

# Trust-Informed Expert Routing for Multimodal Biomedical Classification

## Abstract

Trustworthy Multimodal Biomedical AI is increasingly shaped by the need to reconcile performance with evidence quality, resource limits, and transfer across settings. The present synthesis investigates routing modalities according to conditional reliability rather than assuming every signal is always useful. The source set brings together 1 focal paper with 13 independently retrieved publications validated against DOI-registration metadata. The analysis is organized around modality risk, expert routing, missing data, calibration, and clinical safety. The analysis declines to treat results from heterogeneous studies as exchangeable, the review compares problem definitions, methodological assumptions, and validation boundaries. Across the literature, the comparison suggests that advances in trustworthy multimodal biomedical AI become credible when measurement, model, and clinical decision are evaluated together and when uncertainty about calibration is reported explicitly. The comparative structure connects method selection to decision consequence and recurring validity threats, and proposes a research agenda centered on well-specified controls, robustness tests, and reproducible workflows.

## Keywords

Trustworthy Multimodal Biomedical AI; Modality Risk; Expert Routing; Missing Data; Calibration; Clinical Safety

## 1. Introduction

Contemporary investigation of trustworthy multimodal biomedical AI sits at an interface between scientific ambition and practical constraint. Researchers pursue not only stronger nominal performance but also explanations of the conditions and mechanisms behind success. This difficulty is evident in routing modalities according to conditional reliability rather than assuming every signal is always useful. A conclusion supported by a bounded material, dataset, population, scale, or operating trace can erode beyond the conditions that produced it. Evidence integration should therefore preserve the unit of analysis and the decision context. Equally important is distinguishing a mechanism from a correlation, a benchmark advantage from a deployment advantage, and an interpretable diagnostic from a causal explanation.

Comparison proceeds across five linked dimensions: modality risk, expert routing, missing data, calibration, clinical safety. The lenses were selected because they jointly cover what is designed, how it is measured, and what must remain stable for the conclusion to transfer. The central organizing idea is measurement, model, and clinical decision. Under that principle, innovation must be accompanied by validation; the comparison also evaluates comparator fairness, uncertainty reporting, and real operating constraints. The present article offers conceptual synthesis and is not presented as a new experiment and supplies no synthetic performance claim.

## 2. Clinical and Measurement Background

Four linked elements clarify the problem: the input evidence, the transformation applied to it, the output used for evaluation, and the decision that follows. In trustworthy multimodal biomedical AI, the chain is commonly fragmented even though errors propagate across them. Upstream noise may grow as it passes through a flexible model; an apparently accurate output can still be poorly calibrated for the final decision. A direct requirement is that, external validation should accompany aggregate performance. A strong comparison states which factors remain invariant and which may change, then selects metrics that correspond to the intended use rather than to convenience alone.

A further foundation is explicit scope. Modality risk defines the local design problem, while clinical safety determines whether the design can survive outside that boundary. Connecting the endpoints are expert routing and missing data connect those ends by exposing trade-offs that a single score conceals. Under this view, negative evidence and perturbation analysis reveals the interval in which an explanation can still be defended. For applied work, documentation of patient-level heterogeneity is therefore part of the scientific result, not an administrative afterthought.

## 3. Analytical Approaches

The target study X8-01, TIER-MoE: Trust-Informed Expert Routing via Conditional Modality Risk for Multimodal Fusion in Biomedical Classification [1], addresses a concrete part of trustworthy multimodal biomedical AI through trust-informed conditional routing among modality experts in biomedical classification. It identifies a representation, mechanism, or evaluation boundary needed to compare claims. Against the focus on routing modalities according to conditional reliability rather than assuming every signal is always useful, the paper serves as a design case with explicit limits: transfer may depend on its sensing regime, preparation, workload, population, or benchmark. The synthesis preserves that scope and tests its assumptions against neighboring work.

Across the focal studies, modality risk should be interpreted jointly with expert routing. A design can improve one while imposing a less visible trade-off on the other. The evidence can be organized in a matrix whose rows are operating conditions and whose columns are evidence quality, computation or process cost, robustness, and consequence of failure. The matrix is designed to prevent an attractive mechanism from being detached from the circumstances that support it. The matrix helps determine which ablations are informative: ablation design should target causal attribution instead of numerical proliferation.

Missing data provides the bridge from method to decision. A minimal evaluation must separate familiar inputs from perturbations and retain distributional summaries, and test plausible alternative explanations. Where calibration is central, calibration or sensitivity is as important as ranking. Where costs or hazards are asymmetric, utility-weighted assessment is more revealing than undifferentiated accuracy. These decisions need not homogenize studies; they make the reasons for their differences auditable.

