i)}$ scales with $\sqrt{b}$, a fixed absolute distance to the boundary maps into an increasingly extreme standardized value as b grows, even as the share of affected blocks falls, this inflates raw Skew and Kurt without bearing on the asymptotic normality of $\hat{\vartheta}_n$ itself, established on the full sample by Proposition 2.1. SW and JB rejection rates are per-replication rejection frequencies at the 5% level. qCov and nCov denote quantile-based and normal-approximation CI coverage; confidence intervals use all $n - b + 1$ overlapping blocks when that count is at or below a target of approximately 2,000 (439 blocks at $n = 500$; 902 at $n = 1,000$), and an evenly-spaced thinned subset of approximately 2,000 of them otherwise (1,908 of 9,537 at $n = 10,000$; 1,997 of 97,847 at $n = 100,000$; 1,999 of 493,702 at $n = 500,000$), while the normality diagnostics use the non-overlapping blocks throughout.

Table S.6: Subsampling results (reparameterized) for $n = 1,000$ ($b = 99$, $N_b = 10$, $M = 200$)

Parameter	True	Avg Est.	Std Est.	Dev. Mean	Dev. Std	Skew	Kurt	SW Rej.	Quantile CI		Normal CI		
									JB Rej.	qCov90	qCov95	nCov90	nCov95
ρ_1	0.800	0.803	0.073	−0.05	1.552	−0.32	0.00	0.190	0.000	0.605	0.650	0.785	0.810
ρ_2	0.300	0.269	0.161	0.92	2.849	0.15	−0.51	0.320	0.000	0.455	0.460	0.620	0.690
σ	1.600	1.566	0.157	−0.34	3.828	0.27	−0.08	0.075	0.000	0.685	0.720	0.760	0.815
π_1	0.437	0.429	0.144	−0.87	2.352	0.24	−0.10	0.130	0.010	0.380	0.450	0.565	0.655
α	1.500	1.502	0.079	0.34	2.347	−0.11	−0.51	0.040	0.005	0.790	0.835	0.865	0.920

Notes: See notes to Table S.5.

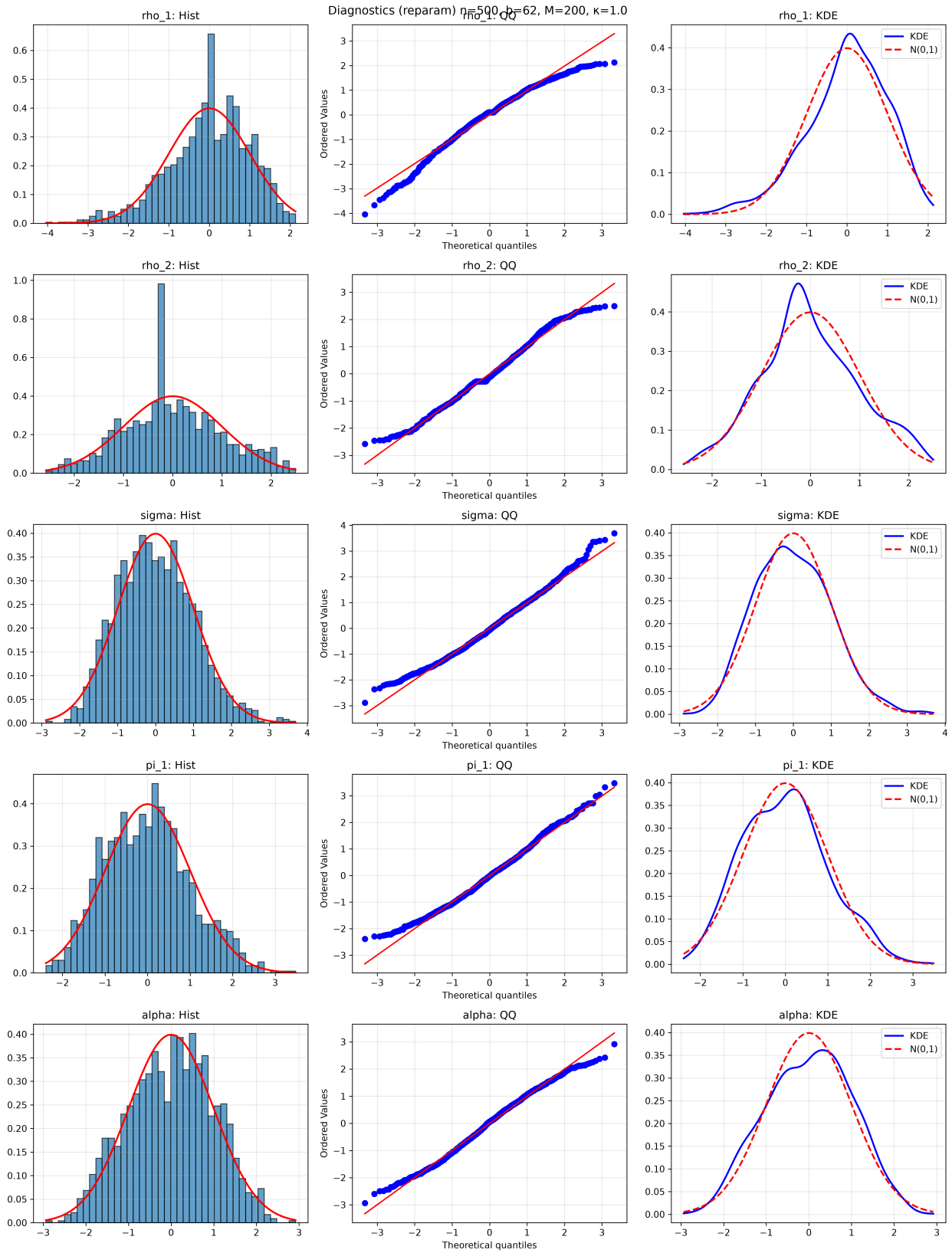


Figure S.1: Diagnostic plots for scaled subsample deviations (reparameterized) at $n = 500$. For each parameter (rows), we display: (left) histogram with standard normal overlay, (center) Q-Q plot against standard normal quantiles, and (right) kernel density estimate compared to the standard normal density.

