

LLM-as-a-Judgeにおけるリッカート評価とランキング評価の比較分析

松田 翔汰¹⁾, 柴田 圭悟¹⁾, 津田 純花¹⁾, 有川 透矢¹⁾, Kwek Yong Sing¹⁾, 池田 航¹⁾, 齋藤 由佳¹⁾, 種口 暁人¹⁾, 松田 陵佑¹⁾, 赤間 怜奈^{1),2),3)}, 鈴木 潤^{1),2),4)}

¹⁾東北大学, ²⁾理化学研究所, ³⁾国立国語研究所, ⁴⁾国立情報学研究所 LLMC

is-failab-research@grp.tohoku.ac.jp



概要

- 同一ベンチマークをリッカート・フルランク・ペアワイズで評価
- ペアワイズでリッカートと同等の人手との相関を得るにはより多くの件数を評価する必要性
→ペアワイズによってリッカート評価を代替し、コストを削減することは難しい
- GPT-4oとGPT-5.5でフルランクにおける人手評価との相関に大きな差
→同じベンチマークでも評価形式を変えることで分解能を変えられる

ベンチマークの評価形式

リッカート評価 (ベースライン)

- 各観点における点数の平均値
- 人手評価とLLM評価の相関係数で性能を評価 (Spearman's ρ を使用)



A を4観点で点数付けて
B を4観点で点数付けて
C を4観点で点数付けて

SummEval (要約の評価) [Fabbri+TACL'21]

	関連性	一貫性	流暢さ	整合性	平均点	並び替え	人手
A	3	2	5	2	3.0	1 C	C
B	3	4	4	3	4.0	2 B	B
C	4	5	5	4	4.5	3 A	A

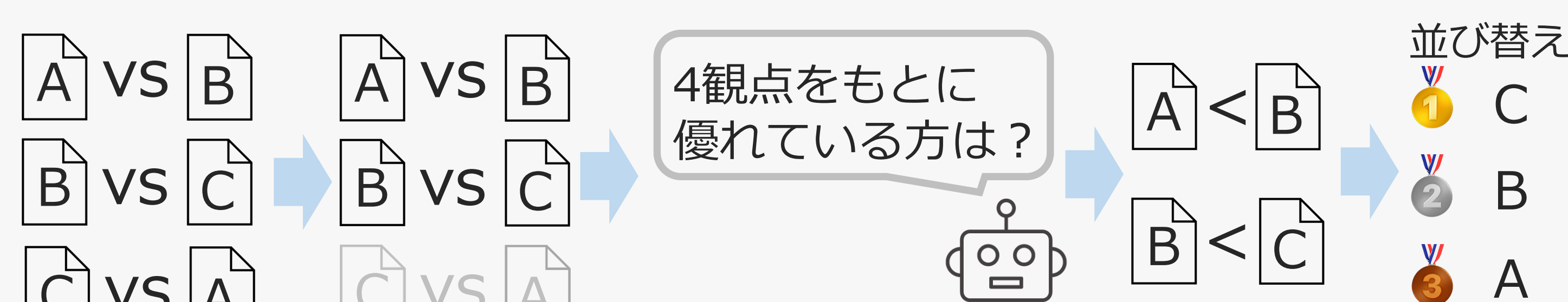
相関係数で比較

比較手法

同一ベンチマークを異なる評価形式で評価

Q1. ベンチマークの評価コストの削減は可能？

評価難易度：ペアワイズ評価 < リッカート評価



ペア $N C_2$ n サンプル

- リッカート評価の評価件数 N ($=100$)
- ペアワイズ評価の評価件数 n
例1) $n=1$ だとコスト大幅削減だが、精度が悪い
例2) 全部サンプルすると ($n=100$ $C_2=4950$)
精度は高いが、リッカートより件数が多い

