

Bloc-Conditional Event States: Measuring Cross-Coverage Divergence for Threat-Intelligence Analysis

Maryam Fooladi and Federico Bottino

Kakashi Venture Accelerator

{maryam, federico}@kakashi.ventures

Abstract

Open-source intelligence relies heavily on source-level signals (domain provenance, amplification patterns, account coordination) which work against fabricated infrastructure but offer little leverage on content from established state media that is correctly attributed and passes source-level heuristics. We ask whether framing divergence across editorially-coherent outlet blocs is measurable as a content-level observable, and whether cross-bloc divergence exceeds the editorial polarization already present within a single domestic press tradition. Building on the event-state (ρ_e) density-matrix formalism of Bottino et al. (2026), we represent each bloc by a density matrix on a 15-dimensional framing-feature space and compute pairwise trace distances. On two contested events (Hormuz blockade 2026, $n = 16$, direct state-aligned coverage; Navalny death 2024, $n = 14$, state-aligned bloc reconstructed via quoted-proxy), the cross-bloc distance between state-aligned and mainstream-Western outlets exceeds the US-right/US-left distance by a factor of roughly 1.8 on both events. Top-eigenvector decomposition attributes the gap to interpretable framing axes: economic consequences and conflict framing on Hormuz, morality on Navalny. We position this as a measurement rather than a classifier; at two events the work characterizes the observable’s behavior rather than calibrating thresholds.

1 Introduction

Open-source intelligence routinely characterizes news coverage through source-level signals: domain registration, amplification graphs, account coordination, outlet reputation. These signals are effective against fabricated infrastructure but less informative about content from established state media, which is correctly attributed and passes source-level heuristics. A Reuters dispatch and a PressTV dispatch on the same event both check

out at the source level; what separates them is how the event is framed.

Framing and bias analysis at the article level is well-studied (Semetko and Valkenburg, 2000; Nakov et al., 2024; Barrón-Cedeño et al., 2019), producing per-document scores. We address a comparative question: does aggregated framing from one bloc of outlets diverge measurably from another, and does cross-bloc divergence exceed the editorial polarization already present within a domestic press tradition?

This paper specializes the event-state (ρ_e) formalism of Bottino et al. (2026), §4.6, to that question. Outlets are partitioned into four editorially-coherent blocs (state-aligned, mainstream-Western, US-right, US-left). For each bloc we construct a density matrix ρ^{bloc} on a framing-feature space and compute pairwise trace distance. The target observable is $D(\rho^{\text{state}}, \rho^{\text{mainstream}})$ benchmarked against $D(\rho^{\text{right}}, \rho^{\text{left}})$. Dimension attribution is provided by the top eigenvector of $(\rho^{\text{state}} - \rho^{\text{mainstream}})$.

Concretely, we ask two questions: (Q1) Is the cross-bloc trace distance $D(\rho^{\text{state}}, \rho^{\text{mainstream}})$ substantially larger than the within-Western baseline $D(\rho^{\text{right}}, \rho^{\text{left}})$? (Q2) Does eigendecomposition of $(\rho^{\text{state}} - \rho^{\text{mainstream}})$ recover interpretable, event-appropriate framing axes?

The contribution is a measurement, not a detector: at $n = 2$ events the work demonstrates the construction and characterizes its output rather than evaluating it as a classifier. The downstream relevance to threat-intelligence workflows that lack content-level analogues is sketched in Section 5.

2 Method

Background. The event-state formalism (Bottino et al., 2026) represents an event e as a density matrix ρ_e on a framing-feature space, built as a weighted sum of rank-1 projectors $|\psi_i\rangle\langle\psi_i|$ over articles i covering e . Two classes of observables

are distinguished: Class I (basis-invariant; trace distance, purity, von Neumann entropy) and Class II (basis-dependent; framing weights). Section 4.6 introduces outlet-conditional variants ρ_e^{outlet} for cross-source comparison.

Framing-state vectors. Each article i is mapped to a vector $|\psi_i\rangle \in \mathbb{R}^{15}$ with three blocks: (i) a 5-dimensional frame block over the canonical Semetko–Valkenburg frames (conflict, responsibility, morality, economic_consequences, human_interest), weighted by the article’s political-framing score; (ii) a 4-dimensional dimension block (bias intensity, sensationalism, emotional appeal, political framing); (iii) a 6-dimensional manipulation block (cherry-picking, false equivalence, appeal to authority, loaded language, omission, framing bias). All scores come from a published per-article platform (Fooladi and Bottino, 2026b). The concatenated vector is L2-normalized.