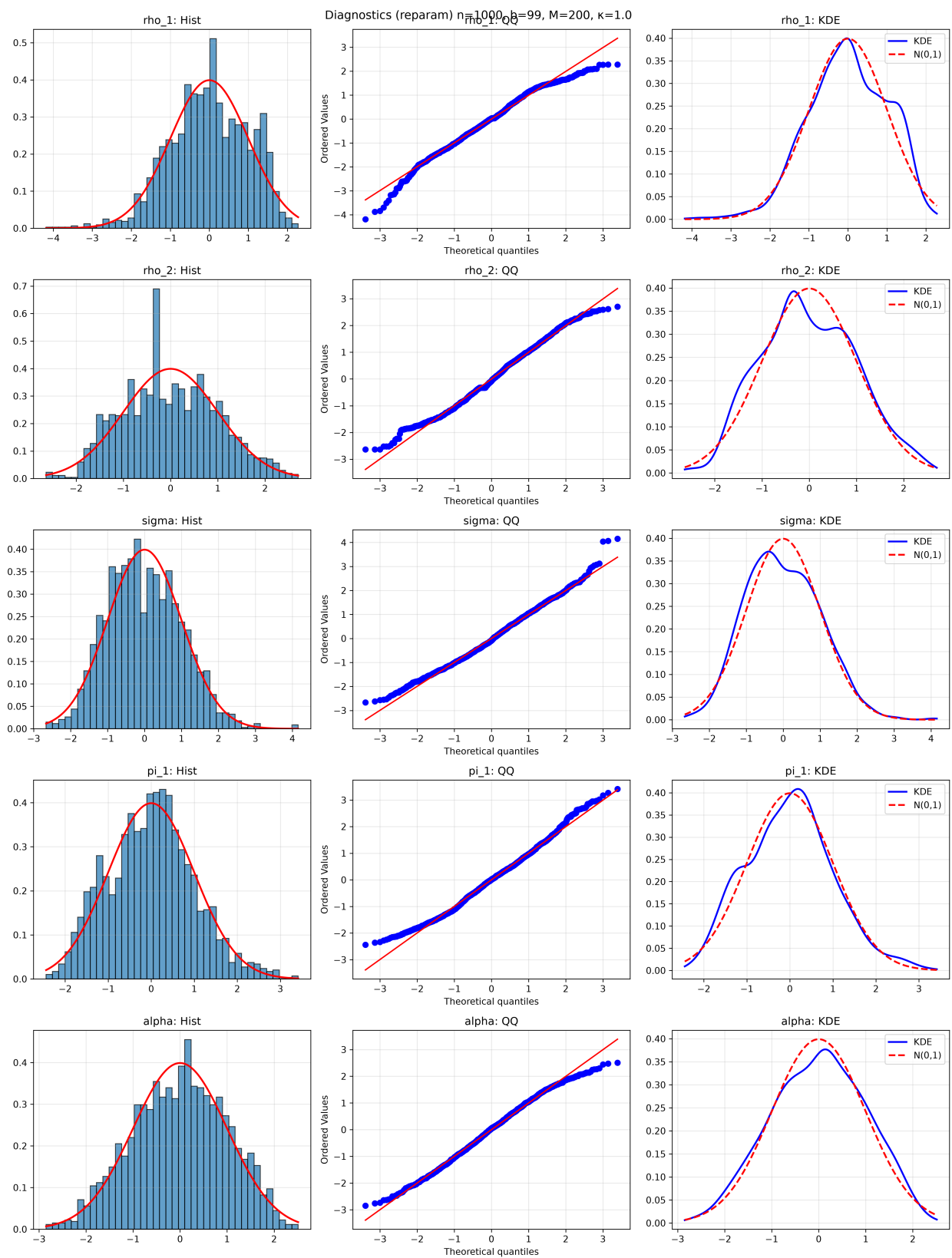


Figure S.2: Diagnostic plots for scaled subsample deviations (reparameterized) at $n = 1,000$. See caption to Figure S.1 for details.

Table S.7: Subsampling results (reparameterized) for $n = 10,000$ ($b = 464$, $N_b = 21$, $M = 200$)

Parameter	True	Avg Est.	Std Est.	Dev. Mean	Dev. Std	Skew	Kurt	SW Rej.	Quantile CI		Normal CI		
									JB Rej.	qCov90	qCov95	nCov90	nCov95
ρ_1	0.800	0.801	0.019	0.21	2.186	0.22	−0.50	0.240	0.005	0.865	0.900	0.920	0.970
ρ_2	0.300	0.293	0.044	0.23	4.403	0.13	−0.77	0.290	0.000	0.780	0.795	0.890	0.945
σ	1.600	1.597	0.051	−1.25	4.718	0.11	−0.44	0.065	0.005	0.830	0.865	0.845	0.920
π_1	0.437	0.437	0.039	−1.05	3.895	−0.31	−0.41	0.230	0.020	0.830	0.855	0.905	0.945
α	1.500	1.501	0.026	0.15	2.593	−0.08	−0.10	0.055	0.030	0.895	0.955	0.900	0.940

Notes: See notes to Table S.5.

Table S.8: Subsampling results (reparameterized) for $n = 100,000$ ($b = 2,154$, $N_b = 46$, $M = 200$)

Parameter	True	Avg Est.	Std Est.	Dev. Mean	Dev. Std	Skew	Kurt	SW Rej.	Quantile CI		Normal CI		
									JB Rej.	qCov90	qCov95	nCov90	nCov95
ρ_1	0.800	0.801	0.006	0.11	2.380	0.94	2.03	0.595	0.590	0.940	0.970	0.935	0.990
ρ_2	0.300	0.301	0.013	−0.55	5.146	−0.65	0.64	0.625	0.330	0.925	0.950	0.955	0.980
σ	1.600	1.599	0.015	−0.82	5.115	−0.24	0.19	0.085	0.095	0.925	0.960	0.935	0.965
π_1	0.437	0.436	0.012	−0.13	4.826	−0.79	1.96	0.475	0.545	0.950	0.995	0.965	0.995
α	1.500	1.501	0.008	0.07	2.692	−0.01	0.02	0.050	0.025	0.905	0.940	0.890	0.950

Notes: See notes to Table S.5.

Table S.9: Subsampling results (reparameterized) for $n = 500,000$ ($b = 6,299$, $N_b = 79$, $M = 200$)

Parameter	True	Avg Est.	Std Est.	Dev. Mean	Dev. Std	Skew	Kurt	SW Rej.	Quantile CI		Normal CI		
									JB Rej.	qCov90	qCov95	nCov90	nCov95
ρ_1	0.800	0.800	0.003	0.03	2.099	0.49	1.53	0.260	0.275	0.925	0.965	0.910	0.950
ρ_2	0.300	0.299	0.006	−0.30	4.779	−0.73	1.63	0.460	0.420	0.940	0.965	0.955	0.980
σ	1.600	1.599	0.006	−0.35	5.045	−0.06	0.14	0.090	0.065	0.940	0.960	0.930	0.955
π_1	0.437	0.437	0.005	0.01	4.251	−0.33	1.44	0.170	0.195	0.925	0.970	0.930	0.965
α	1.500	1.500	0.004	0.04	2.763	0.01	−0.01	0.070	0.045	0.910	0.950	0.900	0.945

Notes: See notes to Table S.5. Confidence intervals use a thinned subset of 1,999 of the 493,702 overlapping blocks (target of approximately 2,000, evenly spaced).

