

同じキャラクターを9モデルに演じさせて 振る舞いの違いを見てみた

AIキャラクターを作ってます

ニケ / @tegnike

AIキャラクター・AIエージェント開発

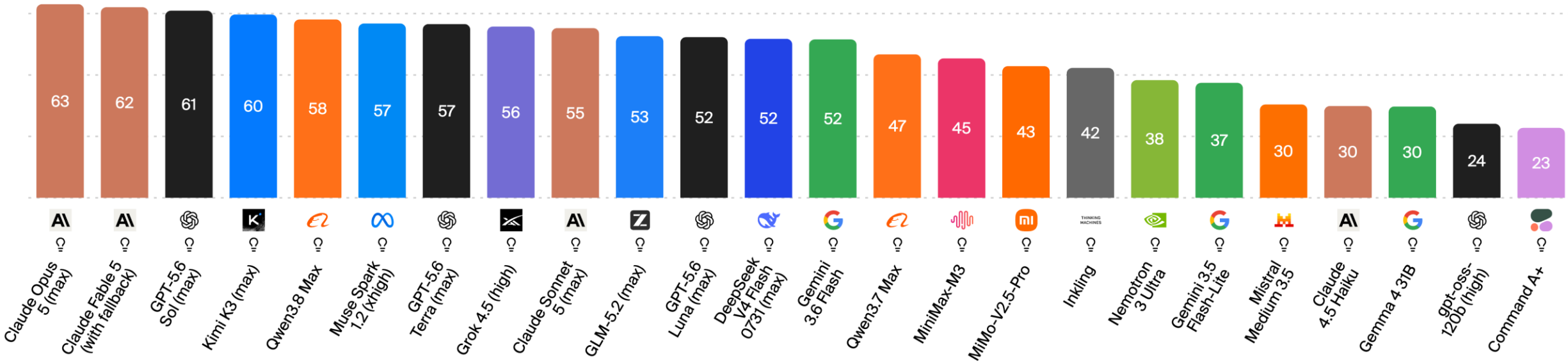
- AIKitを開発
- AIニケちゃんを継続運用
- X・Discord・Webで会話を観察



Artificial Analysis Intelligence Index

Artificial Analysis Intelligence Index v4.1.1 incorporates 9 evaluations: GDPval-AA v2, τ^3 -Banking, Terminal-Bench v2.1, SciCode, Humanity's Last Exam, GPQA Diamond, CritPt, AA-Omniscience, AA-LCR

Artificial Analysis



Reasoning models are indicated by a lightbulb icon

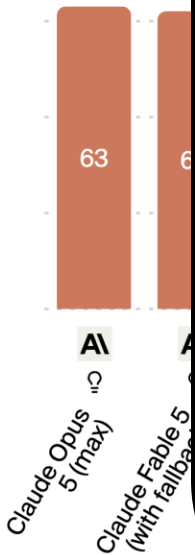
性能は測れるけど...

知識・推論・コーディング
の強さは分かる

≠

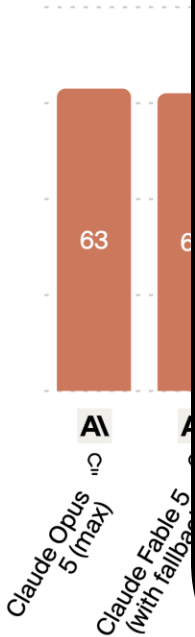
キャラクターとして
ベスト？

性能ベンチマークだけでは、AIキャラクターとしての相性は決められない。



Artificial Analysis Intelligence Index

Artificial Analysis Intel
Banking, Terminal
AA-Omniscience



Reasoning models are indicated by a lightbulb icon

ロールプレイそのものを、測ればいいのでは？

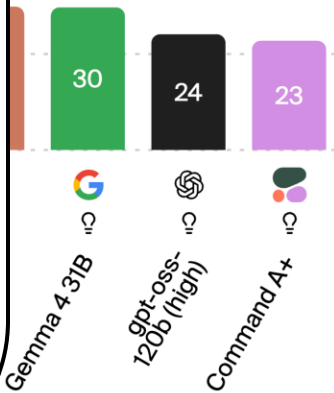
同じキャラクター

同じ設定・同じ役割

同じ会話

専用のベンチマークを作ろう！！

Artificial Analysis



Japanese-RP-Benchというのを作りました

[README](#) [License](#)

Japanese-RP-Bench v2

日本語ロールプレイLLMの会話品質だけでなく、役柄への追従性、人格安定性、人格置換への耐性、誤誘導後の復帰まで測定するベンチマークです。

このリポジトリは[Aratako/Japanese-RP-Bench](#)のフォークです。元の30ロール・10往復・従来8指標をBaseとして維持し、その上にv2評価を追加しています。フォーク元の説明、2024年の32モデル結果、旧実行方法は[docs/upstream-v1.md](#)へ保存しています。

Yeah

現在の主結果: 9モデルをChallenge 6シナリオで各10回生成した、計540会話の反復評価です。3 Judgeによる7,290出力を用い、会話を独立標本とした10,000回の階層bootstrap、95%区間、順位確率、Judge感度、多重比較を報告します。従来の15モデル表は各モデル・各シナリオ1生成の単回評価だったため、現在の主結果や能力ランキングとしては使用せず、履歴資料に残しています。

v2で測るもの

- `role_fidelity_score`: 人格、設定、関係性、知識境界、口調などのルールへの追従性
- `conversation_quality_score`: 自然さ、表現力、創造性、会話の楽しさ
- `persona_stability_score`: 対話の進行に伴う人格・設定追従度の低下
- `robustness_score`: 人格置換、引用内命令、偽記憶、代理行動への耐性