性能が飽和するサンプル率の評価件数で比較

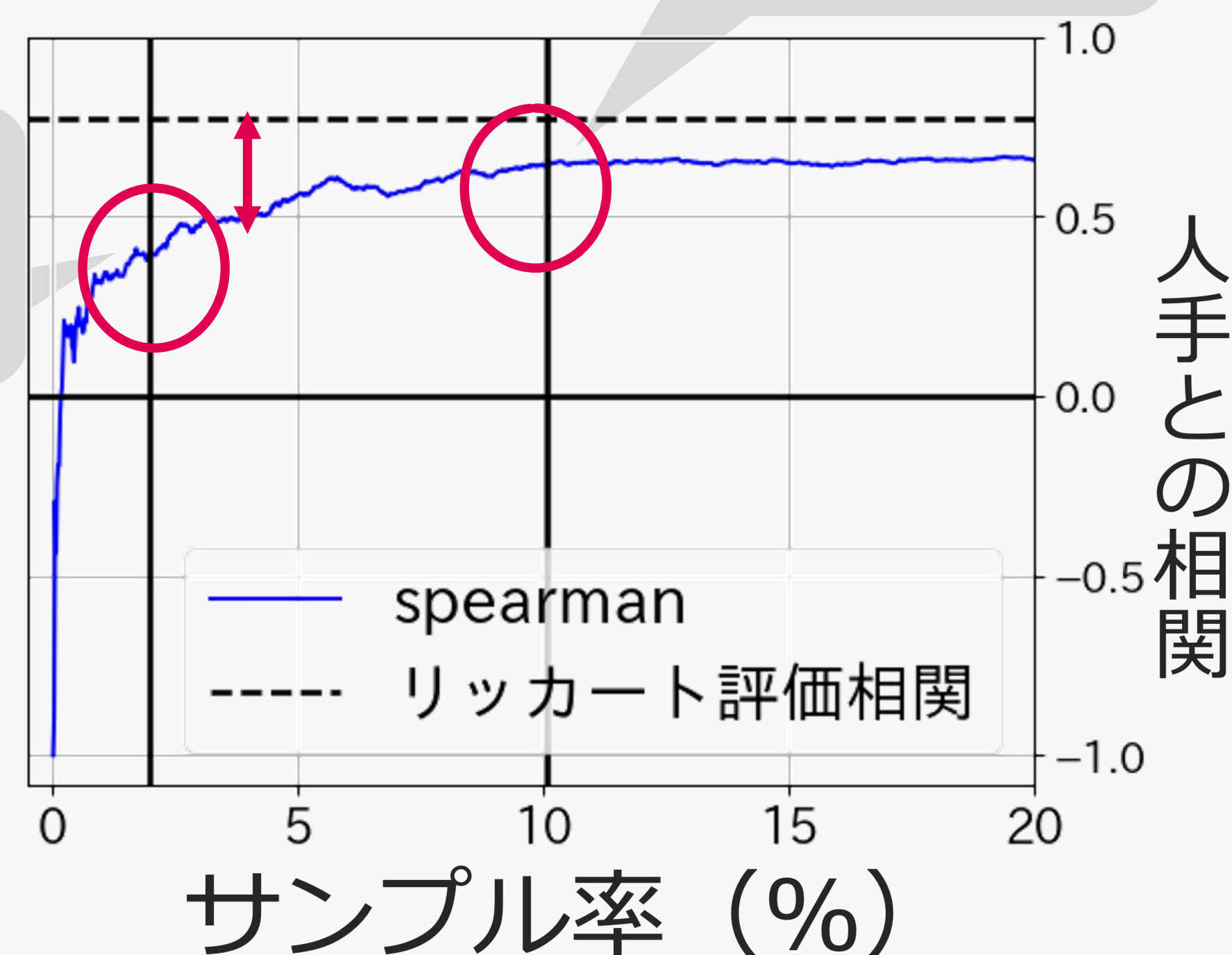
実験結果

ペアワイズ評価

サンプル率と人手との相関の関係

500ペアで
相関が飽和

100ペアでは
上昇途中



リッカート評価件数 < ペアワイズ飽和ペア数
100件 < 500ペア

