

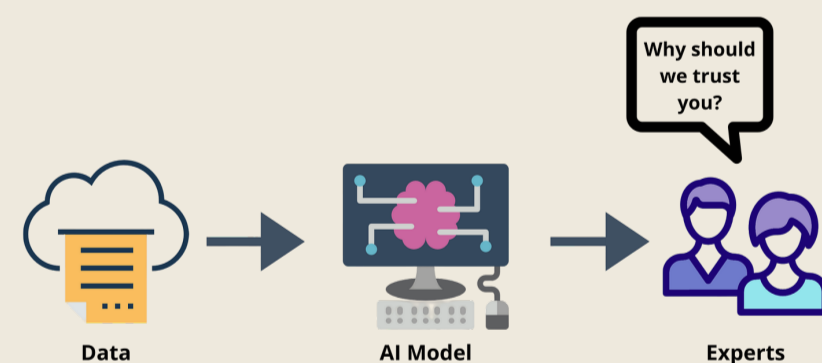
Borislav Kolev Marinov
B.K.Marinov-1@student.tudelft.nl

Pradeep Murukannaiah,
Shubhalaxmi Mukherjee



Evaluating Faithfulness of LLM Generated Explanations for Claims: Are Current Metrics Effective?

Analysing the Capabilities of Automatic Metrics to Represent the Difference in Faithfulness Between Explanations



01. Introduction

- Data Access → Claims and evidence flow into LLMs
- LLM → Produces outputs and explanations
 - Can sound convincing but be misleading
- Experts ask → “Why should we trust this explanation?” → evaluation required
- Our Goal
 - Evaluate if existing metrics (G-Eval, FactCC, UniEval, QAGs) truly capture faithfulness in fact-checking explanations

03. Methodology

Datasets of Claims and Evidence:

- QuanTemp with Journalist explanations from PoliFact
- HoVer with evidence from Wikipedia articles

Step 1: Generate Explanations

Task: You are a fact-checking assistant. Your task is to evaluate the truthfulness of a given claim based on provided evidence.
Label: <True/False/Half-True>
Justification: <your explanation here>
Claim: {claim}, Evidence: {evidence}

Ollama Framework Llama 3.1

LLM Output
Label: True;
Justification: The claim is true as ...

Step 2: Evaluate Faithfulness

Obtain Faithfulness Score with G-Eval, FactCC, QAGs, UniEval
Determine Correlation Between Metrics and Similarity of Generated to Politifact Explanations
Compare Scores to Prediction Accuracy

Step 3: Targeted Analysis via Sentence Masking

Generated Explanation
1 unrelated sentence
2 unrelated sentences
3 unrelated sentences
Faithfulness Score Change

Generated Explanation
1 unsupported
2 unsupported
3 unsupported
Faithfulness Score Change



02. Research Question Terminology

How well do current evaluation metrics reflect the faithfulness of LLM-generated fact-checking explanations compared to journalist-written ones?

Faithfulness Metrics

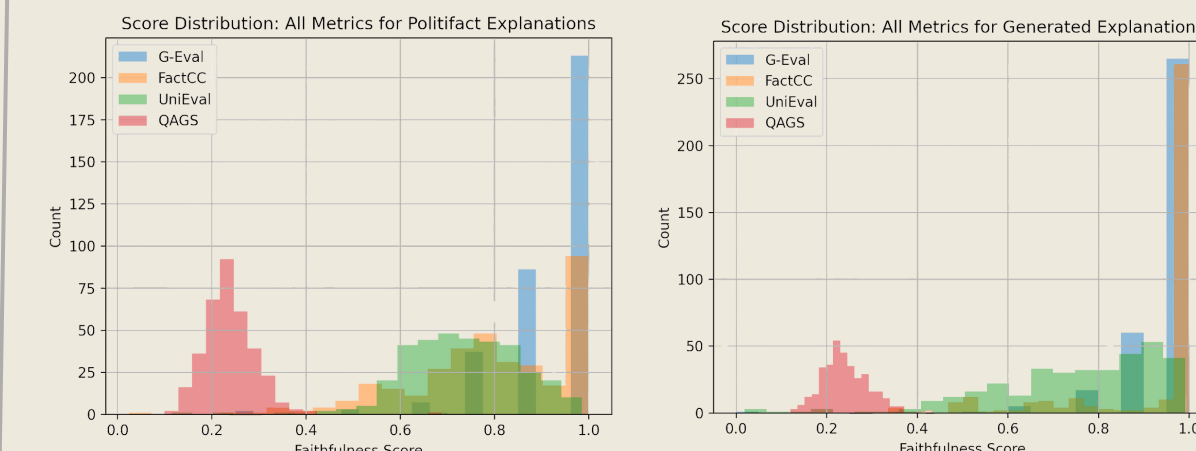
- G-Eval — LLM grades explanations using reasoning prompts with form-filling and Chain-of-Thought
- FactCC — Checks consistency with evidence using BERT
- QAGs — Converts explanation into questions, answers them with evidence
- UniEval — LLM uses boolean QA prompts to judge factual correctness
- Each metric gives a score between 0 (unfaithful) and 1 (faithful)

