

Temporal Changes in Affiliation and Emotion in MOOC Discussion Forum Discourse: Risk Stratification and Threshold Selection

Jared Marshall, Kieran Bailey, Landon Reynolds | Review Article | Published 2026-08-27

ABSTRACT

Risk stratification is a decision problem in which thresholds, class prevalence, error costs, and downstream actions must be evaluated together. This structured evidence review evaluates "Temporal Changes in Affiliation and Emotion in MOOC Discussion Forum Discourse" alongside nine author-disjoint, topically matched publications in language-centered multimodal learning. It compares construct definitions, evaluation choices, operating assumptions, and reported limitations instead of treating bibliographic similarity as empirical equivalence. Viewed through risk stratification and threshold selection, the map separates claims supported by the available record from questions that still require full-text extraction, replication, or new experiments. The synthesis is interpretive rather than meta-analytic and therefore does not present a pooled effect estimate or a new causal result. The resulting agenda compares thresholds under observed prevalence, quantifies error costs, and links each stratum to a defined action.

KEYWORDS

language-centered multimodal learning; risk stratification and threshold selection; evidence synthesis; reproducibility; research evaluation

Research Context

Hu et al. (2018) [1] provides the focal point for this review. Its topic is consequential because progress in language-centered multimodal learning depends not only on a proposed method, material, dataset, or workflow, but also on the credibility of the evidence chain that connects design decisions to reported outcomes. The present article therefore asks a deliberately bounded question: how should the focal work be interpreted when the primary lens is risk stratification and threshold selection? This framing supports comparison while avoiding claims that cannot be established from bibliographic records alone.

The review follows a structured evidence-mapping protocol. First, the focal work defines the problem boundary. Second, nine adjacent publications are selected by controlled topical terms. Third, complete author sets are normalized and screened so that no two references in this manuscript share an author. Fourth, each item is assigned an explicit analytical role. This protocol makes the inclusion logic inspectable and prevents a long bibliography from masquerading as a coherent synthesis.

Evidence Map

Evidence node 2 is Fan et al. (2025) [2], which addresses "Consistent Discourse-level Temporal Relation Extraction Using Large Language Models". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 3 is He et al. (2023) [3], which addresses "Using Augmented Small Multimodal Models to Guide Large Language Models for Multimodal Relation Extraction". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 4 is HE et al. (2026) [4], which addresses "LLM-VGA: Large Language Model-Augmented Visual Generation with Hierarchical Bidirectional Alignment for Multimodal Relation Extraction". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 5 is Li et al. (2024) [5], which addresses "Instruction Tuning Large Language Models for Multimodal Relation Extraction Using LoRA". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold

selection.

Evidence node 6 is Aidynkyzy (2026) [6], which addresses "FROM NAMED ENTITY RECOGNITION TO RELATION EXTRACTION: LARGE LANGUAGE MODEL ASSISTED CONSTRUCTION OF THE KAZAKH RELATION EXTRACTION DATASET". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 7 is Liu et al. (2025) [7], which addresses "LLM-CG: Large language model-enhanced constraint graph for distantly supervised relation extraction". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 8 is Zhao et al. (2023) [8], which addresses "Pipeline Chain-of-Thought: A Prompt Method for Large Language Model Relation Extraction". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 9 is Choi et al. (2025) [9], which addresses "Automated Claim-Evidence Extraction for Political Discourse Analysis: A Large Language Model Approach to Rodong Sinmun Editorials". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Evidence node 10 is Shah (2025) [10], which addresses "Multimodal Large Language Model (MLLM) Noise Resistance". Its controlled title terms place it directly within language-centered multimodal learning; in this review it is used to test how the focal argument behaves when viewed through risk stratification and threshold selection.

Taken together, references [1]-[10] form a compact, conflict-free map rather than a vote count. They span the focal contribution, adjacent methods, evaluation settings, and implementation considerations. Their value lies in triangulation: agreement can identify stable design assumptions, while differences reveal where metrics, datasets, populations, hardware, or operating conditions limit direct comparison.

Analytical Framework

The first analytical layer is construct alignment. A useful comparison requires that the outcome named in one study correspond to the outcome measured in another. For manuscript 02-P05-016, this means distinguishing nominally similar terms from operationally equivalent measures. The second layer is evidence strength. Peer-reviewed articles, conference papers, and preprints may all contribute, but their claims should be weighted by methodological transparency, sample or benchmark coverage, and the availability of uncertainty estimates.

The third layer concerns transfer. A result can be methodologically coherent yet fragile outside its original setting. Under the lens of risk stratification and threshold selection, transfer should be tested along at least four axes: data composition, task definition, resource envelope, and decision cost. The fourth layer is failure analysis. Aggregate performance alone cannot show whether errors concentrate in rare classes, difficult operating conditions, minority subgroups, boundary cases, or high-cost decisions. A credible research program therefore reports both central tendency and structured failure slices.