## 4. Comparative Evidence

For triangulation, the literature includes Dual Routing Mixture-of-Experts for Multi-Scale Representation Learning in Multimodal Emotion Recognition [2]; Uncertainty-aware mixture of experts for robust multimodal sentiment analysis [3]; Intelligent LLM Orchestration: Advanced Mixture of Experts Routing for Large Language Model Systems [4]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of trustworthy multimodal biomedical AI. Together they illuminate modality risk: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

The design space is broadened by Attention-enhanced Mixture of Experts for multimodal breast cancer detection [5]; NeuroMoE++: Patient-Adaptive Multi-Level Multimodal Fusion With Mixture-of-Experts for Neurological Diso... [6]; Spatial-Spectral Operator Mixture Transformer With Mixture-of-Experts Fusion for Endoscopic Classificati... [7]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of trustworthy multimodal biomedical AI. Together they illuminate expert routing: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

A first comparison set connects Mixture of Experts Classification Using a Hierarchical Mixture Model [8]; PAM-MoE-AD: A plane-aware multi-stage mixture-of-experts framework for Alzheimer’s disease classificatio... [9]; Classification of ECG arrhythmia by a modular neural network based on Mixture of Experts and Negatively... [10]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of trustworthy multimodal biomedical AI. Together they illuminate missing data: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

The neighboring evidence base brings together Multi-Scale Gaussian Mixture Model-Gated Mixture of Experts for Fine-Grained Insect Pest Classification [11]; CAG-MoE: Multimodal Emotion Recognition with Cross-Attention Gated Mixture of Experts [12]; Dynamic Language Routing Mixture of Experts Model for Multilingual Speech Recognition [13]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of trustworthy multimodal biomedical AI. Together they illuminate calibration: some emphasize representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

A complementary strand is visible in Multimodal Gated Mixture of Experts Using Whole Slide Image and Flow Cytometry for Multiple Instance Lea... [14]. These publications were retrieved and verified through DOI or publisher metadata, and their titles directly intersect the problem boundary of trustworthy multimodal biomedical AI. Together they illuminate clinical safety: some emphasize

representations or mechanisms, whereas others foreground validation and application conditions. The useful comparison is not a leaderboard across incompatible settings, but a structured reading of what each study controls, leaves variable, and treats as a meaningful outcome. Divergence can then reveal a change in scale, data process, endpoint, or operating constraint.

## 5. Clinical Translation

Deployment decisions benefit from a sequence of checks. First, define the decision and its failure cost. Second, specify the evidence and operating envelope. Third, choose a method whose inductive assumptions match that evidence. Fourth, test calibration and clinical safety under controlled perturbations. Finally, retain provenance so that an unexpected outcome can be traced to data, configuration, material batch, environmental condition, or user interaction. This ordering holds routing modalities according to conditional reliability rather than assuming every signal is always useful connected to observable requirements.

Across applications, responsible progress in trustworthy multimodal biomedical AI is determined by coupling as well as local design. Interfaces determine how uncertainty moves between measurement and inference, between algorithm and operator, or between microscopic design and system behavior. Teams spanning disciplines should predefine acceptance rules and triggers for revising conclusions. When those agreements are explicit, efficiency can be improved without concealing losses in validity, safety, or interpretability.

## 6. Limitations and Future Directions

This synthesis is limited by heterogeneity in terminology, study design, and reporting granularity. Verified authorship and venue do not substitute for independent empirical replication. Rapid publication cycles can also alter venues and versions. The supplied focal strings remain preserved while questionable normalization is shown only as a candidate.

Subsequent studies need to preregister comparison protocols where feasible, retain machine-readable setup details while testing at least one nontrivial shift in patient-level heterogeneity. Research should also connect failure frequency to an explanatory taxonomy. For trustworthy multimodal biomedical AI, a valuable shared benchmark would combine ordinary performance, operational burden, uncertainty calibration, and resilience. Such a benchmark would reward designs that remain useful when evidence is incomplete or conditions change.

## 7. Conclusion

Research across trustworthy multimodal biomedical AI is most informative when its studies are read relationally rather than a collection of isolated scores. The focal studies show how routing modalities according to conditional reliability rather than assuming every signal is always useful; the additional literature reveals the assumptions required to interpret those contributions. Across modality risk, expert routing, missing data, calibration, and clinical safety, a durable conclusion is the need for alignment: the representation or mechanism must match the problem, assessment must correspond to the decision and deployment claims to the evaluated envelope. Progress built on that alignment is easier to reproduce, safer to transfer, and more informative when it fails.

