

Introduction Large language models (LLMs) are increasingly used for information extraction, biomedical concept normalization, summarization, question answering, and decision support. Despite strong performance, they remain difficult to deploy reliably due to their output variability across repeated queries, disagreement across models, and instability under prompt perturbation. One approach to improve reliability is to use ensemble methods. This paper develops a hierarchical uncertainty framework that integrates Dempster–Shafer (DS) evidence theory with Interval Type-3 fuzzy sets to model structured, multi-layer uncertainty in LLM ensembles. Within-model variability 1 A single model may give different answers for the same repeated prompt. Between-model disagreement 2 Different LLMs encode different priors, training corpora, and reasoning behaviours. Prompt-level meta-uncertainty 3 Prompt phrasing alters the semantic framing and stability of the answer.	Evidence Representation from LLM Responses For a fixed model L_i, prompt P_r, let n_y, n_o, n_u be the counts of Yes, No, Unsure in k repeated trials. The DS mass function is: $m_{ir}(Y) = n_y/k, \quad m_{ir}(N) = n_o/k, \quad m_{ir}(U) = n_u/k$ This separates support, contradiction, and ignorance directly. Evidence across models is combined iteratively with Dempster's rule to obtain a fused mass function m^*_r. The fused belief and plausibility for proposition Y are: $\text{Bel}_r(Y) = m^*_{ir}(Y), \quad \text{Pl}_r(Y) = m^*_{ir}(Y) + m^*_{ir}(U)$ Prompt P_r produces a belief–plausibility interval equivalent to a Type-2 fuzzy membership: $I_r(Y) = [\text{Bel}_r(Y), \text{Pl}_r(Y)]$ <i>Stacking these intervals across R prompts creates the higher-order (Type-3) uncertainty object — the core of our framework.</i>	Experiments & Results Three local models were evaluated $M = 3$ LLMs: MedGemma, LLaMA 3.2, Mistral 3B. For each item, $R = 4$ prompt formulations $\times k = 10$ repeated queries were used. DS masses were fused per prompt, then aggregated with the Interval Type-3 operator. Table I: UMLS Concept Verification <table><tr><th>Method</th><th>Acc</th><th>Prec</th><th>Rec</th><th>F1</th><th>Cvg</th><th>Rej</th></tr><tr><td>MedGemma</td><td>0.90</td><td>0.83</td><td>1.00</td><td>0.90</td><td>1.00</td><td>0</td></tr><tr><td>LLaMA 3.2</td><td>0.43</td><td>0.46</td><td>0.82</td><td>0.59</td><td>1.00</td><td>0</td></tr><tr><td>Mistral 3B</td><td>0.98</td><td>0.99</td><td>0.97</td><td>0.98</td><td>1.00</td><td>0</td></tr><tr><td>DS+T3 (ours)</td><td>1.00</td><td>1.00</td><td>1.00</td><td>1.00</td><td>0.84</td><td>31</td></tr></table> Table II: Gene–Gene Interaction Verification <table><tr><th>Method</th><th>Acc</th><th>Prec</th><th>Rec</th><th>F1</th><th>Cvg</th><th>Rej</th></tr><tr><td>MedGemma</td><td>0.58</td><td>0.54</td><td>0.99</td><td>0.70</td><td>1.00</td><td>0</td></tr><tr><td>LLaMA 3.2</td><td>0.49</td><td>0.49</td><td>1.00</td><td>0.66</td><td>1.00</td><td>0</td></tr><tr><td>Mistral 3B</td><td>0.49</td><td>0.49</td><td>1.00</td><td>0.66</td><td>1.00</td><td>0</td></tr><tr><td>DS+T3 (ours)</td><td>1.00</td><td>1.00</td><td>1.00</td><td>1.00</td><td>0.07</td><td>182</td></tr></table>	Method	Acc	Prec	Rec	F1	Cvg	Rej	MedGemma	0.90	0.83	1.00	0.90	1.00	0	LLaMA 3.2	0.43	0.46	0.82	0.59	1.00	0	Mistral 3B	0.98	0.99	0.97	0.98	1.00	0	DS+T3 (ours)	1.00	1.00	1.00	1.00	0.84	31	Method	Acc	Prec	Rec	F1	Cvg	Rej	MedGemma	0.58	0.54	0.99	0.70	1.00	0	LLaMA 3.2	0.49	0.49	1.00	0.66	1.00	0	Mistral 3B	0.49	0.49	1.00	0.66	1.00	0	DS+T3 (ours)	1.00	1.00	1.00	1.00	0.07	182
Method	Acc	Prec	Rec	F1	Cvg	Rej																																																																		
MedGemma	0.90	0.83	1.00	0.90	1.00	0																																																																		
LLaMA 3.2	0.43	0.46	0.82	0.59	1.00	0																																																																		
Mistral 3B	0.98	0.99	0.97	0.98	1.00	0																																																																		
DS+T3 (ours)	1.00	1.00	1.00	1.00	0.84	31																																																																		
Method	Acc	Prec	Rec	F1	Cvg	Rej																																																																		
MedGemma	0.58	0.54	0.99	0.70	1.00	0																																																																		
LLaMA 3.2	0.49	0.49	1.00	0.66	1.00	0																																																																		
Mistral 3B	0.49	0.49	1.00	0.66	1.00	0																																																																		
