

LLMs as Span Annotators: A Comparative Study of LLMs and Humans

Zdeněk Kasner¹, Vilém Zouhar², Patrícia Schmidtová¹, Ivan Kartáč¹, Kristýna Onderková¹,
Ondřej Plátek¹, Dimitra Gkatzia³, Saad Mahamood⁴, Ondřej Dušek¹, Simone Balloccu⁵

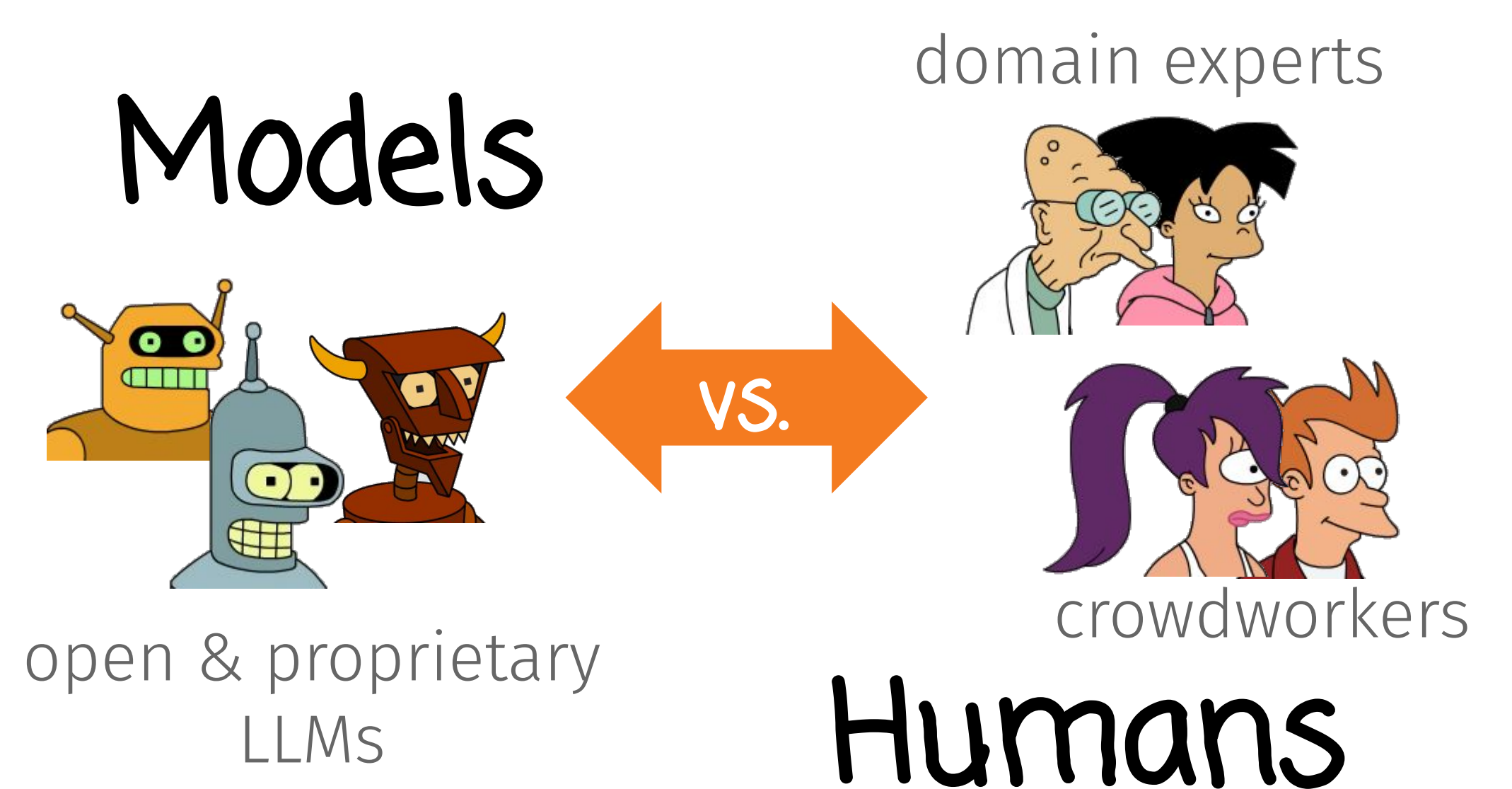


LLMs can anotate text spans as well as skilled crowdworkers, but human experts are better.

