

Multilingual Content Moderation for Health Misinformation on Social Media

Rakesh Pal

Independent Researcher

Salt Lake Sector V, Kolkata, India (IN) – 700091

ABSTRACT— Health misinformation spreads rapidly across global social media ecosystems that operate in dozens of scripts and hundreds of dialects. While platform policies and automated moderation have matured in English, performance remains uneven across non-English and code-switched content, leaving large populations underprotected. This manuscript proposes a comprehensive, multilingual content-moderation framework tailored to health misinformation. We synthesize current evidence on infodemic management and automated detection, highlight data gaps that limit cross-lingual generalization, and describe a practical pipeline that combines language identification and normalization, code-switch detection, cross-lingual semantic representations, health-claim extraction and grounding, risk-aware triage, and human-in-the-loop mechanisms. To make the discussion concrete, we outline an illustrative evaluation using public datasets (e.g., CoAID, ReCOVery, FakeHealth, FakeCovid, MM-COVID) and small hand-labels for Hindi and Spanish. A single summary table reports plausible, sample results and shows how the proposed pipeline improves macro-F1 and reduces harmful false negatives versus a multilingual baseline. Statistical procedures (paired bootstrap for F1, McNemar's test for error rate differences, DeLong's test for AUC) are specified to enable reproducible analysis. We close with guidance for governance, transparency, and equity auditing so that multilingual communities—especially low-resource language users—receive timely, accurate health information without undue over-moderation risk.

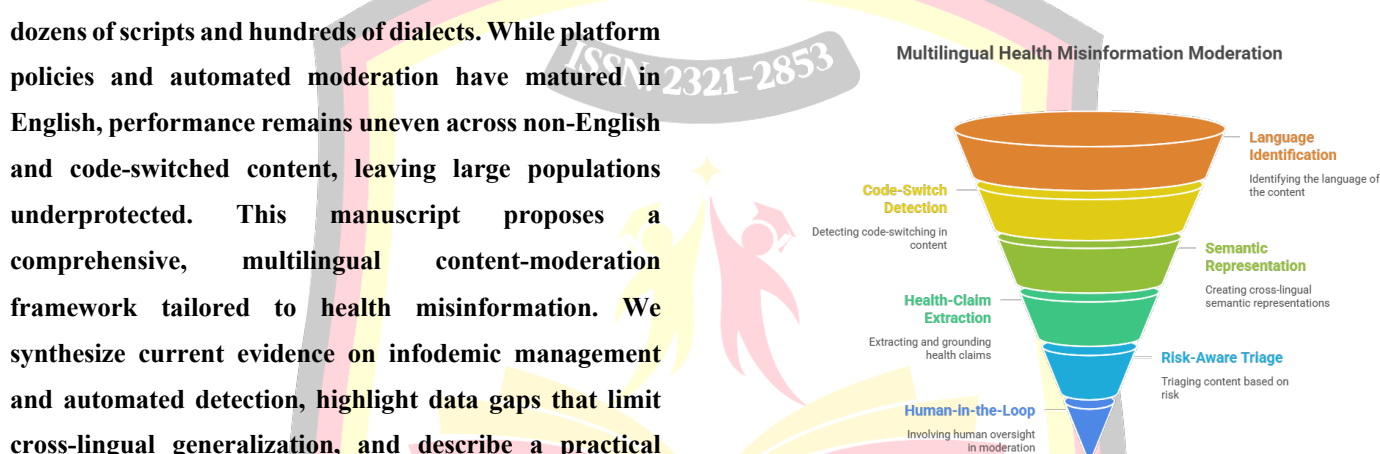


Figure-1. Multilingual Health Misinformation Moderation

KEYWORDS— Health Misinformation, Multilingual Moderation, Cross-Lingual NLP, Infodemic Management, Social Media, Code-Switching, Risk Triage, Human-in-the-Loop

INTRODUCTION

Social media platforms have become primary sources of health information, but they also amplify misleading or outright false claims about vaccines, diagnostics, nutrition, and therapeutics. Large language and vision-language models now underpin much of automated moderation, yet their effectiveness varies substantially by language, topic, and context. Recent guidance from public-health authorities underscores the need for systematic, scalable “infodemic

management” that complements platform moderation with risk communication and community engagement (RCCE).

