

Review Response

July 22, 2026

Dear NAI Editor and Reviewers:

We thank the reviewers for their second careful evaluation. Below, we address the two final points and summarize our revisions.

- **Energy Function, Symbolic Component, Symbolic Parameter Distinction.** We agree the purpose of our notations was not fully motivated, and we have further revised the paper to directly address this. First, regarding symbolic potentials, ψ , we have kept symbolic potentials as a defined object in Appendix C and further explained its purpose. Symbolic potentials organize the arguments of the symbolic component by the role they play in formulating the prediction program. This organization is what formally distinguishes the three modeling paradigms and is not explicit in the signature of g_{sy} or E . Further, thank you for the careful eye, we removed the reference to ψ in the figure given the definition of ψ was moved to the appendix in the latest revision. Symbolic potentials ψ appear only where they are needed in Appendix C to formalize the modeling paradigms and Appendix D to express other NeSy approaches as NeSy-EBMs.

Second, regarding E and g_{sy} , though the energy function E is defined as a composition of the symbolic component g_{sy} and the neural component $\mathbf{g}_{nn}$ they are distinct objects with technical benefits. This distinction and motivation has been added after Definition 1 on page 11. The symbolic component g_{sy} is defined at every vector $\mathbf{u} \in \mathbb{R}^{d_{nn}}$, independent of whether $\mathbf{u}$ is realized by any neural component composed with it. The explicit $g_{sy}, \mathbf{g}_{nn}$ composition notation is necessary for the learning theory developed in Section 5 where gradients of value-functions with respect to the neural weights (Theorem 6) are obtained by perturbing the argument of g_{sy} at the realized neural output, a perturbation that cannot be expressed in terms of E alone. The energy function, E , abstracts away the neural output as an intermediate layer and succinctly expresses inference and learning objectives without relying on explicitly expressing the machinery of neural symbolic composition.

- **Learning Loss Separability** We thank the reviewer for the precise pointer and agree the specific referenced example (Example 6, Appendix D, page 66) deserves further clarification. In the example, the generalized mean aggregation defines the symbolic component. The empirical risk minimization framework requires separability across training samples, not among variables within a training sample. For this reason, when the generalized mean aggregation ranges over an entire training dataset, the dataset can be treated as a single training sample in our NeSy-EBM formulation of LTNs. In this case, the learning results of section 5 are still applicable. However, if a modeler wishes to train with mini batches of the entity set, $\mathcal{U}$, with an objective that does not decompose into a sum of losses, for instance $p \neq 1$ in the generalized mean aggregation function, this is indeed not within the general empirical risk minimization framework studied in Section 5. We updated the related statement on scope and empirical risk minimization in Section 5 page 24 and the discussion following Example 6 in Appendix D to clarify this distinction and the scope of the learning results.