OVERGENERALIZATION

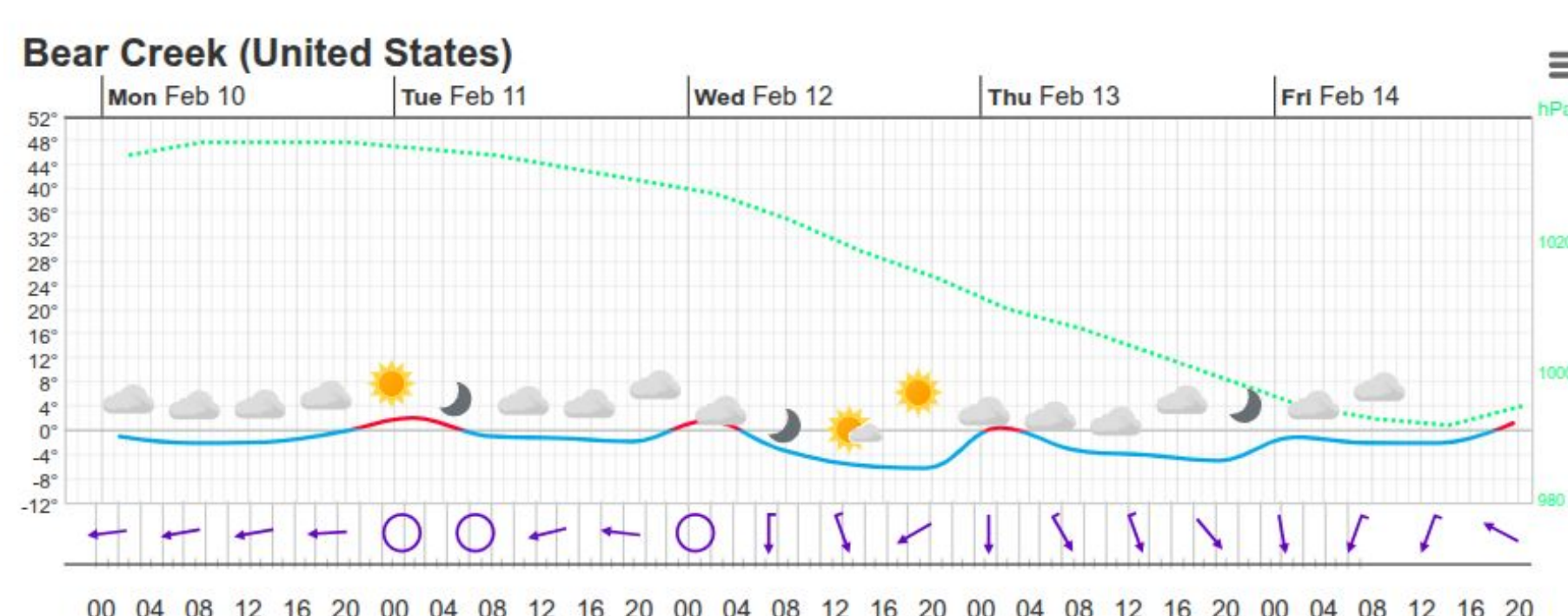
Task: Given a text & list of categories, **locate and categorize the spans** in the text.

Task	Description	Input X	Text Y with annotations A (category, span, reason)
D2T-Eval	Detecting errors in data-to-text. Custom annotation from Prolific.	Mon Tue Wed   	Skies will be mostly clear , but winds will remain strong . <i>Rain on Mon & Wed</i> <i>Wind speed data is missing.</i>
MT-Eval	Detecting errors in machine translation. WMT 2024: English → 9 langs.	Der schnelle braune Fuchs springt über den faulen Hund.	The quick brown fox jump over the lazy fox . <i>Third person singular</i> <i>'Hund' translates to 'dog'</i>
Propaganda	Detecting propaganda techniques. Da San Martino et al. (2019): various online sources	∅	Study Finds That Driving Car Is More Efficient than Biking <i>Appeal to a 'study'</i>



Research question: Can we reliably **automate span annotation with LLMs**?

Input: weather data (JSON for models, visual chart for humans) & generated forecast