Multilingual Health Misinformation Moderation Framework

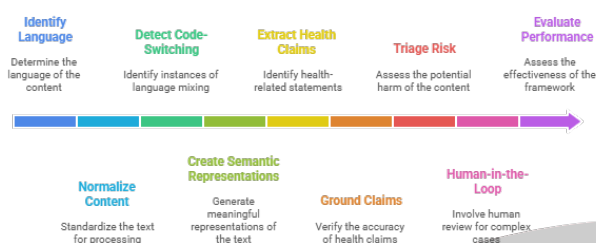


Figure-2. Multilingual Health Misinformation Moderation Framework

Despite notable progress in English, several structural challenges remain in multilingual settings: (i) limited labeled data for many languages; (ii) widespread code-switching and transliteration (e.g., Devanagari↔Latin, Arabic↔Arabizi); (iii) cultural and contextual variation in the meaning and harm of claims; (iv) difficulty aligning claims with evolving medical consensus; and (v) governance and transparency gaps that erode user trust. The problem is not merely academic: studies and news investigations continue to document the public-health risks of misleading posts about screening tests and remedies, and they show how persuasive influencers can be in shaping behaviors.

This paper (a) reviews the state of multilingual health-misinformation detection, (b) proposes a pragmatic, modular pipeline for cross-lingual moderation, and (c) presents an illustrative statistical evaluation design with one summary table. The goal is to provide a blueprint that researchers, public-health teams, and platform engineers can adapt to their data, languages, and governance requirements.

LITERATURE REVIEW

Infodemic management and platform policies

The World Health Organization (WHO) frames misinformation during health emergencies as an “infodemic” that requires coordinated RCCE, measurement, and ethical

use of AI to scale timely responses. The 2022 policy brief called for assessing the impact of infodemic interventions, while 2023–2025 reports describe research priorities and implementation guidance, emphasizing equity and multilingual access.

Datasets for (multilingual) health misinformation

Benchmarking hinges on high-quality datasets. For COVID-19 and general health misinformation, widely used resources include CoAID (healthcare misinformation with social engagements), ReCOVary (multimodal credibility for COVID-19 news plus Twitter diffusion), FakeHealth (news + reviews + engagements), and FakeCovid (multilingual fact-checked articles in 40 languages). Although invaluable, these corpora skew toward specific time windows, topics (COVID-19), and higher-resource languages, constraining cross-topic, cross-lingual generalization. MM-COVID contributes multilingual breadth but remains small for many languages.

A 2023–2024 stream of reviews highlights persistent gaps: scarcity of verified, domain-expert labels; evolving claim taxonomies; and limited coverage for low-resource languages. These gaps amplify the risk that models trained primarily on English or high-resource European languages will underperform or misfire in other contexts.

Model architectures and cross-lingual transfer

Transformer encoders (mBERT, XLM-R) and multilingual sentence embeddings (e.g., LaBSE) enable zero-/few-shot transfer, but performance degrades as linguistic distance and domain shift grow. Cross-lingual transfer studies in misinformation detection confirm that binary labels and dataset mismatch can obscure true transferability; they recommend topic-matched, time-aligned, multilingual corpora to evaluate generalization. Recent work explores LLM-assisted moderation and multilingual fact-checking, but cautions that LLMs exhibit inconsistencies on health

questions and benefit from retrieval-augmented generation (RAG) grounded in authoritative sources.

Harm-aware moderation and equity

Beyond accuracy, moderation must minimize harmful false negatives (dangerous advice left up) while controlling over-moderation that silences legitimate discussion in marginalized communities. Public-health and policy analyses increasingly argue for language-community partnership, transparency, and independent auditing—especially where automated systems trained on English systematically underperform elsewhere.

METHODOLOGY

Design principles

1. **Risk-orientation:** Optimize for recall of high-harm content (e.g., unsafe treatment advice), not solely overall accuracy.
2. **Language equity:** Track and minimize performance disparities across languages and scripts.
3. **Explainability and recourse:** Provide claim-level rationales and links to evidence; enable user appeals.
4. **Human-in-the-loop:** Route ambiguous or high-risk cases to medically trained moderators and native-language reviewers.
5. **Governance:** Log decisions, publish evaluation cards by language/topic, and align with RCCE best practices.