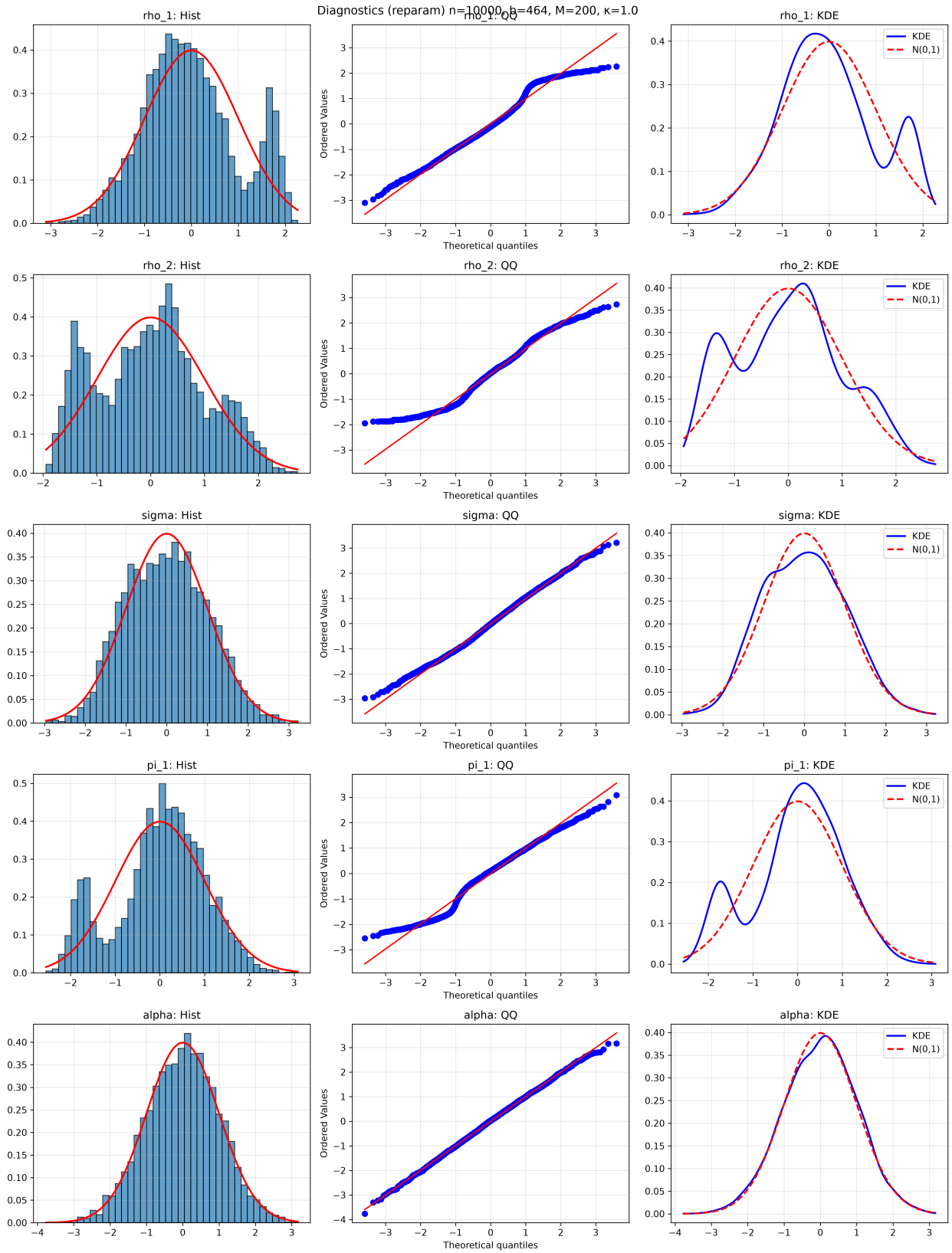


Figure S.3: Diagnostic plots for scaled subsample deviations (reparameterized) at $n = 10,000$. See caption to Figure S.1 for details.