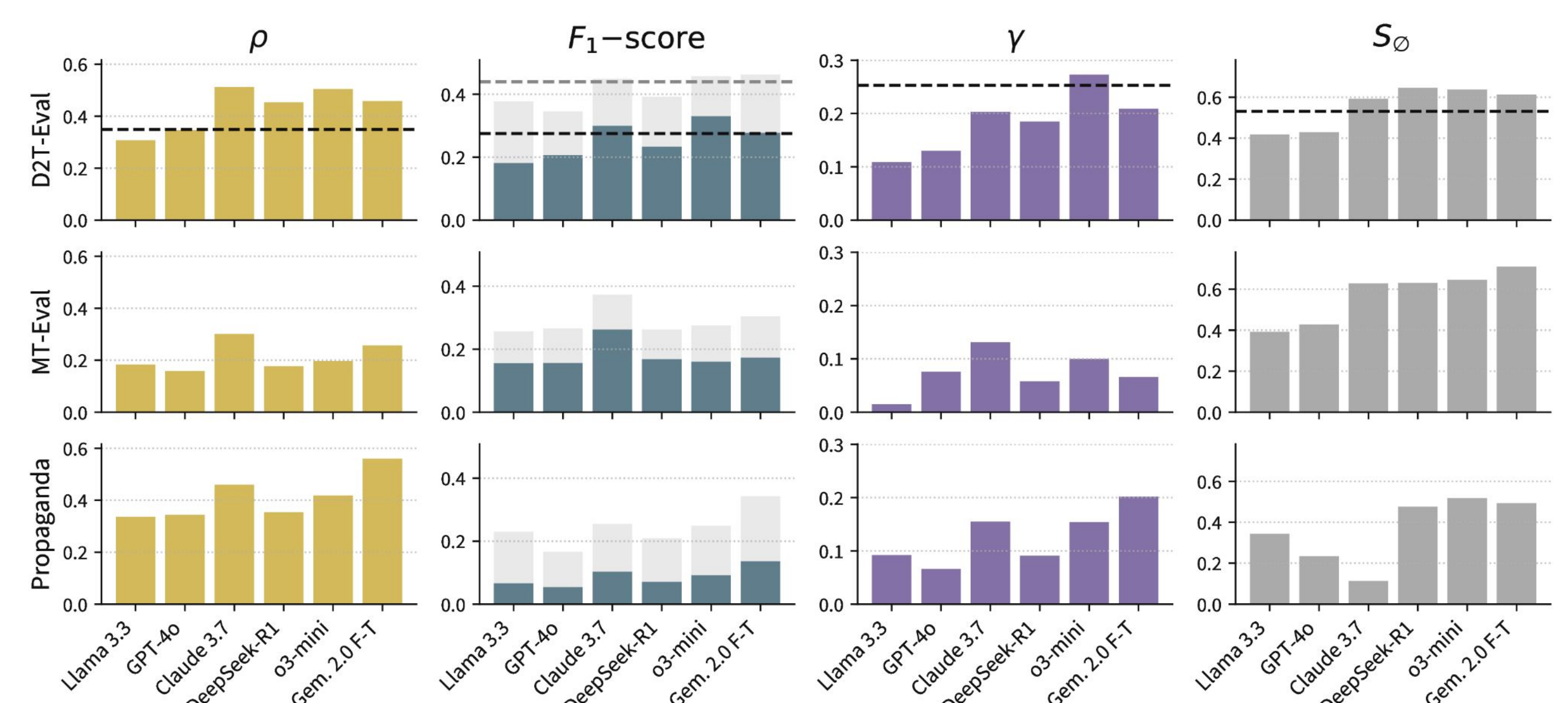
Example

human	Claude 3.7
Bear Creek will experience mostly overcast skies for the next five days, with a chance of occasional showers on Wednesday and Thursday . Temperatures will remain relatively cool, hovering around freezing point or slightly below. Winds will be light to moderate throughout the period, with some gusts expected on Friday . The overall weather pattern suggests a generally cloudy and potentially wet week .	Bear Creek will experience mostly overcast skies for the next five days, with a chance of occasional showers on Wednesday and Thursday . Temperatures will remain relatively cool, hovering around freezing point or slightly below. Winds will be light to moderate throughout the period, with some gusts expected on Friday . The overall weather pattern suggests a generally cloudy and potentially wet week .

Output: errors in the generated forecast
(**Contradictory**, **Misleading**)

- **Pearson p** : span counts
 - **γ score** (Mathet et al., 2015)
 - **F1** (w/overlaps)
 - **S_ϕ -score**: ex. with no spans
- position & category

How well are the spans **aligned**?



Evaluation

How **efficient** is each approach?

Are the spans **correct**?

Model	correct	partially correct	wrong category	incorrect	undecidable
Llama 3.3	7	2	8	9	2
GPT-4o	5	2	10	9	3
Claude 3.7	7	2	3	6	9
DeepSeek	11	2	4	8	6
o3-mini	10	3	2	3	15
Gemini 2 F-T	12	4	7	3	2

Annotator	Cost \$/1k out	Time s/out
crowdworkers	500	129.1
Llama 3.3 70B	-	21.6
DeepSeek-R1 70B	-	227.5
Claude 3.7 Sonnet	10.5	9.8
o3-mini	3.6	21.8

Low-to-moderate agreement
→ LLMs **cannot** replace human annotators straight away.

However, depends on the **domain and annotator expertise**.

Takeaways

In some domains, LLMs **can replace skilled annotators**.

LLMs are also more **cost- and time-efficient** than human annotators.

Paper & data explorer available at llm-span-annotators.github.io

Presented at Multilingual and Multicultural Evaluation (MME) Workshop, EMNLP 2026, Rabat, Morocco

Funded by the European Union (ERC, NG-NLG, 101039303). Supported by the National Recovery Plan funded project MPO 60273/24/21300/21000 CEDMO 2.0 NPO and the Charles University Research Centre program No.-24/SSH/009. It used resources of the LINDAT/CLARIAH-CZ Research Infrastructure (Czech Ministry of Education, Youth, and Sports project No. LM2018101).