DS+T3 (ours)	1.00	1.00	1.00	1.00	0.07	182																																																																		
Problem Formulation Consider a proposition H (e.g., "term t exists in UMLS" or "gene A up-regulates gene B"). Let L_i, $i = 1 \dots M$, denote an ensemble of M LLMs; let P_r, $r = 1 \dots R$, denote prompt R configurations. Under each model–prompt condition, the proposition is queried k times, producing labels from $A = \{\text{Yes, No, Unsure}\}$ Goal: Produce a structured uncertainty representation for the truth of H that captures: 1. Within-model response variation over k repeated trials 2. Between-model disagreement under a fixed prompt P_r 3. Between-prompt meta-uncertainty about the resulting belief interval	Type-3 Fusion Operator Definition 1 (Weighted Type-3 Fusion Operator) Given prompt-level intervals $I_r(Y) = [\text{Bel}_r(Y), \text{Pl}_r(Y)]$ and weights w_r: $I''(Y) = [\sum w_r \text{Bel}_r(Y), \sum w_r \text{Pl}_r(Y)]$ Reliability Weighting Width: $W_r = \text{Pl}_r(Y) - \text{Bel}_r(Y)$ Reliability: $R_r = 1 - W_r$ Norm. weight: $w_r = R_r / \sum R_r$ Narrower intervals (more stable prompts) receive higher weight. A case is accepted when the Type-3 interval is sufficiently narrow; otherwise rejected (abstain).	Cross-Dataset Interpretation <table><tr><th>UMLS Verification</th><th>Gene–Gene Interaction</th></tr><tr><td>Coverage: 84.5% 31 abstentions</td><td>Coverage: 7.6% 182 abstentions</td></tr><tr><td>High-coverage reliability enhancer</td><td>High-precision uncertainty filter</td></tr><tr><td>Framework accepts most items, abstains on prompt-sensitive or conflict-heavy cases.</td><td>All models over-predict. Framework becomes extremely conservative to preserve perfect precision.</td></tr></table>	UMLS Verification	Gene–Gene Interaction	Coverage: 84.5% 31 abstentions	Coverage: 7.6% 182 abstentions	High-coverage reliability enhancer	High-precision uncertainty filter	Framework accepts most items, abstains on prompt-sensitive or conflict-heavy cases.	All models over-predict. Framework becomes extremely conservative to preserve perfect precision.																																																														
UMLS Verification	Gene–Gene Interaction																																																																							
Coverage: 84.5% 31 abstentions	Coverage: 7.6% 182 abstentions																																																																							
High-coverage reliability enhancer	High-precision uncertainty filter																																																																							
Framework accepts most items, abstains on prompt-sensitive or conflict-heavy cases.	All models over-predict. Framework becomes extremely conservative to preserve perfect precision.																																																																							
Framework Overview LLM Responses k queries per model–prompt pair $\rightarrow \{n_y, n_o, n_u\}$ DS Mass Functions $m_{ir}(Y/N/U)$ per model i, prompt r DS Fusion (per prompt) Combine across M models $\rightarrow I_r = [\text{Bel}_r, \text{Pl}_r]$ Reliability Weights $W_r = \text{interval width} \rightarrow R_r = 1 - W_r \rightarrow w_r$ Type-3 Aggregation Weighted avg across R prompts $\rightarrow I''(Y)$ Accept / Reject Narrow $I'' \rightarrow \text{Accept}$ Wide $I'' \rightarrow \text{Reject}$	Conclusions Hierarchical fusion: DS mass functions capture within-model and between-model uncertainty; Type-3 intervals capture between-prompt meta-uncertainty. Adaptive behavior: The framework operates as a high-coverage enhancer (UMLS) or high-precision filter (gene–gene) depending on task difficulty — automatically. Perfect classification: Both tasks achieve Acc/Prec/Rec/F1 = 1.000 on accepted items, at the cost of selective abstention. Key insight: Higher-order evidence fusion should be judged by the appropriateness of its confidence structure, not solely by raw accuracy.	References [1] Banerjee et al. "LLMs Will Always Hallucinate." ArXiv 2409.05746, 2024. [2] G. Shafer, A Mathematical Theory of Evidence. Princeton Univ. Press, 1976. [3] J. M. Mendel, Uncertain Rule-Based Fuzzy Systems. Prentice Hall, 2001. [4] Mendel & John, "Type-2 fuzzy sets made simple," IEEE Trans. Fuzzy Syst., 2002. [7] P. Smets, "Combination of evidence in the TBM," IEEE TPAMI, 1990. [8] Sellergren et al., "MedGemma Technical Report," arXiv 2507.05201, 2025. [9] Grattafiori et al., "The Llama 3 Herd of Models," arXiv 2407.21783, 2024. [10] Jiang et al., "Mistral 7B," arXiv 2310.06825, 2023.																																																																						