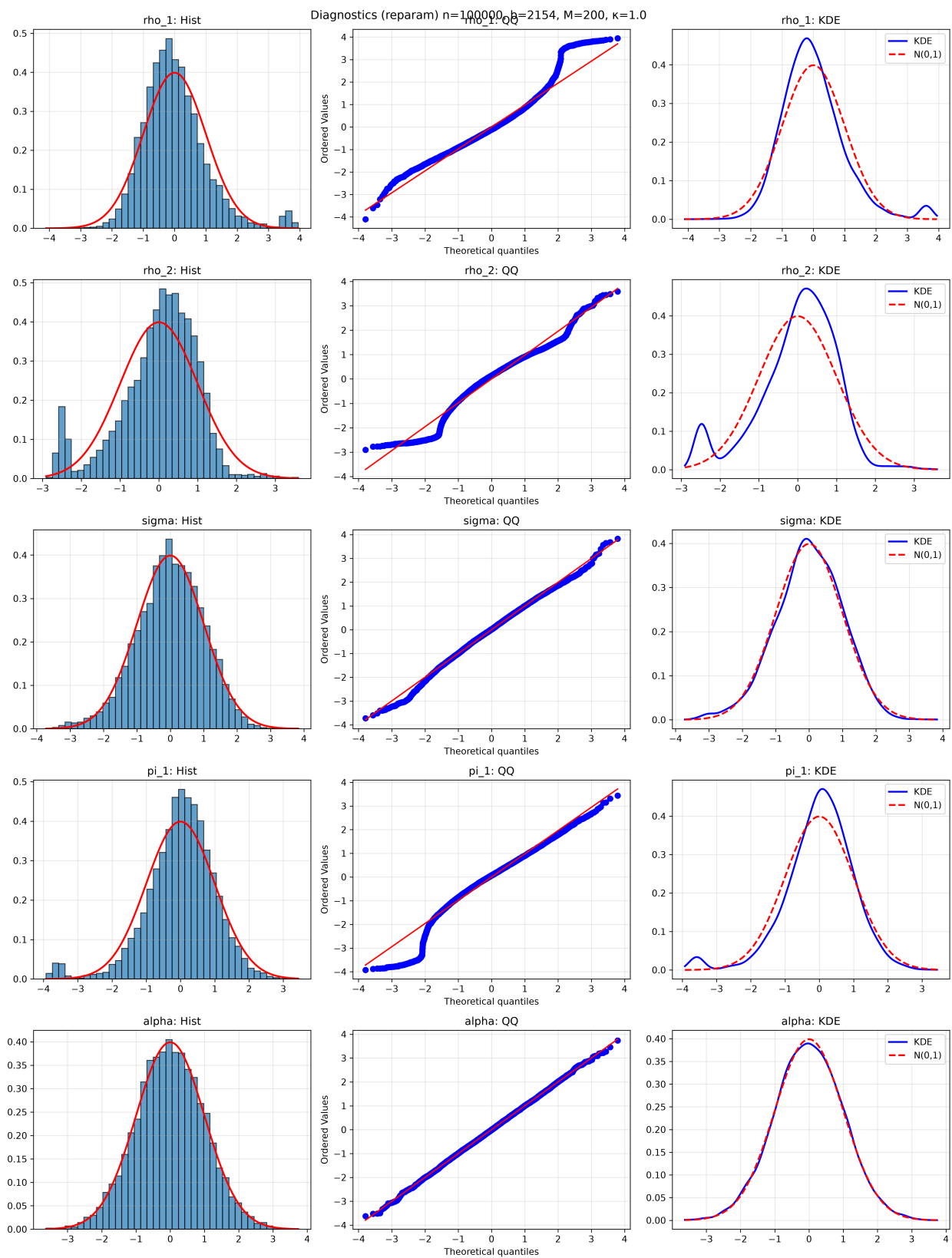


Figure S.4: Diagnostic plots for scaled subsample deviations (reparameterized) at $n = 100,000$. See caption to Figure S.1 for details.

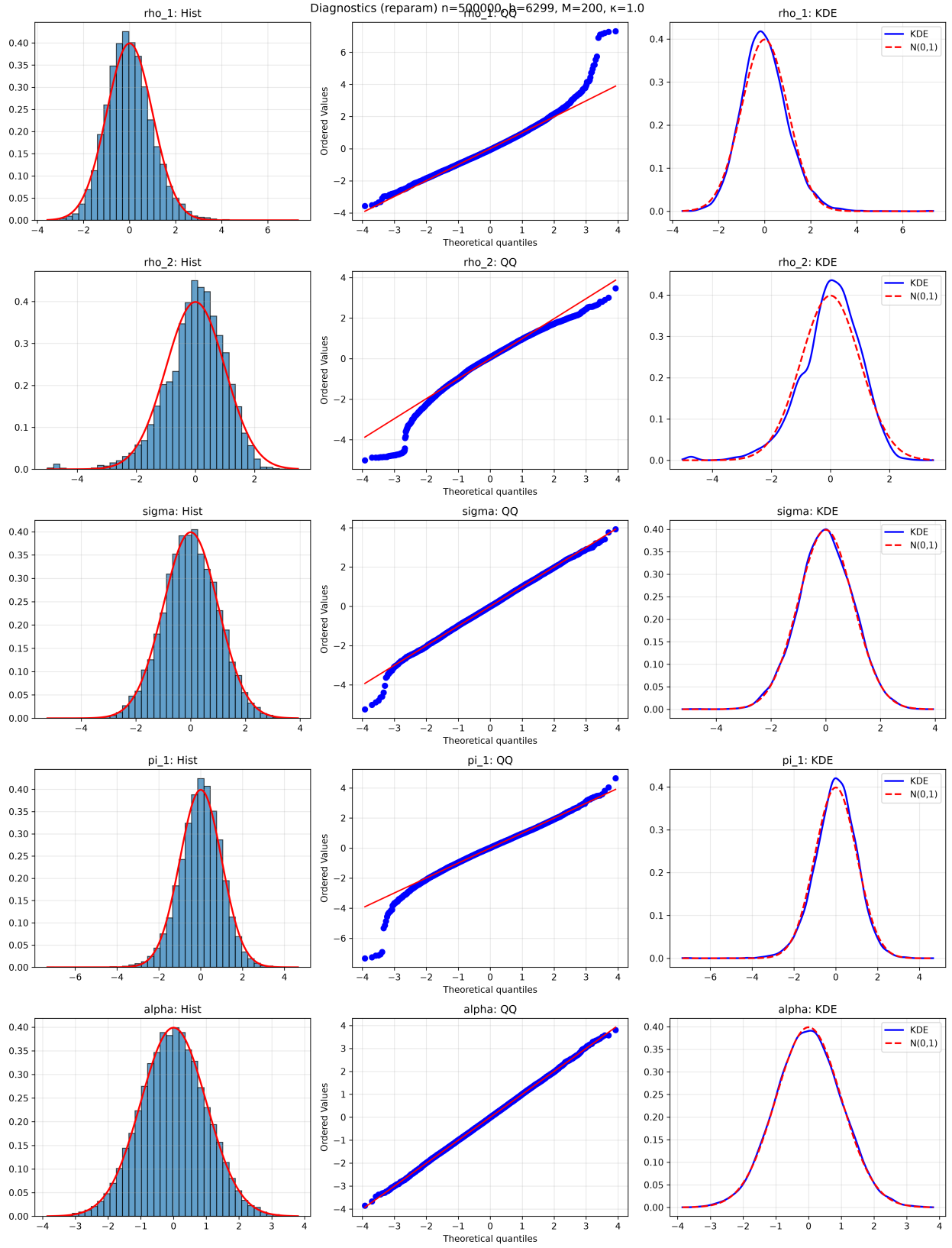


Figure S.5: Diagnostic plots for scaled subsample deviations (reparameterized) at $n = 500,000$. See caption to Figure S.1 for details.

S.3. Illustration on the OVX Index

This section provides the empirical analysis of the CBOE Crude Oil ETF Volatility Index (OVX) that supports the summary presented in Section 5 of the main paper. The analysis proceeds in four stages. Section S.3.1 reports the estimation results under three distributional specifications and discusses the identification of the latent components. Section S.3.2 presents the deconvolution of the observed series into its two constituent components using the filtering methodology of [de Truchis et al. \(2026b\)](#), revealing the heterogeneous dynamics that characterize oil market volatility. Section S.3.3 develops the in-sample forecasting exercise for the 2020 oil market disruption, detailing the pattern-matching algorithm, the crash probability computations, and the role of risk thresholds in generating predictions. Section S.3.4 collects the per-component graphical diagnostics, including crash probability profiles and forecast trajectories across multiple risk thresholds.

