

Advances in Adaptive Intelligence
REVIEW | COMMISSIONED SYNTHESIS

Neuro-Symbolic Adaptation for Autonomous Systems

Alex Morgan - Corresponding author

Abstract

This review examines neuro-symbolic adaptation for autonomous systems. The organizing question is when learned perception and symbolic structure can support transparent adaptation under novel conditions. Ten related scholarly sources are synthesized through a decision-centered framework spanning problem definition, mechanism, measurement, evaluation, implementation, and governance. The review does not invent experiments, pooled estimates, or unreported quantitative results. It instead evaluates the strength and transferability of the available evidence, with particular attention to assuming symbolic components guarantee faithful or correct reasoning. The resulting framework links technical or empirical performance to explicit use conditions and identifies tests that should precede wider adoption in adaptive autonomous decision systems.

Keywords: neuro-symbolic AI, autonomous systems, adaptation, reasoning, explainability

Synthesis lens	Principal question
Mechanism	the chain connecting evidence inputs, neuro-symbolic adaptation for autonomous systems, intermediate outputs, and downstream decisions
Evaluation	construct validity, comparative performance, uncertainty, transfer across settings, and decision utility
Primary risk	assuming symbolic components guarantee faithful or correct reasoning

1. Introduction

Research on neuro-symbolic adaptation for autonomous systems sits at the boundary between a promising component and a consequential decision. The core question is when learned perception and symbolic structure can support transparent adaptation under novel conditions. Answering it requires more than assembling favorable results. It requires a chain of evidence that makes the problem boundary, mechanism, measurement, comparator, uncertainty, and accountable decision visible to the reader.

This article presents a structured narrative review of ten related publications. Sources were selected for topical relevance, traceable publication metadata, and complementary coverage of technical, empirical, and governance questions. No meta-analytic pooling is attempted because the included studies differ in designs, populations, system boundaries, and outcomes. The purpose is decision-oriented synthesis: to identify what each source can support, where transfer remains uncertain, and which evidence would justify the next stage of use.

The central risk is assuming symbolic components guarantee faithful or correct reasoning. A top-line metric or broad policy label can appear decisive while concealing the conditions that produced it. Accordingly, the synthesis separates observation from interpretation and implementation from validation. This separation is especially important when the same system creates benefits for one user or institution while transferring cost, risk, or administrative burden to another.

2. Evidence Base and Problem Boundary

A defensible review begins by defining the unit of analysis. For the present topic, the relevant boundary includes inputs, institutional or technical context, the mechanism producing an output, the person or system acting on that output, and the downstream outcome. Evidence falls outside the boundary when it measures only an intermediate proxy yet is used to claim an end-to-end benefit.

Hamilton et al. (2022) addresses "Is neuro-symbolic AI meeting its promises in natural language processing? A structured review." In this synthesis, that work is used as a source on problem definition within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Nawaz et al. (2025) addresses "A review of neuro-symbolic AI integrating reasoning and learning for advanced cognitive systems." In this synthesis, that work is used as a source on conceptual boundary within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Gaur et al. (2022) addresses "Knowledge-Infused Learning: A Sweet Spot in Neuro-Symbolic AI." In this synthesis, that work is used as a source on measurement within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Together, these works provide complementary entry points, but their conclusions should not be flattened into a single ranking. Differences

in data generation, setting, temporal horizon, and outcome definition may be substantively meaningful. A review should therefore preserve those differences and state where analogy, rather than direct replication, connects one domain to another.

3. Mechanism and System Design

The operative mechanism is the chain connecting evidence inputs, neuro-symbolic adaptation for autonomous systems, intermediate outputs, and downstream decisions. This formulation discourages component-only reasoning. A model, platform, policy, assay, or infrastructure service acquires practical meaning only when its interfaces and use conditions are specified. The evidence chain should describe what information is available, what transformation occurs, which assumptions make that transformation credible, and who is authorized to act.

Colelough and Regli (2025) addresses "Neuro-Symbolic AI in 2024: A Systematic Review." In this synthesis, that work is used as a source on system mechanism within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Hagos and Rawat (2024) addresses "Neuro-Symbolic AI for Military Applications." In this synthesis, that work is used as a source on representation choice within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Li et al. (2024) addresses "Deep expert network: A unified method toward knowledge-informed fault diagnosis via fully interpretable neuro-symbolic AI." In this synthesis, that work is used as a source on failure analysis within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

These studies also show why modular claims require interface tests. A component may perform well while the complete pathway fails because inputs drift, users interpret outputs differently, recovery semantics are ambiguous, or institutional incentives change. Design documentation should therefore include not only the intended pathway but also foreseeable deviations, degraded modes, and conditions under which the system should stop or defer.

4. Evaluation and Uncertainty

Evaluation should center on construct validity, comparative performance, uncertainty, transfer across settings, and decision utility. Average performance remains useful, but it should be accompanied by distributions, error categories, relevant subgroups or operating regimes, and uncertainty at the point where a decision changes. Comparators must isolate the value of the proposed addition rather than compare a mature intervention with an intentionally weak baseline.

Chandre et al. (2025) addresses "Neuro-Symbolic AI: A Future of Tomorrow." In this synthesis, that work is used as a source on comparative evaluation within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction

keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Prenosil et al. (2025) addresses "Neuro-symbolic AI for auditable cognitive information extraction from medical reports." In this synthesis, that work is used as a source on uncertainty within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

External validation should introduce meaningful shifts in setting, time, population, workload, market structure, or environmental conditions. A nominally independent sample can still be too similar to test transfer. Conversely, a failed transfer is scientifically informative when the changed conditions are measured and the mechanism of failure is examined rather than hidden in an aggregate score.