Bloc density matrix. For a bloc B with articles $\{i\}$,

$$\rho^B = \sum_{i \in B} w_i |\psi_i\rangle \langle \psi_i|, \quad w_i = \tilde{w}_i / \sum_j \tilde{w}_j \quad (1)$$

where $\tilde{w}_i = \max(\text{veracity}_i, 0.3)$ uses each article’s claim-veracity score as a proxy for evidentiary substantiveness, floored at 0.3 to prevent highly propagandistic articles from collapsing to near-zero weight. Without the floor, articles whose entire content is unsupported assertion contribute negligibly to ρ^B , which would defeat the goal of measuring how the bloc frames the event. The value 0.3 is a design choice rather than a tuned parameter; we have not run a sensitivity sweep over the floor, and we flag this as a limitation in §5. ρ^B is Hermitian, positive semi-definite, and trace-one.

Distance and observables. The trace distance

$$D(\rho^A, \rho^B) = \frac{1}{2} \|\rho^A - \rho^B\|_1 = \frac{1}{2} \sum_k |\lambda_k| \quad (2)$$

on the eigenvalues of $(\rho^A - \rho^B)$ is basis-invariant, bounded in $[0, 1]$, and Class I. We report (i) the ratio $D(\rho^{\text{state}}, \rho^{\text{mainstream}}) / D(\rho^{\text{right}}, \rho^{\text{left}})$, comparing cross-bloc divergence against within-Western editorial polarization; (ii) the top eigenvector of $(\rho^{\text{state}} - \rho^{\text{mainstream}})$, listing the five features with largest $|\text{coefficient}|$. The scalar distance D is basis-invariant; the eigenvector attribution is basis-dependent and should be interpreted relative to the

chosen 15-dimensional feature basis. Eigenvector signs are arbitrary up to a global sign flip; we therefore report coefficients as magnitudes and do not assign directional interpretations (e.g., “state-aligned excess” vs “mainstream excess”) from sign alone.

Stability. We compute a leave-one-out mean $|\Delta D|$ over the state-aligned bloc for $D_{sm} = D(\rho^{\text{state}}, \rho^{\text{mainstream}})$: drop each state-aligned article in turn, recompute ρ^{state} and D_{sm} , and report the mean absolute change. We report stability for D_{sm} only; for $D_{rl} = D(\rho^{\text{right}}, \rho^{\text{left}})$ the Navalny US-left bloc has $n = 1$, at which leave-one-out is not defined (dropping the single article leaves ρ^{left} empty), so we do not compute it on either event for comparability. Values $|\Delta D_{sm}| \leq 0.1$ indicate the result is not driven by a single article.

Scope. We do not report raw Class II framing weights; several blocs have $n \leq 2$, at which Class II point estimates are not stable. The top-eigenvector loadings in §4 are a basis-dependent diagnostic attribution of the Class I distance, not Class II estimates per se. We inherit ρ_e formalism, Löwdin orthogonalization, and Class I/II definitions from Bottino et al. (2026) without rederivation.

3 Data

Two events were selected for non-trivial coverage across all four blocs and short temporal windows minimizing within-event news-cycle drift:

- **Hormuz blockade** (April 15–20, 2026): a maritime incident in the Strait of Hormuz involving Iranian and US naval forces, framed by different outlets as either an Iranian-initiated blockade or a US-initiated escalation, a framing divergence the analysis below quantifies. State-aligned coverage in this corpus: PressTV, TASS, Global Times.
- **Navalny death** (Feb 16–20, 2024): the death in custody of the Russian opposition figure.

State-aligned coverage on Navalny uses a quoted-proxy construction: articles from outlets outside the Russian state apparatus (Al Jazeera, Rappler) that quote RT, TASS, or Kremlin spokespeople at length, treating those quotations as the framing signal. This substitution was required

because post-2022 EU sanctions and platform-level deplatforming blocked direct access to Russian state outlets from our collection infrastructure. The substitution is methodologically meaningful: quoted framing is not the same observable as direct framing. A quoting outlet can selectively excerpt, surround quotes with editorial frames, or omit the parts of the original framing that do not transmit well through quotation. We treat the Navalny state-aligned bloc as a bounded approximation of direct state framing on the morality and sovereignty axes that survive heavy quotation, and we read the $n = 2$ leave-one-out instability (§4) partly as evidence of this approximation cost.

