

STRICTA: Structured Reasoning In Critical Text Assessment for Peer Review and Beyond



Nils Dycke, Matej Zečević, Ilia Kuznetsov, Beatrix Suess, Kristian Kersting, Iryna Gurevych

UKP Lab, AIML Lab, Synthetic Biology Lab
at Technical University of Darmstadt



1 Contributions

3-step framework for modeling structured reasoning with uncertainty

Routed in causality using Pearl's SCMs

Expert multi-modal reasoning dataset with *>4k reasoning steps* on *>20 biomedical papers*

Insights in human and LLM assessment behavior e.g. background knowledge is a major source of variance

2 Motivation

I observe X in the text... Overall, the quality... Due to X and Y, I think... hidden reasoning process

- Critical text assessment is ubiquitous: peer review, essay grading, fact checking
- Unveiling the black box
- Structured reasoning because...
 - + fine-granular evaluation
 - + human behavior analysis
 - + human-LLM collaboration

3 STRICTA

I. Design Structured Causal Model

Discover reasoning steps and links e.g. by expert interviews

$M = (U, V, F, P_M)$

II. Populate SCM

Collect human reasoning as annotations

III. Analysis + LLMs

Study causal reasoning Test LLMs

4 Biomedical Paper Assessment

Biomedical Assessment Workflow based on interviews

paper → extract → infer → final verdict

Expert-generated Dataset

Dataset Statistics	
# step answers	4371
# workflow ans.	93
# papers	22
Krippendorff's α	0.42

Background knowledge invokes dissimilarity

5 Causal Analysis + LLM Experiments

Which factors shape the assessment most?

Steps' Average Causal Effect on Verdict	
Consistency of conclusions to RQs	0.37
Relevance of conclusions to science	0.20
Clarity of the paper	0.20

► Humans show a positive bias towards well-written papers.
... and a lot more analysis including counterfactuals etc.

Can LLMs explain human's final verdicts?

LLMs perform best on extraction and weight factors differently

	BERT-F1 ↑	SummaC ↑	TRUE ↑	F1 ↑
human* [†]	0.799±.06	-0.151±.30	0.151±.27	0.801
Llama3 Prg [†]	0.752±.10	-0.274±.36	0.098±.30	0.170
Mixtral Prg [†]	0.761±.09	-0.149±.26	0.120±.32	0.559
GPT3.5t Prg [†]	0.759±.10	-0.178±.35	0.163±.37	0.531
GPT4o Prg [†]	0.780±.07	-0.186±.30	0.139±.35	0.720
majority ^{io}				0.854
human ^{io}	0.799±.06	-0.158±.29	0.150±.27	0.801
Llama3 ^{io}	0.786±.07	-0.141±.30	0.145±.35	0.657
Mixtral ^{io}	0.794±.07	-0.077±.27	0.161±.37	0.822
GPT3.5t ^{io}	0.805±.07	-0.125±.34	0.214±.41	0.789
GPT4o ^{io}	0.795±.07	-0.094±.28	0.194±.40	0.876
GPT4o [§]	0.776±.07	-0.154±.28	0.188±.39	0.828