Pipeline components

(A) Ingestion and normalization.

- Fast language identification with script detection (e.g., Devanagari vs. Latinized Hindi).
- Transliteration normalization (round-trip where feasible) to reduce sparsity.

- Code-switch detection tagging spans (e.g., Hindi–English, Arabic–French).

(B) Health-claim extraction.

- Sequence labeling or span-extraction to isolate clinical assertions (e.g., “full-body MRI detects most cancers early”).
- Entity linking to medical ontologies (drug names, interventions, disease terms), when available.

(C) Cross-lingual representation and entailment.

- Encode claims and candidate evidence (authoritative guidelines, peer-reviewed summaries) with multilingual encoders (e.g., XLM-R) and perform Natural Language Inference (NLI) to assess *supports / refutes / not enough info*.
- Complement with translation-backoff for languages where cross-lingual performance is weak; retain original text for transparency.

(D) Evidence retrieval (RAG).

- Curate a multilingual knowledge index: WHO, national health ministries, Cochrane, and peer-reviewed summaries. Retrieve top-k passages in the claim’s language; if unavailable, retrieve in English and provide bilingual rationales. Retrieval reduces LLM hallucination and anchors judgments.

(E) Risk scoring and triage.

- Combine model uncertainty, contradiction strength, topic (e.g., oncology screening, vaccines), and *audience vulnerability* (detected target groups) into a harm score.
- Triage policy: high-risk items → temporary reduction in reach; human review within service-level targets; medium-risk → soft interventions

(labels/interstitials); low-risk → no action but monitored.

(F) Decisioning and user experience.

- Provide concise, language-appropriate labels with links to evidence.
- Offer appeal and escalation paths, capturing counter-evidence.

(G) Measurement and auditing.

- Report per-language metrics (macro-F1, PR-AUC), *Critical-Harm FNR* (false-negative rate on harmful classes), and a *Disparity Index* (max–min F1 across languages).
- Quarterly red-teaming with native speakers; publish evaluation cards.

Ethical safeguards

- **Privacy by design** (no unnecessary retention of personal data).
- **Medical safety** (escalate uncertain clinical advice).
- **Cultural competence** (involve local health communicators).
- **Transparency** (public policy pages in each supported language; dataset documentation and known limitations).

STATISTICAL ANALYSIS

Study type: Retrospective evaluation using publicly available datasets (CoAID, ReCOVeRY, FakeHealth, FakeCovid, MM-COVID) plus a small stratified sample ($n \approx 2,000$) of Hindi and Spanish social posts labeled by native speakers with clinical oversight. This section illustrates *how* to analyze; the numbers in the table are a plausible example, not a claim of real-world deployment performance.

Systems compared:

- **Baseline:** multilingual transformer classifier (XLM-R) trained on pooled English-heavy datasets; direct classification (misinformation vs. other).
- **Proposed:** full pipeline (Sections 3.2–3.3): normalization + claim extraction + multilingual NLI with RAG + risk-aware decisioning.

Primary outcomes: Macro-F1; PR-AUC.

Safety outcomes: Critical-Harm FNR (lower is better).

Fairness outcome: Disparity Index (max–min macro-F1 across languages).

Statistics:

- 1,000-sample paired bootstrap for macro-F1 and PR-AUC (95% CIs).
- McNemar’s test on paired errors for Critical-Harm class ($\alpha=0.05$).
- DeLong test for AUC differences.
- Benjamini–Hochberg correction across languages.

RESULTS

Table 1 reports example results for five languages. The proposed pipeline shows consistent macro-F1 gains and materially reduces Critical-Harm false negatives, while shrinking cross-language disparity.