- Claim — Statement to be verified
- Evidence — Factual context for claim
- Label — Verdict (True / Half-True / False)
- Faithfulness — Explanation aligns only with evidence (no extra or false info)

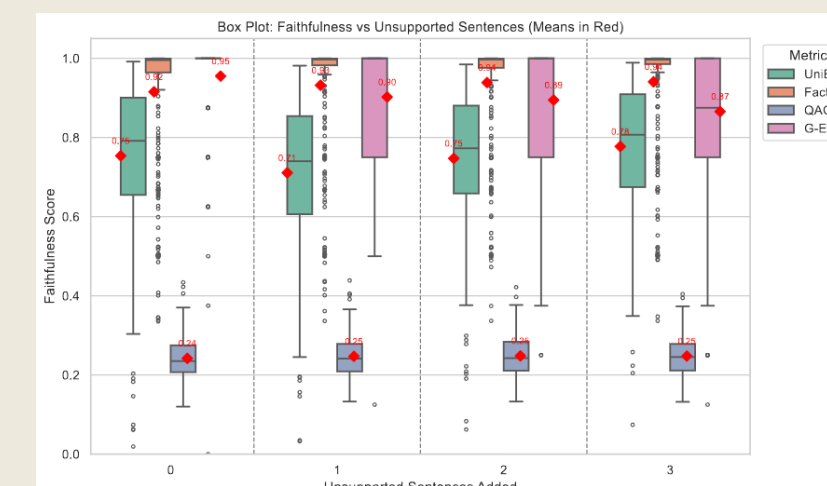
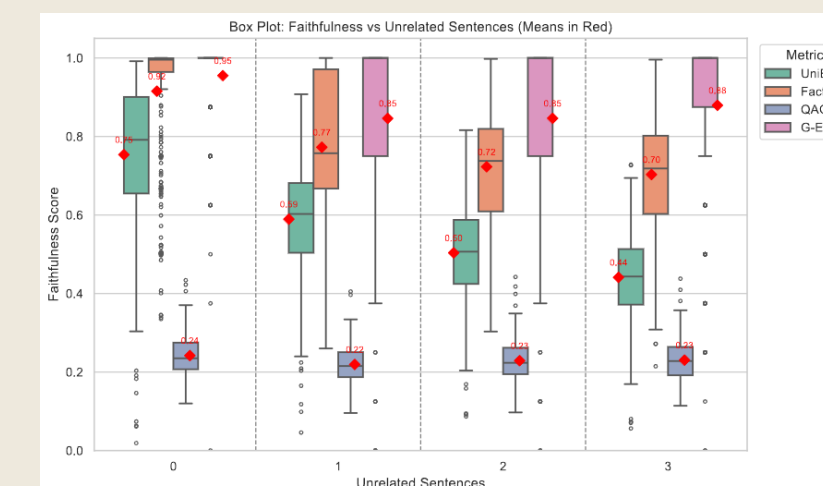
05. Analysis

- Metric Correlation with Semantic Similarity
 - Weak correlations overall ($p \leq 0.23$)
 - FactCC & QAGs: Best (though limited) alignment with expert explanations
 - UniEval & G-Eval: Low or even negative correlations
- Score Bias
 - All metrics assign slightly higher scores to incorrect predictions
 - G-Eval shows the strongest bias favoring LLM outputs
- Perturbation Sensitivity
 - Unrelated Sentences: FactCC & UniEval show clear score drops
 - Unsupported Sentences: Metrics often fail to detect hallucinations, though some score reduction observed

04. Results



Metric Correlations: Metrics favor LLMs – Normal distribution of scores for expert, Normal distribution skewed towards higher scores for generated



Targeted Tests: Some metrics detect hallucinations (FactCC, UniEval); others (G-Eval) fail to penalize unfaithful content consistently, while QAGs is robust to noise, but low in scores.

06. Conclusion

- Existing metrics are unreliable and inconsistent
- Different metrics perform better under the experiments, but no metric consistently reflects faithfulness
- Have some positive aspects that show latent potential
- Further Work**
 - Use more capable LLMs (e.g., GPT-4) for generation and evaluation.
 - Fine-tune metrics on the actual dataset (e.g., PoliFact QuanTemp) to improve alignment.
 - Test new perturbations — insert noise or hallucinations into the middle of explanations.