S.3.1. Estimation

This subsection reports the minimum distance estimation results for the OVX index under three distributional specifications: general α -stable ($\mathcal{G}\alpha\mathcal{S}$), symmetric α -stable ($\mathcal{S}\alpha\mathcal{S}$), and Cauchy ($\mathcal{S}1\mathcal{S}$). The observed series is the CBOE Crude Oil ETF Volatility Index, collected from the FRED website at a weekly frequency over the period May 23, 2015–April 27, 2025 ($T = 518$), and linearly detrended to avoid high-frequency noise contamination. Table S.10 presents the parameter estimates, asymptotic standard errors computed using the subsampling methodology of Section S.2.1, and t -statistics for each specification.

Table S.10: Estimation results for the OVX index under three specifications

$\hat{\theta}$	$\mathcal{G}\alpha\mathcal{S}$			$\mathcal{S}\alpha\mathcal{S}$			$\mathcal{S}1\mathcal{S}$		
	Estimate	Std.	t -stat	Estimate	Std.	t -stat	Estimate	Std.	t -stat
ψ_1	0.7989	0.0673	11.862	0.2507	0.0077	32.477	0.9226	0.0082	112.824
ψ_2	0.8470	0.0668	12.678	0.9865	0.0040	244.560	0.9346	0.0074	126.404
α	1.4686	0.0995	14.764	1.2405	0.0084	147.613	—	—	—
β	−0.1275	0.0684	−1.863	—	—	—	—	—	—
σ	2.0932	0.2400	8.723	0.8964	0.0212	42.226	0.1966	0.0294	6.692
π_1	0.2790	0.0403	6.930	0.8915	0.0052	171.084	0.5029	0.0511	9.833
π_2	0.7210	0.0409	17.622	0.1085	0.0182	5.957	0.4971	0.0545	9.121

Notes: Asymptotic standard errors are computed using the subsampling methodology described in Section S.2.1. The t -statistics test the null hypothesis that each parameter equals zero. For the $\mathcal{S}1\mathcal{S}$ specification, $\alpha = 1$ is imposed.

The $\mathcal{G}\alpha\mathcal{S}$ specification provides strong evidence of anticipative dynamics: both AR coefficients are highly significant ($\hat{\psi}_1 = 0.80$, $\hat{\psi}_2 = 0.85$). The first component captures more abrupt volatility bursts (lower weight $\hat{\pi}_1 = 0.28$), while the second drives more persistent explosive episodes ($\hat{\pi}_2 = 0.72$). The estimated tail index $\hat{\alpha} = 1.47$ confirms the presence of heavy tails well beyond Gaussian accommodation. The asymmetry

parameter $\hat{\beta} = -0.13$ is not significant at the 5% level, yet numerically distorts the parameter structure of the $\mathcal{S}\alpha\mathcal{S}$ model considerably ($\hat{\psi}_1 = 0.25$, $\hat{\psi}_2 = 0.99$). The $\mathcal{S}1\mathcal{S}$ Cauchy restriction ($\alpha = 1$) appears overly binding given the estimated $\hat{\alpha}$ values in the other specifications.

S.3.2. Filtration and deconvolution

This subsection presents the deconvolution of the OVX index into its two latent components under the $\mathcal{G}\alpha\mathcal{S}$ specification. The filtering is performed using the dual MCMC procedure of [de Truchis et al. \(2026b\)](#), which recovers the unobserved component trajectories from the aggregate signal. Figure S.6 displays the observed detrended series alongside the filtered components.

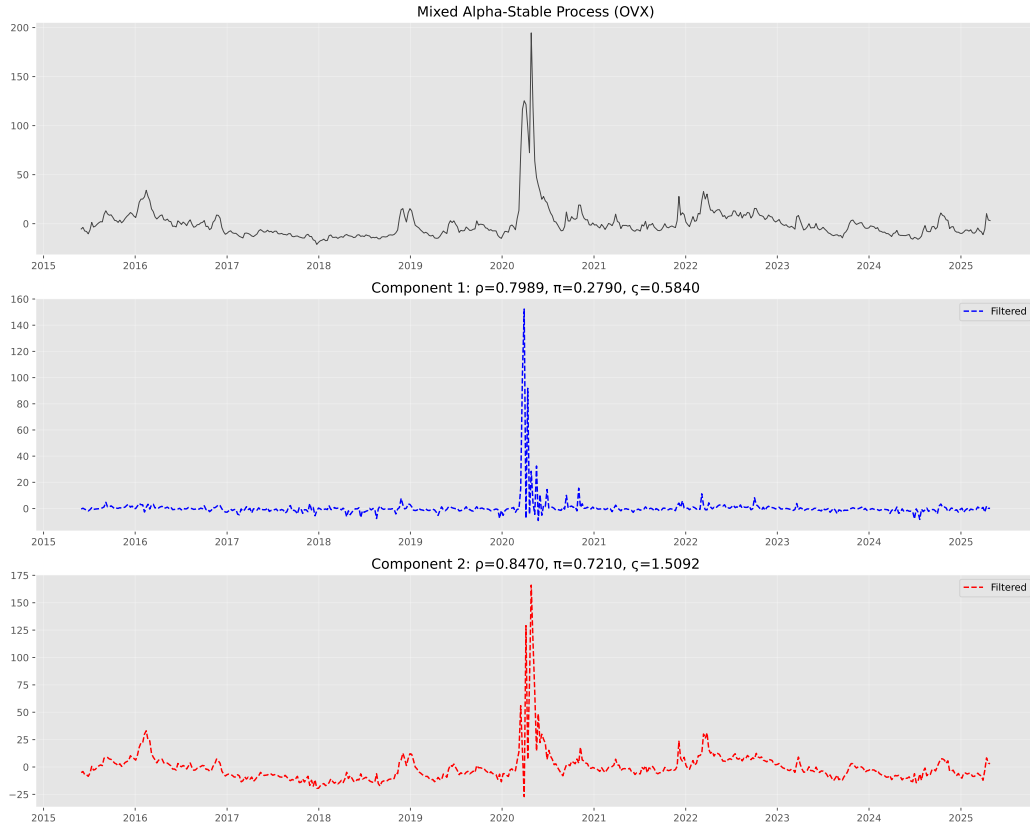


Figure S.6: Deconvolution of the OVX index from the $\mathcal{G}\alpha\mathcal{S}$ two-component model, filtered using [de Truchis et al. \(2026b\)](#). Top panel: observed detrended OVX series. Middle panel: filtered first component ($\hat{\psi}_1 = 0.7989$, $\hat{\pi}_1 = 0.2790$). Bottom panel: filtered second component ($\hat{\psi}_2 = 0.8470$, $\hat{\pi}_2 = 0.7210$).