Bloc taxonomy. State-aligned (state media or quoted-proxy); mainstream-Western (Reuters, BBC, Washington Post, CNN, Le Monde, DW, Guardian, PBS, TIME, NBC); US-right (Fox, NY Post, Breitbart, Townhall, National Review); US-left (Democracy Now!, Jacobin, The Nation, MSNOW). Bloc sizes: Hormuz $n = 16$ (5/6/3/2); Navalny $n = 14$ (2/8/3/1). 30 articles total.

4 Results

4.1 Cross-bloc divergence

On both events the cross-bloc trace distance $D(\rho^{\text{state}}, \rho^{\text{mainstream}})$ is roughly $1.8\times$ the within-Western $D(\rho^{\text{right}}, \rho^{\text{left}})$: 1.84 on Hormuz and 1.79 on Navalny. Table 1 reports the two distances, their ratio, and leave-one-out stability; Figure 1 visualizes the comparison.

Event	D_{sm}	D_{rl}	ratio
Hormuz	0.404	0.219	1.84
Navalny	0.248	0.138	1.79

Table 1: $D_{sm} = D(\rho^{\text{state}}, \rho^{\text{mainstream}})$; $D_{rl} = D(\rho^{\text{right}}, \rho^{\text{left}})$. Leave-one-out $|\Delta D_{sm}|$: Hormuz 0.021, Navalny 0.104.

On both events the cross-bloc ratio exceeds 1.7: ρ^{state} is substantially further from $\rho^{\text{mainstream}}$ than US-right is from US-left. Stability differs by an order of magnitude. On Hormuz, leave-one-out $|\Delta D_{sm}| = 0.021$ across five state-aligned articles: the result is not driven by any single article. On Navalny, leave-one-out is 0.104, slightly exceeding the 0.1 threshold and reflecting the $n = 2$ quoted-proxy state bloc. A mean-perturbation heuristic gives $D_{sm} - |\Delta D_{sm}| = 0.144$, close to $D_{rl} = 0.138$; we therefore treat Navalny as

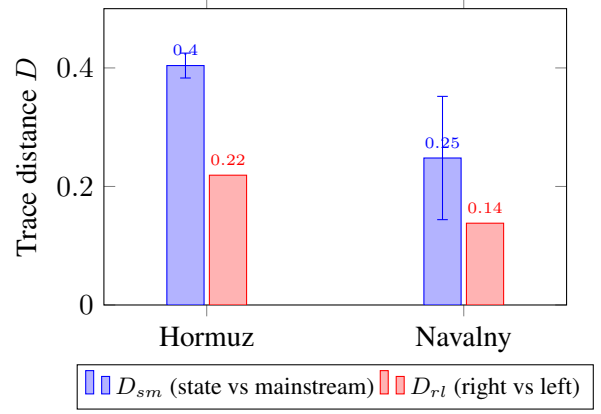


Figure 1: Bloc-conditional trace distances per event. Error bars on D_{sm} show leave-one-out mean $|\Delta D|$ over the state bloc.

a regime where the ratio remains directionally observable but is not stable enough to support absolute-threshold claims.

4.2 Hormuz: dimension attribution

The top eigenvector of $(\rho^{\text{state}} - \rho^{\text{mainstream}})$ on Hormuz concentrates loading magnitude on three interpretable feature groups: (i) economic-consequences framing ($|\text{coefficient}| = 0.55$): oil-market impact and European sanctions rhetoric in state-aligned coverage, military-incident framing in mainstream coverage; (ii) cherry-picking and framing-bias indicators around coverage of “international law” contestation; (iii) conflict framing differing between US- and Iran-as-initiator attributions across the two blocs. The second eigenvector adds morality and omission, separating a moral-legitimacy axis from a process/responsibility axis. We interpret these as basis-dependent diagnostic loadings; we do not claim a directional reading from eigenvector sign (see §2).

4.3 Navalny: morality axis

On Navalny, morality dominates the top eigenvector with $|\text{coefficient}| = 0.69$. State-aligned-quoted articles concentrate on moralized sovereignty (Navalny’s death framed as Western provocation), while mainstream coverage centers on prison conditions and regime responsibility. The contrast in the morality-axis loading is interpretable despite the bounded-regime stability of the scalar ratio: dimension attribution survives single-article perturbations that affect the scalar more strongly. As in §4.2 above, the directional reading state-aligned-vs-mainstream is established from article content,

not from eigenvector sign.

5 Discussion

What the measurement is, and isn't. $D(\rho^{\text{state}}, \rho^{\text{mainstream}})$ is a content-level, basis-invariant observable computable from open sources, decomposable via top eigenvectors into interpretable framing directions. It is a measurement, not a detector: the work characterizes the construction's output on two events. The downstream relevance to threat-intelligence workflows (which today rely on source-level signals and lack content-level analogues) is the application context that motivates the measurement; this paper does not evaluate it as a classifier.