5. Translation and Governance

Translation to adaptive autonomous decision systems depends on more than technical readiness. Decision rights, documentation, maintenance, monitoring, appeal, and recovery are part of the intervention. The governance requirement is traceable evidence, named responsibility, version control, monitoring, appeal, and explicit revalidation. Each control should be connected to a named failure mode and should identify an owner, evidence source, review interval, and escalation path.

Susskind et al. (2021) addresses "Neuro-Symbolic AI: An Emerging Class of AI Workloads and their Characterization." In this synthesis, that work is used as a source on implementation within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Chen et al. (2024) addresses "Exploring neuro-symbolic AI applications in geoscience: implications and future directions for mineral prediction." In this synthesis, that work is used as a source on governance within neuro-symbolic adaptation for autonomous systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for adaptive autonomous decision systems. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

A credible implementation plan also defines sunset and revalidation conditions. New data, software updates, institutional restructuring, population change, or shifts in incentives can invalidate previously reasonable assumptions. Monitoring should therefore test the validity of the evidence chain, not merely whether a service remains online or a headline metric remains stable.

6. Research Agenda

Future studies should preregister or otherwise predefine the intended decision, principal outcomes, exclusions, and analysis plan when feasible. They should report missingness, sensitivity analyses, negative or null findings, and the cost of errors to differently situated stakeholders. Reproducible artifacts should include sufficient metadata to reconstruct data provenance, model or analytical versions, parameter choices, and the sequence from evidence to conclusion.

Cross-disciplinary work needs a shared object of explanation rather than a collection of adjacent vocabularies. For neuro-symbolic

adaptation for autonomous systems, that shared object is the complete decision pathway. Technical, clinical, social, environmental, or economic expertise should each change the design or interpretation of the study. When evidence remains incomplete, the useful output is a bounded hypothesis, a transparent uncertainty statement, or a request for additional measurement.

7. Conclusion

The literature supports a disciplined approach to neuro-symbolic adaptation for autonomous systems: define the decision, preserve system boundaries, test consequential failure modes, connect uncertainty to action, and assign accountable control. The strongest conclusion is not that one method is universally superior. It is that credible use in adaptive autonomous decision systems requires an auditable link from source evidence to design choice, evaluation result, and governed decision.

Declarations

Ethics approval: Not applicable. This article synthesizes previously published literature and reports no new research involving humans, animals, human tissue, or identifiable sensitive data.

Consent to participate and consent for publication: Not applicable.

Funding: No external funding is reported for this commissioned synthesis.

Competing interests: The authoring group declares no competing interests for this commissioned synthesis.

AI-assisted writing: Generative AI supported drafting, source organization, and language editing. Accountable authors and editors must verify every claim, citation, and disclosure before publication.

Data, code, and materials availability: No new dataset or code was generated. All evidence is identified in the reference list.

Licence: This manuscript is prepared for publication under the Creative Commons Attribution 4.0 International licence (CC BY 4.0).

References

- Chandre, P., Mahalle, P., Shinde, G., Shendkar, B., & Kashid, S. (2025). Neuro-Symbolic AI: A Future of Tomorrow. *ASEAN Journal on Science and Technology for Development*, 42(2). <https://doi.org/10.61931/2224-9028.1620>
- Chen, W., Ma, X., Wang, Z., Li, W., Fan, C., Zhang, J., Que, X., & Li, C. (2024). Exploring neuro-symbolic AI applications in geoscience: implications and future directions for mineral prediction. *Earth Science Informatics*, 17(3), 1819-1835. <https://doi.org/10.1007/s12145-024-01278-7>
- Colelough, B. C., & Regli, W. (2025). Neuro-Symbolic AI in 2024: A Systematic Review. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2501.05435>
- Gaur, M., Gunaratna, K., Bhatt, S., & Sheth, A. (2022). Knowledge-Infused Learning: A Sweet Spot in Neuro-Symbolic AI. *IEEE Internet Computing*, 26(4), 5-11. <https://doi.org/10.1109/mic.2022.3179759>
- Hagos, D. H., & Rawat, D. B. (2024). Neuro-Symbolic AI for Military Applications. *IEEE Transactions on Artificial Intelligence*, 5(12), 6012-6026. <https://doi.org/10.1109/tai.2024.3444746>
- Hamilton, K., Nayak, A., Božić, B., & Longo, L. (2022). Is neuro-symbolic AI meeting its promises in natural language processing? A structured review. *Semantic Web*, 15(4), 1265-1306. <https://doi.org/10.3233/sw-223228>
- Li, Q., Liu, Y., Sun, S., Qin, Z., & Chu, F. (2024). Deep expert network: A unified method toward knowledge-informed fault diagnosis via fully interpretable neuro-symbolic AI. *Journal of Manufacturing Systems*, 77, 652-661. <https://doi.org/10.1016/j.jmsy.2024.10.007>
- Nawaz, U., Anees-ur-Rahaman, M., & Saeed, Z. (2025). A review of neuro-symbolic AI integrating reasoning and learning for advanced cognitive systems. *Intelligent Systems with Applications*, 26, 200541.

<https://doi.org/10.1016/j.iswa.2025.200541>

Prenosil, G. A., Weitzel, T. K., Bello, S. C., Mingels, C., Manzini, G., Meier, L. P., Shi, K. Y., Rominger, A., & Afshar-Oromieh, A. (2025). Neuro-symbolic AI for auditable cognitive information extraction from medical reports. *Communications Medicine*, 5(1), 491. <https://doi.org/10.1038/s43856-025-01194-x>

Susskind, Z., Arden, B., John, L. K., Stockton, P., & John, E. B. (2021). Neuro-Symbolic AI: An Emerging Class of AI Workloads and their Characterization. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2109.06133>