

MotiMem - Motion-Aware Approximate Memory for Energy-Efficient Neural Perception in Autonomous Vehicles: Risk Stratification and Threshold Selection

Brett Campbell, Bryce Hayes, Carson Coleman | 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 "MotiMem: Motion-Aware Approximate Memory for Energy-Efficient Neural Perception in Autonomous Vehicles" alongside nine author-disjoint, topically matched publications in resource-efficient autonomous perception. 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

resource-efficient autonomous perception; risk stratification and threshold selection; evidence synthesis; reproducibility; research evaluation

Research Context

Que et al. (2026) [1] provides the focal point for this review. Its topic is consequential because progress in resource-efficient autonomous perception 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 Costa et al. (2025) [2], which addresses "PerceptNet-V2X: Perception Network for Vehicle to Everything Scenarios in Autonomous Driving". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Chandrashekhar et al. (2025) [3], which addresses "Reinforcement learning in Autonomous vehicle communication under V2X Framework using edge computing and Privacy-Preserving AI". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Guo et al. (2025) [4], which addresses "When Autonomous Vehicle Meets V2X Cooperative Perception: How Far Are We?". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Asaju et al. (2025) [5], which addresses "A Survey of Multi-Layered Cybersecurity Threats in Connected and Autonomous Vehicles: Risks from Edge Computing and V2x Protocols". Its controlled title terms place it

directly within resource-efficient autonomous perception; 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 Garcia et al. (2024) [6], which addresses "Edge Learning Optimization in Task-Oriented NOMA Communications for Autonomous Vehicle Perception". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Wang (2025) [7], which addresses "Edge Intelligence for V2X Communications: Advances in Collaborative Perception and Privacy Preservation". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Richards et al. (2025) [8], which addresses "Edge-Enabled Collaborative Object Detection for Real-Time Multi-Vehicle Perception". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Roy (2025) [9], which addresses "Beyond V2X: A Review of Large Language Models in In-Vehicle Autonomous Driving Systems". Its controlled title terms place it directly within resource-efficient autonomous perception; 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 Soorchaee et al. (2025) [10], which addresses "Extensible Heterogeneous Collaborative Perception in Autonomous Vehicles with Codebook Compression". Its controlled title terms place it directly within resource-efficient autonomous perception; 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-P10-016, 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. Que, H., Liu, M., Xie, J., Gao, H., Sun, J., Xu, H., Yao, H., & Qiao, F. (2026). MotiMem: Motion-Aware Approximate Memory for Energy-Efficient Neural Perception in Autonomous Vehicles. *arXiv*. <https://doi.org/10.48550/arXiv.2603.27108>
2. Costa, B.-T., Pereira, C., & Maykol Pinto, A. (2025). PerceptNet-V2X: Perception Network for Vehicle to Everything Scenarios in Autonomous Driving. *IEEE Access*, 13, 182645-182660. <https://doi.org/10.1109/access.2025.3624285>
3. Chandrashekhar, B., & Shakkeera, L. (2025). Reinforcement learning in Autonomous vehicle communication under V2X Framework using edge computing and Privacy-Preserving AI. 2025 2nd International Conference on New Frontiers in Communication, Automation, Management and Security (ICCAMS), 1-8. <https://doi.org/10.1109/iccams65118.2025.11234562>
4. Guo, A., Zhang, S., Tang, E., Gao, X., Pang, H., Tian, H., Mu, Y., Wen, W., Fang, C., & Chen, Z. (2025). When Autonomous Vehicle Meets V2X Cooperative Perception: How Far Are We?. 2025 40th IEEE/ACM International Conference on Automated Software Engineering (ASE), 1169-1181. <https://doi.org/10.1109/ase63991.2025.00101>
5. Asaju, B.-J., Jung, W., & Wakili, A. (2025). A Survey of Multi-Layered Cybersecurity Threats in Connected and Autonomous Vehicles: Risks from Edge Computing and V2x Protocols. <https://doi.org/10.2139/ssrn.5260375>
6. Garcia, C.-E., Camana, M.-R., Mohammad, A.-B., Querol, J., & Chatzinotas, S. (2024). Edge Learning Optimization in Task-Oriented NOMA Communications for Autonomous Vehicle Perception. 2024 IEEE Wireless Communications and Networking Conference (WCNC), 1-6. <https://doi.org/10.1109/wcnc57260.2024.10570514>
7. Wang, S. (2025). Edge Intelligence for V2X Communications: Advances in Collaborative Perception and Privacy Preservation. *Journal of Big Data and Computing*, 3(4), 73-81. <https://doi.org/10.62517/jbdc.202501409>
8. Richards, E., Thapa, B., & Mashayekhy, L. (2025). Edge-Enabled Collaborative Object Detection for Real-Time Multi-Vehicle Perception. 2025 IEEE International Conference on Edge Computing and Communications (EDGE), 13-22. <https://doi.org/10.1109/edge67623.2025.00011>
9. Roy, A. (2025). Beyond V2X: A Review of Large Language Models in In-Vehicle Autonomous Driving Systems. <https://doi.org/10.36227/techrxiv.175571972.28517288/v1>
10. Soorchaei, B.-E., Raftari, A., & Fallah, Y.-P. (2025). Extensible Heterogeneous Collaborative Perception in Autonomous Vehicles with Codebook Compression. *Robotics*, 14(12), 186. <https://doi.org/10.3390/robotics14120186>