The coefficient estimates of Section S.3.1 already point to these same qualitative differences between the two components; the filtering step lets us observe each component's own trajectory and characteristics directly, and the same filtered components feed the forecasting algorithm of Section S.3.3. The first component captures abrupt, short-lived volatility bursts: its trajectory shows sharp spikes that decay rapidly, consistent with its lower persistence ($\hat{\psi}_1 = 0.80$) and smaller contribution to the aggregate ($\hat{\pi}_1 = 0.28$). The

second component tracks more sustained explosive patterns: its evolution is smoother but more persistent, reflecting its higher autoregressive coefficient ($\hat{\psi}_2 = 0.85$) and dominant weight ($\hat{\pi}_2 = 0.72$).

Periods of extreme oil market stress, most visibly the 2020 disruption triggered by the COVID-19 pandemic and the Russia-Saudi Arabia oil price war, feature a superposition of both dynamics. During this episode, both components contribute substantially to the observed volatility spike, but with different temporal profiles: the first component captures the initial shock and rapid escalation, while the second sustains the elevated volatility over a longer period. This heterogeneity in persistence properties is precisely what motivates our aggregate modeling framework and enables more nuanced forecasting relative to single-component approaches.

S.3.3. Forecasting

This subsection develops the in-sample forecasting exercise for the 2020 oil market disruption. We first describe the pattern-matching algorithm that underlies our prediction strategy, then present the complete forecast table with crash and survival probabilities for both components, and finally discuss the role of risk thresholds in generating predictions.

S.3.3.1. Algorithm

Our prediction strategy rests on the theoretical result of Proposition 3.3: during an extreme event, the trajectory of the aggregate process concentrates around a specific normalized pattern $\vartheta \mathbf{d}_{j_0,k} / \|\mathbf{d}_{j_0,k}\|$. These pattern structures are defined by four essential elements: the shape derived from the coefficient sequence $\mathbf{d}_{j,k}$; the component index $j \in \{1, \dots, J\}$ identifying which latent process is driving the event; the time shift $k_0 \in \mathbb{Z}$ indicating the position within the pattern; and the sign $\vartheta \in \{-1, +1\}$ reflecting upward or downward movements. In the non-aggregated setting ($J = 1$), this same pattern-matching approach applies directly to a single component, which is exactly what running the algorithm below component by component amounts to; a complete, self-contained description of the algorithm in that single-component setting is provided in the online supplement of [de Truchis et al. \(2026a\)](#). Focusing on the aggregate case, Algorithm 1 summarizes the procedure.

Algorithm 1 Pattern-matching forecast of an emerging extreme event

Require: Observed path of length $m+1$ ($m = 20$ throughout); estimated $\mathcal{G}\alpha\mathcal{S}$ parameters $(\psi_j, \pi_j, \alpha)_{j=1,\dots,J}$; risk threshold p

Ensure: Predicted crash horizon h^* and forecast trajectory $(\hat{X}_{t_0+1}, \dots, \hat{X}_{t_0+h^*})$

```
1: Observation: record the initial  $m + 1$  observations of the emerging event
2: Pattern matching: match the observed path to the theoretical patterns of the estimated model, identifying the component index  $j_0$ , the time shift  $k_0$ , and the sign  $\vartheta_0$ 
3: for  $h = 1, 2, \dots$  do
4:    $\mathbb{P}(\text{crash by } h) \leftarrow 1 - |\psi_{j_0}|^{\alpha h}$  ▷ Proposition 3.3(a)
5:    $\mathbb{P}(\text{survive beyond } h) \leftarrow |\psi_{j_0}|^{\alpha h}$ 
6:    $\hat{X}_{t_0+h} \leftarrow \hat{\psi}_{j_0}^{-(h-k_0)} \cdot X_{t_0}$  ▷ forecast trajectory
7:   if  $\mathbb{P}(\text{crash by } h) > p$  then
8:      $h^* \leftarrow h$ ; break ▷ crash predicted at horizon  $h^*$ 
9:   end if
10: end for
```

For the 2020 episode, the cut-off date is January 2020. Pattern matching yields $k_0 = 3$ for Component 1 ($\hat{\psi}_1 = 0.7989$) and $k_0 = 1$ for Component 2 ($\hat{\psi}_2 = 0.8470$).

S.3.3.2. Prediction table

Table S.11 reports the predicted trajectory values, crash probabilities, and survival probabilities at each weekly horizon for both components, from the January 2020 cut-off through May 2020.

For Component 1 ($k_0 = 3$), crash probabilities become relevant only at $h = 3$; the bubble trajectory escalates from 11.27 to 208.73 over 13 periods before the 99% crash threshold is crossed at $h = 14$. For Component 2 ($k_0 = 1$), crash probabilities are immediately non-negligible (0.22 at $h = 1$), yet the lower growth rate sustains the trajectory through 18 periods, reaching 223.9 before the 99% threshold is crossed at $h = 19$.

S.3.3.3. Risk thresholds

The risk threshold parameter allows practitioners to customize predictions according to their risk tolerance. When the cumulative crash probability $1 - |\psi_{j_0}|^{\alpha h}$ exceeds the threshold, the model predicts a crash; otherwise, it anticipates continued growth. Higher thresholds generate more extreme bubble projections before predicting a crash, while lower thresholds produce more conservative forecasts with earlier predicted crash points. For Component 1, the predicted crash horizon shifts from $h = 7$ at the 90% threshold to $h = 14$ at the 99% threshold. For Component 2, the shift is from $h = 10$ to $h = 19$. In both cases, the 99% forecast trajectory most closely tracks the realized 2020 OVX path.