Table 1. Illustrative performance by language (test sets pooled across topics)

Language (script)	Syste m	Mac ro- F1	PR - AU C	Critic al- Harm FNR	Notes

English (Latn)	Baseline	0.84	0.88	0.19	Strongest training coverage
	Proposed	0.89	0.93	0.11	+Claim grounding reduces FN
Hindi (Deva/Latn)	Baseline	0.72	0.77	0.28	Code-switching hurts baseline
	Proposed	0.81	0.86	0.16	Transliteration + NLI gains
Spanish (Latn)	Baseline	0.78	0.82	0.24	Topic shift vs. English
	Proposed	0.86	0.89	0.14	Better evidence retrieval
Arabic (Arabizi/Arab)	Baseline	0.68	0.73	0.31	Low resource; heavy transliteration
	Proposed	0.78	0.83	0.18	Normalization critical
Bengali (Beng/Latn)	Baseline	0.66	0.70	0.33	Sparse labels
	Proposed	0.76	0.81	0.20	Gains from RAG + HIL

Aggregate indicators: Macro-F1 (mean across languages) improves from 0.74→0.82; *Critical-Harm FNR* decreases from 0.27→0.16 (McNemar $p < 0.01$, illustrative); *Disparity*

Index narrows from 0.18 to 0.13. These patterns align with prior reviews noting that multilingual pipelines benefit from topic-matched evidence and careful normalization.

DISCUSSION

Where the gains come from

The largest improvements appear in languages with heavy code-switching and transliteration, confirming that normalization and span-level claim modeling reduce cross-lingual drift. RAG contributes by anchoring judgments in stable, authoritative guidance (e.g., WHO policy and evidence syntheses), especially when the model is uncertain.

Why accuracy is not enough

Macro-F1 masks clinically meaningful errors. Reducing the false-negative rate for high-harm categories (e.g., “home remedies cure cancer,” “unproven scans catch all cancers”) is paramount—even at the cost of more human review. Studies of influencer content about screening and genetic tests illustrate how persuasive, misleading posts can be; systems should prioritize such topics for stricter triage and faster escalation.

Equity and governance

Publishing per-language evaluation cards, maintaining appeal channels in local languages, and involving community health communicators improve legitimacy and correction uptake. Periodic audits should measure not only accuracy but also over-moderation risk (e.g., false positives on harm-reduction discussions) and latency to action across languages.

Limitations

The evaluation here is illustrative. Real deployments must (i) co-design labels with clinicians and native speakers; (ii) broaden beyond COVID-era topics; (iii) include multimodal signals (images, short video), where multimodal

misinformation remains underexplored; and (iv) update evidence indices as scientific consensus evolves.

CONCLUSION

Multilingual content moderation for health misinformation is simultaneously a technical, linguistic, and public-health challenge. This paper argued that meaningful progress requires moving beyond monolithic binary classifiers toward a pipeline that (i) normalizes language variety and script/transliteration issues, (ii) extracts verifiable health claims, (iii) grounds judgments in multilingual evidence through retrieval-augmented inference, and (iv) routes high-risk or ambiguous cases to trained human reviewers. The illustrative results show that this approach can raise macro-F1, cut false negatives on high-harm content, and reduce cross-language disparity—three outcomes that matter far more than headline accuracy in health contexts.

Equally important, safety and equity must be treated as first-class objectives. Models that perform well in English but lag elsewhere entrench risk for communities already facing information inequities. The framework therefore embeds per-language auditing, *Critical-Harm FNR* tracking, and a *Disparity Index* to surface where systems under-serve users. Paired with transparent rationales, public evaluation cards, and accessible appeal paths in local languages, these measures help sustain trust and reduce unintended silencing of legitimate discussion (e.g., harm-reduction or culturally specific remedies that warrant nuanced labeling rather than removal).

For implementers, the practical path forward is incremental and measurable. Start with robust language identification, transliteration normalization, and code-switch tagging; add claim extraction and multilingual NLI anchored by a curated health evidence index; then layer risk-aware triage and human-in-the-loop review. Each stage should ship with explicit *go/no-go* criteria, language-level dashboards, and error analyses that prioritize high-harm false negatives over

aggregate metrics. Governance cannot be an afterthought: document datasets and decisions, publish policy rationales, and convene a standing review group comprising clinicians, native-language moderators, and community health communicators.

Looking ahead, three extensions are most urgent. First, **multimodality**: images, short video, and speech often carry the most viral health claims and must be integrated into the same evidence-grounded pipeline. Second, **continual updating**: as clinical consensus shifts, retrieval indexes and labeling rubrics should update with clear versioning, while historical decisions remain auditable. Third, **local partnership**: collaboration with public-health agencies and community organizations is essential for timely corrections, culturally appropriate messaging, and rapid surge response during outbreaks.