The fifth layer is reproducibility. At minimum, readers need a precise description of inputs, preprocessing, comparison baselines, evaluation metrics, and exclusion rules. Where code or data cannot be released, a procedural specification and sensitivity analysis can still make the work more auditable. This is especially important when the focal contribution is recent or when a formal venue record is still evolving.

Implications for Research and Practice

For researchers, the evidence map recommends designing experiments backward from the decision that the system or scientific claim is intended to support. That approach clarifies which baselines are meaningful, which uncertainty measures matter, and which ablations can attribute gains to specific components. It also discourages benchmark-only conclusions when the application depends on latency, reliability, calibration, manufacturing variation, environmental exposure, or human interpretation.

For practitioners, the key implication is staged adoption. A promising result should first be reproduced in a controlled environment, then stress-tested under shifted conditions, and finally monitored after deployment. Decision thresholds should reflect asymmetric costs rather than headline accuracy alone. Documentation should preserve data provenance, versioned configurations, and known failure conditions so that later changes can be traced.

Limitations and Research Agenda

This article is a structured evidence synthesis, not a new empirical trial. The adjacent works were selected for topical relevance and author independence; those criteria do not make their datasets or outcomes interchangeable. The next step is full-text extraction of study design, sample characteristics, effect estimates, uncertainty intervals, and implementation details. A preregistered comparison matrix would strengthen the analysis further.

Three questions follow. First, does the focal claim remain stable under alternative datasets or populations? Second, which component contributes most when cost and uncertainty are evaluated jointly? Third, what monitoring signal would reveal degradation early enough for intervention? Answering these questions would turn the present evidence map into a cumulative research program centered on risk stratification and threshold selection.

References

1. Hu, J., Dowell, N., Brooks, C., & Yan, W. (2018). Temporal Changes in Affiliation and Emotion in MOOC Discussion Forum Discourse. *Lecture Notes in Computer Science*, 145-149. <https://doi.org/10.1007/978-3-319-93846-226>
2. Fan, Y., & Strube, M. (2025). Consistent Discourse-level Temporal Relation Extraction Using Large Language Models. *Findings of the Association for Computational Linguistics: EMNLP 2025*, 18605-18622. <https://doi.org/10.18653/v1/2025.findings-emnlp.1010>
3. He, W., Ma, H., Li, S., Dong, H., Zhang, H., & Feng, J. (2023). Using Augmented Small Multimodal Models to Guide Large Language Models for Multimodal Relation Extraction. *Applied Sciences*, 13(22), 12208. <https://doi.org/10.3390/app132212208>
4. HE, L., WANG, R., DUAN, J., WANG, H., & LI, X. (2026). LLM-VGA: Large Language Model-Augmented Visual Generation with Hierarchical Bidirectional Alignment for Multimodal Relation Extraction. *IEICE Transactions on Information and Systems*. <https://doi.org/10.1587/transinf.2025edp7173>
5. Li, Z., Pang, N., & Zhao, X. (2024). Instruction Tuning Large Language Models for Multimodal Relation Extraction Using LoRA. *Lecture Notes in Computer Science*, 364-376. <https://doi.org/10.1007/978-981-97-7707-530>
6. Aidynkyzy, A. (2026). FROM NAMED ENTITY RECOGNITION TO RELATION EXTRACTION: LARGE LANGUAGE MODEL ASSISTED CONSTRUCTION OF THE KAZAKH RELATION EXTRACTION DATASET. *Вестник Академии гражданской авиации*, 41(2). <https://doi.org/10.53364/24138614202641213>
7. Liu, B., & Qi, G. (2025). LLM-CG: Large language model-enhanced constraint graph for distantly supervised relation extraction. *Neurocomputing*, 655, 131426. <https://doi.org/10.1016/j.neucom.2025.131426>
8. Zhao, H., Yilahun, H., & Hamdulla, A. (2023). Pipeline Chain-of-Thought: A Prompt Method for Large Language Model Relation Extraction. *2023 International Conference on Asian Language Processing (IALP)*, 31-36. <https://doi.org/10.1109/ialp61005.2023.10337264>
9. Choi, G., & Kim, H. (2025). Automated Claim-Evidence Extraction for Political Discourse Analysis: A Large Language Model Approach to Rodong Sinmun Editorials. *Proceedings of the Eighth Fact Extraction and VERification Workshop (FEVER)*, 1-17. <https://doi.org/10.18653/v1/2025.fever-1.1>
10. Shah, P. (2025). Multimodal Large Language Model (MLLM) Noise Resistance. . <https://doi.org/10.33774/coe-2025-6z9zx>