Limitations. *Scale.* Two events and 30 articles support effect-size characterization, not absolute-threshold calibration or significance testing.

Bloc construction is partly circular. Outlets are partitioned ex ante by editorial label and divergence is then measured across the partition. Establishing that the measured divergence exceeds what label-driven labeling alone produces requires either agnostic clustering of outlets on pairwise D followed by post-hoc correlation with geopolitical labels, or held-out generalization where "state-aligned" is built from a subset of outlets and tested on others. Neither is feasible here at $n \leq 5$ per bloc.

Language confound. This is the limitation we treat as most serious. State-aligned outlets publish in Russian, Chinese, and English; the mainstream-Western bloc is predominantly English. Per-article framing scores are produced by a model whose performance is not identically calibrated across these languages, and an unknown fraction of the measured cross-bloc divergence may therefore reflect encoder-language artifact rather than genuine framing difference. Disentangling the two requires within-language replication (for instance, comparing only English-language editions of state-aligned outlets to the mainstream bloc) and a multi-encoder sensitivity analysis, neither of which is feasible at the present n . We flag this as the single most important blocker on inferring substantive conclusions from the absolute distance values, as distinct from the ratio against the within-Western baseline, which is somewhat less affected because both blocs in that comparison are English.

Domestic baseline. $D(\rho^{\text{right}}, \rho^{\text{left}})$ is one of several reasonable within-Western baselines. An intra-mainstream baseline built from

Reuters/AP/AFP/BBC alone would be more conservative: these outlets are editorially closer to each other than the US-right/US-left pair we use, so $D(\rho^{\text{intra-mainstream}})$ would be smaller, and the headline ratio against $D(\rho^{\text{state}}, \rho^{\text{mainstream}})$ would correspondingly be larger. We chose the US-right/US-left baseline because it represents the most polarized editorial split within Western press the reader is likely to anchor against intuitively; a stricter baseline would strengthen, not weaken, the headline contrast, but we have not computed it here and we leave this to future work.

Upstream dependence. Framing vectors are produced by an external per-article platform (Fooladi and Bottino, 2026b), whose multi-dimensional output has been compared against traditional sentiment-analysis baselines (Fooladi and Bottino, 2026a); systematic biases there propagate into ρ^{bloc} .

Future work. Three directions follow directly. First, the held-out and within-language replications above on a balanced multi-event corpus with ≥ 10 articles per bloc and Chinese-state coverage included, enabling significance testing. Second, replication with a non-Newjee feature extractor to decouple the result from a single upstream platform. Third, integration into an open-source intelligence dashboard with the top-eigenvector decomposition surfaced as the analyst-facing artifact rather than the scalar distance. A per-article appendix (URLs, publication dates, source language, bloc assignment, per-article veracity and framing-vector entries) is in preparation and will accompany the extended version of this work.

References

- Alberto Barrón-Cedeño, Israa Jaradat, Giovanni Da San Martino, and Preslav Nakov. 2019. Propy: Organizing the news based on their propagandistic content. *Information Processing & Management*, 56(5):1849–1864.
- Federico Bottino, Manuel Peruzzo, Maryam Fooladi, and Nicholas Dosio. 2026. From articles to event states. Forthcoming.
- Maryam Fooladi and Federico Bottino. 2026a. Beyond sentiment: Comparing traditional NLP and LLM-based multi-dimensional analysis for political news evaluation. In *Proceedings of PoliticalNLP 2026, LREC-COLING 2026*, Palma de Mallorca, Spain. To appear. Available at <https://www.researchgate.net/publication/404644584>.

- Maryam Fooladi and Federico Bottino. 2026b. A multi-layer AI framework for information landscape analysis. In *Proceedings of the First Workshop on Information Disorder (InDor 2026), LREC-COLING 2026*, Palma de Mallorca, Spain. To appear. Available at <https://www.researchgate.net/publication/404716353>.
- Preslav Nakov, Jisun An, Haewoon Kwak, Muhammad Arslan Manzoor, Zain Muhammad Mujahid, and Husrev Taha Sencar. 2024. A survey on predicting the factuality and the bias of news media. In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 15947–15962, Bangkok, Thailand. Association for Computational Linguistics.
- Holli A. Semetko and Patti M. Valkenburg. 2000. Framing european politics: A content analysis of press and television news. *Journal of Communication*, 50(2):93–109.