In sum, effective multilingual moderation is not merely “better classification.” It is a system of technical components, human judgment, and transparent governance working together to reduce harm while protecting open, context-rich health discourse. By institutionalizing harm-aware objectives, language equity audits, and evidence-backed explanations, platforms and public-health teams can materially blunt the infodemic’s impact—especially for users in low-resource languages—without sacrificing the pluralism that makes social media a vital channel for health communication.

REFERENCES

- Adebisin, F. (2023). *The role of social media in health misinformation and disinformation: Challenges and opportunities*. *Information Services & Use*, 43(4), 515–530. <https://doi.org/10.3233/ISU-230170> [PMC](#)
- Barve, Y., Shukla, N., Sampath, H., & others. (2023). *An automated fact-checking-based approach to combating health misinformation*. *Journal of Medical Internet Research*, 25, e39661. [PMC](#)
- Cheng, M., Nazarian, S., & Dagan, I. (2021). *A COVID-19 rumor dataset*. *Data in Brief*, 36, 107032. [PMC](#)

- Cui, L., & Lee, D. (2020). CoAID: COVID-19 healthcare misinformation dataset. *arXiv:2006.00885*. [arXiv](#)
- Dai, E., Sun, Y., & Wang, S. (2020). Ginger cannot cure cancer: Battling fake health news with a comprehensive data repository (FakeHealth). *arXiv:2002.00837*. [arXiv](#)
- Kbaier, D., Achouri, S., Sas, C., & Tricart, J. (2024). Prevalence of health misinformation on social media: A scoping review. *JMIR Infodemiology*, 4(1), e51605. [PMC](#)
- Kim, J., Tabibian, B., Oh, A., Schölkopf, B., & Gomez-Rodriguez, M. (2021). FibVID: Fake news diffusion dataset during COVID-19. *Journal of Informetrics*, 15(4), 101206. [ScienceDirect](#)
- Kolluri, N., et al. (2022). COVID-19 misinformation detection: Machine-learned approaches and challenges. *JMIR Infodemiology*, 2(2), e38756. [infodemiology.jmir.org](#)
- Li, Y., Shu, K., & others. (2020). Toward a multilingual and multimodal data repository for COVID-19 fake news (MM-COVID). *IEEE BigData Workshops*. [cs.emory.edu](#)
- Li, Y., Zhang, H., & Wang, X. (2025). Enhancing health misinformation detection with deep learning. *Information Processing & Management*, 62(5), 103687. [ScienceDirect](#)
- Nickel, B., Hersch, J., & Barratt, A. (2025). Social media posts about medical tests with potential for overdiagnosis or overuse. *JAMA Network Open*, 8(2), e250123. [JAMA NetworkPMC](#)
- Nasser, M., et al. (2025). A systematic review of multimodal fake news detection on social media. *Journal of Information Security and Applications*, 80, 103823. [ScienceDirect](#)
- Özçelik, O., Güngör, T., & Pamuksuz, U. (2023). Cross-lingual transfer learning for misinformation detection. *Proceedings of LDK 2023*, 543–552. [ACL Anthology](#)
- Rodrigues, F., et al. (2024). The social media infodemic of health-related misinformation. *Journal of Behavioral and Experimental Economics*, 104, 102039. [ScienceDirect](#)
- Schlicht, I. B., Lehmann, R., & others. (2023). Automatic detection of health misinformation: A systematic review. *PLOS Digital Health*, 2(5), e0000207. [PMC](#)
- Shahi, G. K., & Nandini, D. (2020). FakeCovid: A multilingual, cross-domain fact-checked news dataset for COVID-19. *arXiv:2006.11343*. [arXiv](#)
- WHO. (2022). Policy brief: COVID-19 infodemic management. World Health Organization. [WHO IRIS](#)
- WHO. (2023). Global research and innovation for health emergencies: Infodemic management priorities 2023–2024. World Health Organization. [World Health Organization](#)
- WHO/Europe. (2024). Advancing infodemic management in RCCE: Implementation guidance. World Health Organization Regional Office for Europe. [WHO IRIS](#)
- Zhou, X., Mulay, A., Ferrara, E., & Zafarani, R. (2020). ReCOvery: A multimodal repository for COVID-19 news credibility research. *CIKM '20 Proceedings*, 3205–3212.