## References

1. Chang, Y., Cheng, A., Wu, C., Wang, Z., Chen, J., Chattopadhyay, T., ... & Bogdan, P. (2026). TIER-MoE: Trust-Informed Expert Routing via Conditional Modality Risk for Multimodal Fusion in Biomedical Classification. arXiv preprint arXiv:2607.27289.

2. Chae, D. E., & Lee, S. P. (2025). Dual Routing Mixture-of-Experts for Multi-Scale Representation Learning in Multimodal Emotion Recognition. *Electronics*, 14(24), 4972. <https://doi.org/10.3390/electronics14244972>
3. Xiao, X., Wang, X., Qi, S., & Li, H. (2026). Uncertainty-aware mixture of experts for robust multimodal sentiment analysis. *Pattern Recognition*, 179, 113417. <https://doi.org/10.1016/j.patcog.2026.113417>
4. Reddy Bora, S., & Kishan Anapu, S. (2025). Intelligent LLM Orchestration: Advanced Mixture of Experts Routing for Large Language Model Systems. *International Journal of Science and Research (IJSR)*, 1833-1838. <https://doi.org/10.21275/sr25627232328>
5. Abdullakutty, F., Akbari, Y., Al-Maadeed, S., Saady, R. A., & Rasheed, H. M. A. (2026). Attention-enhanced Mixture of Experts for multimodal breast cancer detection. *Biomedical Signal Processing and Control*, 118, 109576. <https://doi.org/10.1016/j.bspc.2026.109576>
6. Raza, W. H., Wen, Y., Shah, A. B., Shen, Y., Martinez-Lemus, J., Schiess, M. C., et al. (2026). NeuroMoE++: Patient-Adaptive Multi-Level Multimodal Fusion With Mixture-of-Experts for Neurological Disorder Classification. *IEEE Transactions on Biomedical Engineering*, 73(8), 2784-2794. <https://doi.org/10.1109/tbme.2026.3691657>
7. Alhaj, Z., & Altairi, A. (2026). Spatial-Spectral Operator Mixture Transformer With Mixture-of-Experts Fusion for Endoscopic Classification. *IEEE Access*, 14, 49928-49944. <https://doi.org/10.1109/access.2026.3679531>
8. Titsias, M. K., & Likas, A. (2002). Mixture of Experts Classification Using a Hierarchical Mixture Model. *Neural Computation*, 14(9), 2221-2244. <https://doi.org/10.1162/089976602320264060>
9. Kumar, H., & Agarwal, R. (2026). PAM-MoE-AD: A plane-aware multi-stage mixture-of-experts framework for Alzheimer's disease classification from sMRI. *Biomedical Physics & Engineering Express*, 12(3), 035092. <https://doi.org/10.1088/2057-1976/ae7c0a>
10. Javadi, M., Arani, S. A. A. A., Sajedin, A., & Ebrahimpour, R. (2013). Classification of ECG arrhythmia by a modular neural network based on Mixture of Experts and Negatively Correlated Learning. *Biomedical Signal Processing and Control*, 8(3), 289-296. <https://doi.org/10.1016/j.bspc.2012.10.005>
11. Şahin, N., Alpaslan, N., & Hanbay, D. (2026). Multi-Scale Gaussian Mixture Model-Gated Mixture of Experts for Fine-Grained Insect Pest Classification. *Electronics*, 15(11), 2268. <https://doi.org/10.3390/electronics15112268>
12. Mengara Mengara, A. G., & Moon, Y. k. (2025). CAG-MoE: Multimodal Emotion Recognition with Cross-Attention Gated Mixture of Experts. *Mathematics*, 13(12), 1907. <https://doi.org/10.3390/math13121907>
13. Wang, J., Wang, Y., Bao, F., & Gao, G. (2025). Dynamic Language Routing Mixture of Experts Model for Multilingual Speech Recognition. *Data Intelligence*, 8(2), 20250139. <https://doi.org/10.3724/2096-7004.di.2025.0139>
14. Hashimoto, N., Hanada, H., Miyoshi, H., Nagaishi, M., Sato, K., Hontani, H., et al. (2024). Multimodal Gated Mixture of Experts Using Whole Slide Image and Flow Cytometry for Multiple Instance Learning Classification of Lymphoma. *Journal of Pathology Informatics*, 15, 100359. <https://doi.org/10.1016/j.jpi.2023.100359>