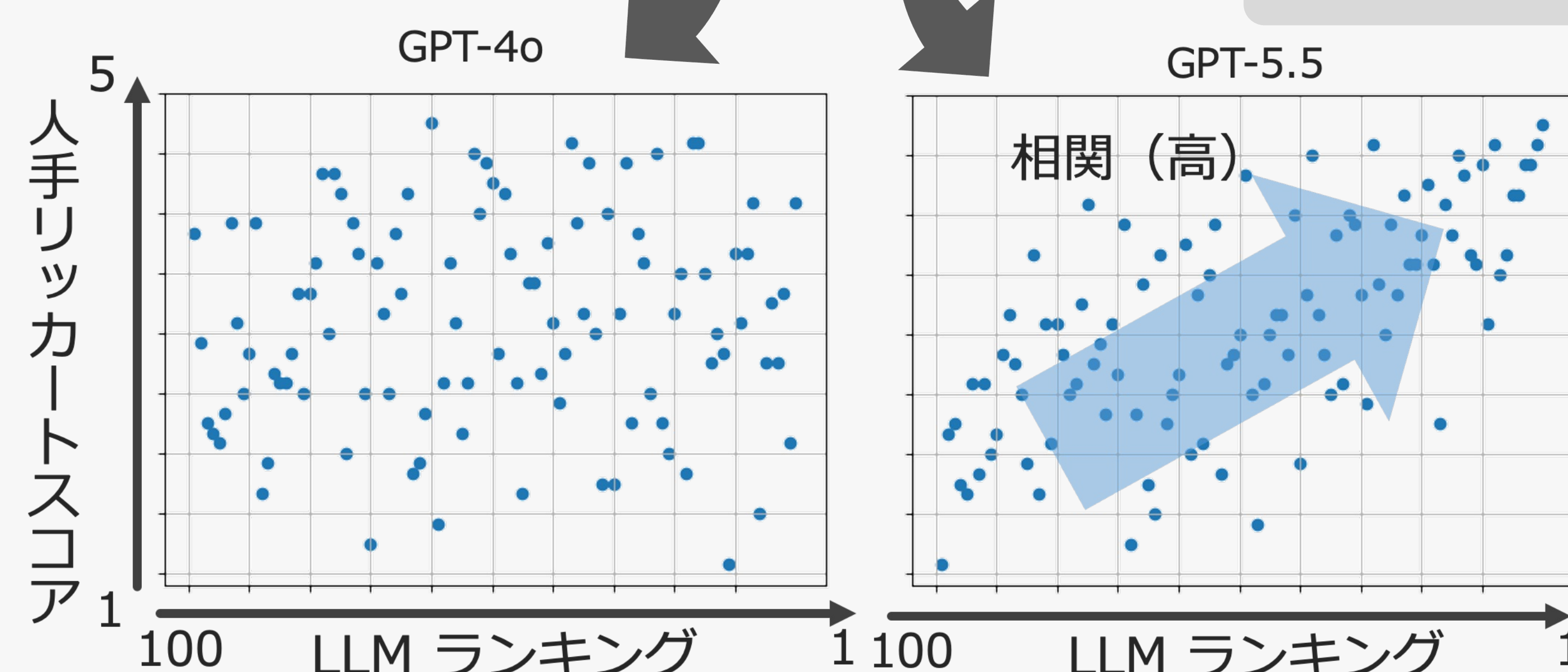
！ペアワイズによるコスト削減は難しい

フルランク評価

差が見えない

評価形式	GPT-4o	GPT-5.5	相関差
リッカート	0.77	0.72	0.05
フルランク	0.095	0.66	0.565

差が見える



！リッカートで見られない性能差をフルランクで検出可能