S.3.4. Per-component graphical diagnostics

This subsection collects the per-component crash probability profiles and forecast trajectories across the three risk thresholds (90%, 95%, 99%). Figures S.7 and S.8 display, for each component separately, the complete diagnostic output. Each figure is organized in three rows: the top row shows the crash/survival/cumulative

Table S.11: In-sample bubble forecast: predicted values and crash/survival probabilities (cut-off: January 2020, risk threshold: 99%)

Date	h	Component 1 ($k_0 = 3, \hat{\psi}_1 = 0.7989$)			Component 2 ($k_0 = 1, \hat{\psi}_2 = 0.8470$)		
		Forecast	Crash at h	Survive at h	Forecast	Crash at h	Survive at h
2020-01-05	0	11.2714	—	—	11.2714	—	—
2020-01-12	1	14.1086	—	—	13.3074	0.2164	0.7836
2020-01-19	2	17.6601	—	—	15.7113	0.3860	0.6140
2020-01-26	3	22.1055	0.6281	0.3719	18.5493	0.5189	0.4811
2020-02-02	4	27.6699	0.7326	0.2674	21.9000	0.6230	0.3770
2020-02-09	5	34.6350	0.8077	0.1923	25.8560	0.7046	0.2954
2020-02-16	6	43.3534	0.8617	0.1383	30.5265	0.7685	0.2315
2020-02-23	7	54.2664	0.9006	0.0994	36.0408	0.8186	0.1814
2020-03-01	8	67.9264	0.9285	0.0715	42.5511	0.8579	0.1421
2020-03-08	9	85.0249	0.9486	0.0514	50.2374	0.8886	0.1114
2020-03-15	10	106.4274	0.9630	0.0370	59.3122	0.9127	0.0873
2020-03-22	11	133.2175	0.9734	0.0266	70.0262	0.9316	0.0684
2020-03-29	12	166.7511	0.9809	0.0191	82.6755	0.9464	0.0536
2020-04-05	13	208.7259	0.9862	0.0138	97.6098	0.9580	0.0420
2020-04-12	14	0.0000*	0.9901	0.0099	115.2418	0.9671	0.0329
2020-04-19	15	0.0000*	0.9929	0.0071	136.0588	0.9742	0.0258
2020-04-26	16	0.0000*	0.9949	0.0051	160.6361	0.9798	0.0202
2020-05-03	17	0.0000*	0.9963	0.0037	189.6530	0.9842	0.0158
2020-05-10	18	0.0000*	0.9974	0.0026	223.9115	0.9876	0.0124
2020-05-17	19	0.0000*	0.9981	0.0019	0.0000*	0.9903	0.0097

Notes: Crash at h denotes $1 - |\psi_{j_0}|^{\alpha h}$; Survive at h denotes $|\psi_{j_0}|^{\alpha h}$. *Crash predicted before this horizon at the 99% threshold (Component 1: $h = 14$; Component 2: $h = 19$). Forecast values after the predicted crash are set to zero. Pattern matching window: $m = 20$.

probability curves as a function of horizon h ; the middle row shows zoomed forecast trajectories; the bottom row situates the forecast within the full 2015–2025 sample.

The first component ($\hat{\psi}_1 = 0.7989$), with its slightly lower persistence, captures more abrupt movements: its contribution to the 2020 crisis is sharp but relatively brief, and the 99% threshold trajectory escalates from 11.27 to 208.73 before the predicted crash. The second component ($\hat{\psi}_2 = 0.8470$), with its higher persistence and dominant weight, tracks the more sustained explosive patterns: its trajectory reaches higher absolute values (up to 223.9) before the predicted collapse, reflecting the longer-lived nature of the underlying market anxiety.

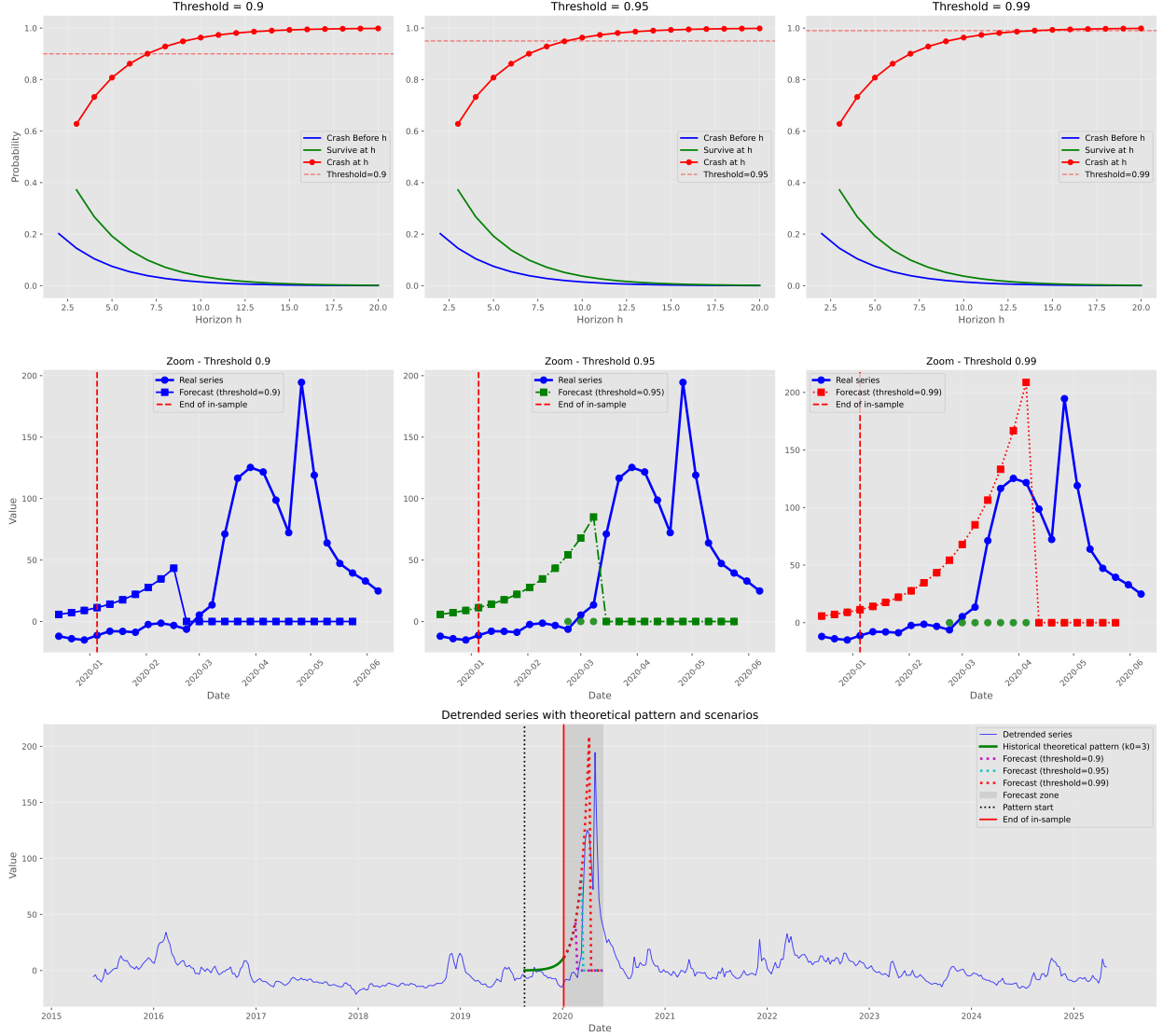


Figure S.7: In-sample forecast for the 2020 OVX bubble: Component 1 ($\hat{\psi}_1 = 0.7989$, $k_0 = 3$). *Top row*: crash probability profiles for risk thresholds 0.90 (left), 0.95 (center), and 0.99 (right). Each panel shows the probability of crashing at a given date (red line with circles), surviving beyond that date (green line), the cumulative crash probability (blue line), and the threshold (horizontal dashed red line). *Middle row*: zoomed-in forecast trajectories for each threshold; colored squares are predicted values; vertical dashed red line marks the cut-off (January 2020). *Bottom row*: full-sample view with the matched historical theoretical pattern ($k_0 = 3$) in green; forecasts at 0.90 (yellow), 0.95 (green dashed), and 0.99 (red dashed) thresholds; shaded area is the forecast zone. Pattern matching window: $m = 20$.

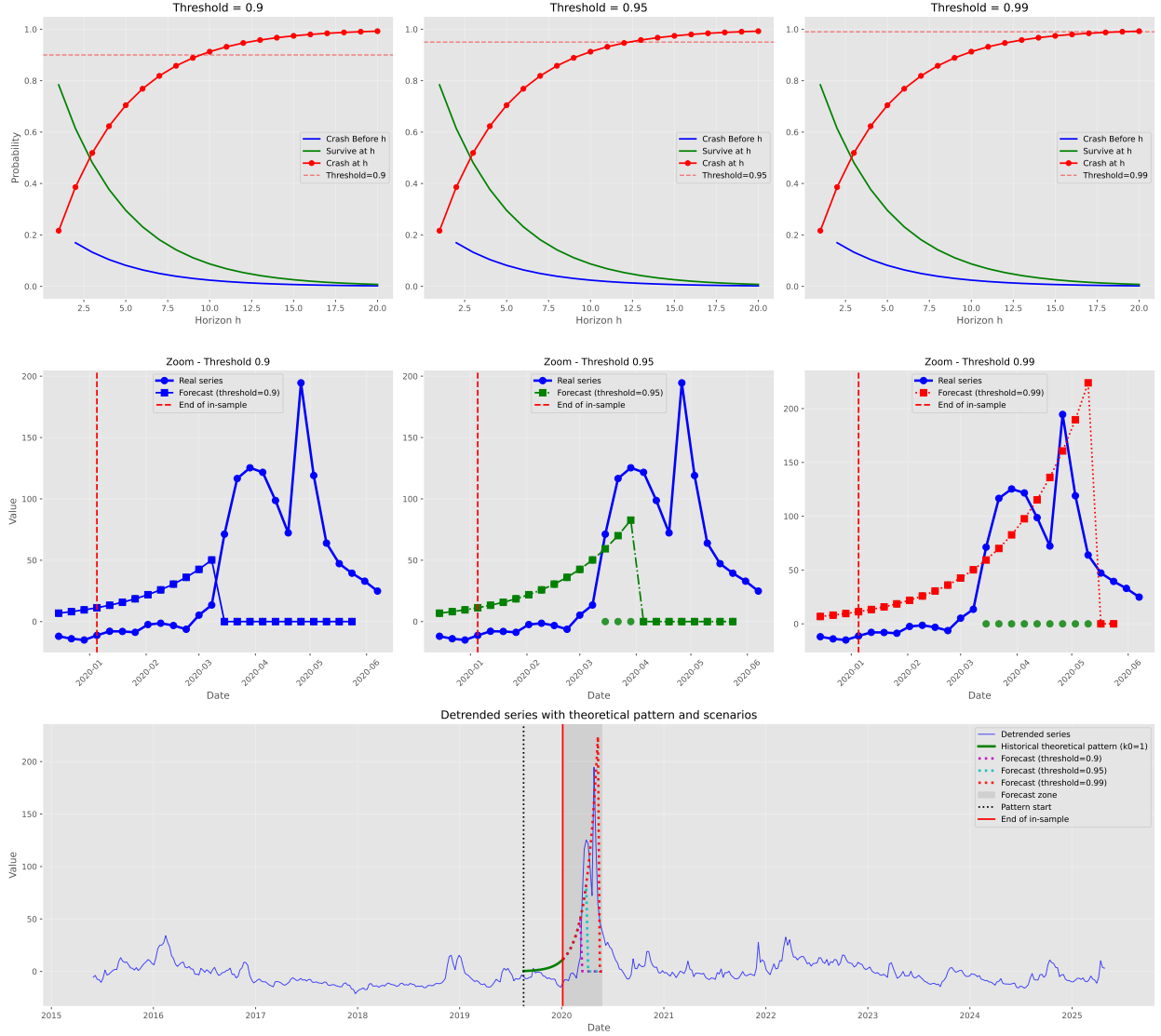


Figure S.8: In-sample forecast for the 2020 OVX bubble: Component 2 ($\hat{\psi}_2 = 0.8470$, $k_0 = 1$). Layout identical to Figure S.7. The historical theoretical pattern ($k_0 = 1$) is shown in green in the bottom row. Compared to Component 1, crash probabilities build more gradually but the forecasted trajectory reaches higher absolute values before the predicted crash, reflecting the higher persistence and dominant weight ($\hat{\pi}_2 = 0.72$) of this component. Pattern matching window: $m = 20$.

References

- Politis, D.N., and J.P., Romano. 1994. Large Sample Confidence Regions Based on Subsamples under Minimal Assumptions. *The Annals of Statistics*, 22(4), 2031–2050.
- Politis, D.N., Romano, J.P., and M., Wolf. 1999. *Subsampling*. Springer.
- de Truchis, G., Fries, S., Thomas, A. (2026). Forecasting Extreme Trajectories Using Seminorm Representations *Working Paper 2025-06*, Chaire Économie du Climat, Paris.
- de Truchis, G., Thomas, A., L. Vaudrée (2026). Deconvolution and Filtering of Non-Causal Alpha-Stable Processes. *On going paper*.