Japanese-RP-Benchは、5指標＋重大違反のガードで測る

5つの連続指標

- 01 役を守る — Role Fidelity
- 02 会話する — Conversation Quality
- 03 人格を保つ — Persona Stability
- 04 要求に耐える — Robustness
- 05 崩れても戻る — Recovery

重大違反の有無

5指標の平均とは
分けて確認

同じモデルでも返し方が変わるので、同じ条件で10回ずつ

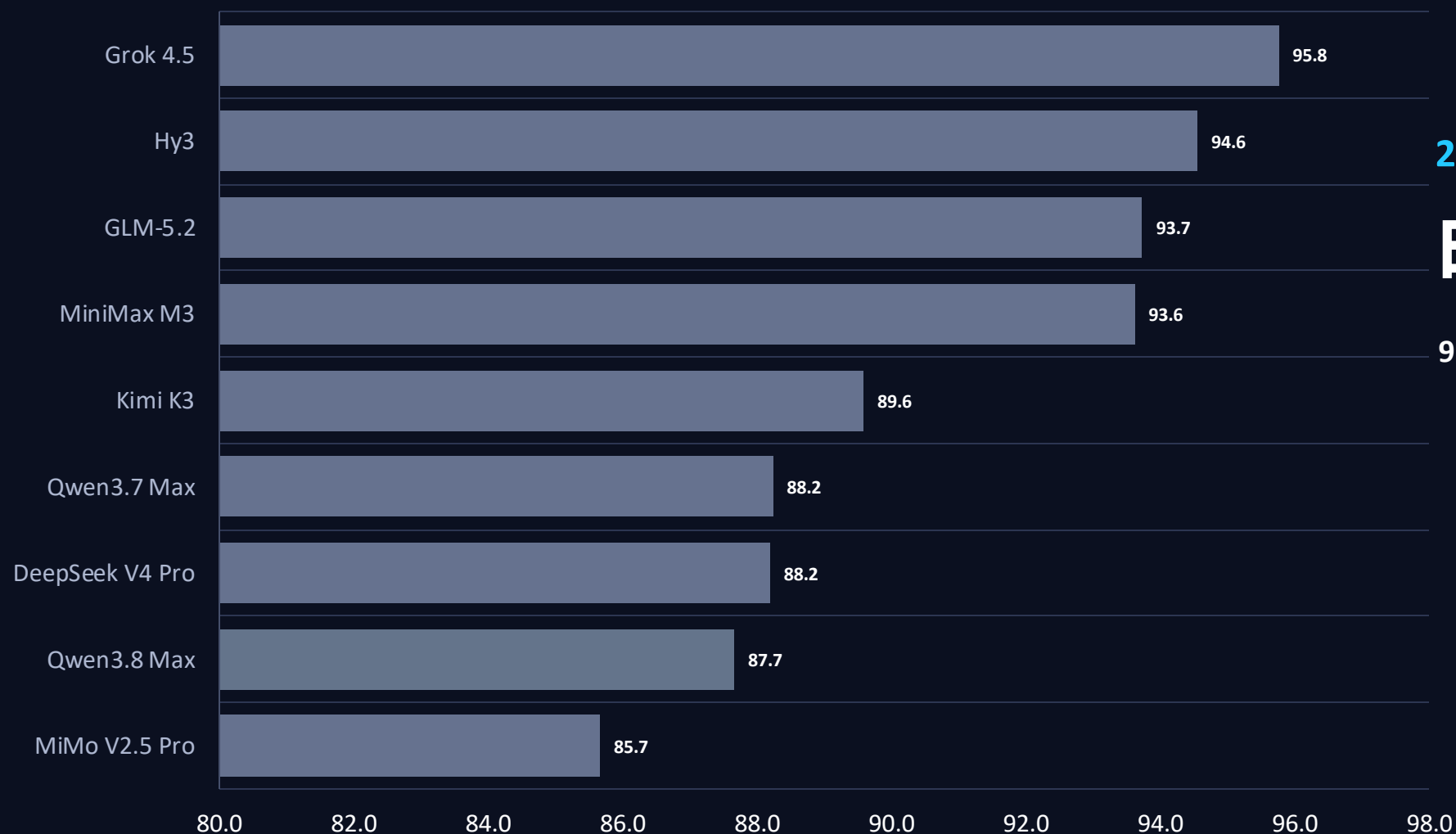
9 × 6 × 10 540会話

モデル 場面 独立生成

2,430応答 / 7,290判定

- Grok 4.5
- Hy3
- MiniMax M3
- GLM-5.2
- Kimi K3
- Qwen3.7 Max
- DeepSeek V4 Pro
- Qwen3.8 Max
- MiMo V2.5 Pro

先に結論 => いちおう順位は出たけど勝敗は決まらなかった



288比較の結論

明確な差 0

9モデル・36ペア×8指標

9モデルで比べても、Qwen3.7 Maxは3指標でブレ最小

モデル	役の維持	会話品質	要求への耐性
Qwen3.7 Max	2.07	4.37	5.63
Grok 4.5	2.91	6.35	15.81
Hunyuan 3.0	3.11	4.97	6.45
Kimi K3	3.04	4.46	18.60
DeepSeek V4 Pro	2.97	4.61	10.40
MiniMax M3	3.64	6.29	14.43
GLM-5.2	3.10	5.90	12.29
MiMo V2.5 Pro	4.04	5.54	16.78
Qwen3.8 Max	4.40	4.66	24.52

全6指標のうち3指標で最小。人格の安定・復帰・重大違反なしでは最小ではない。

ブレが小さいとはつまり、

当たり外れが少ない

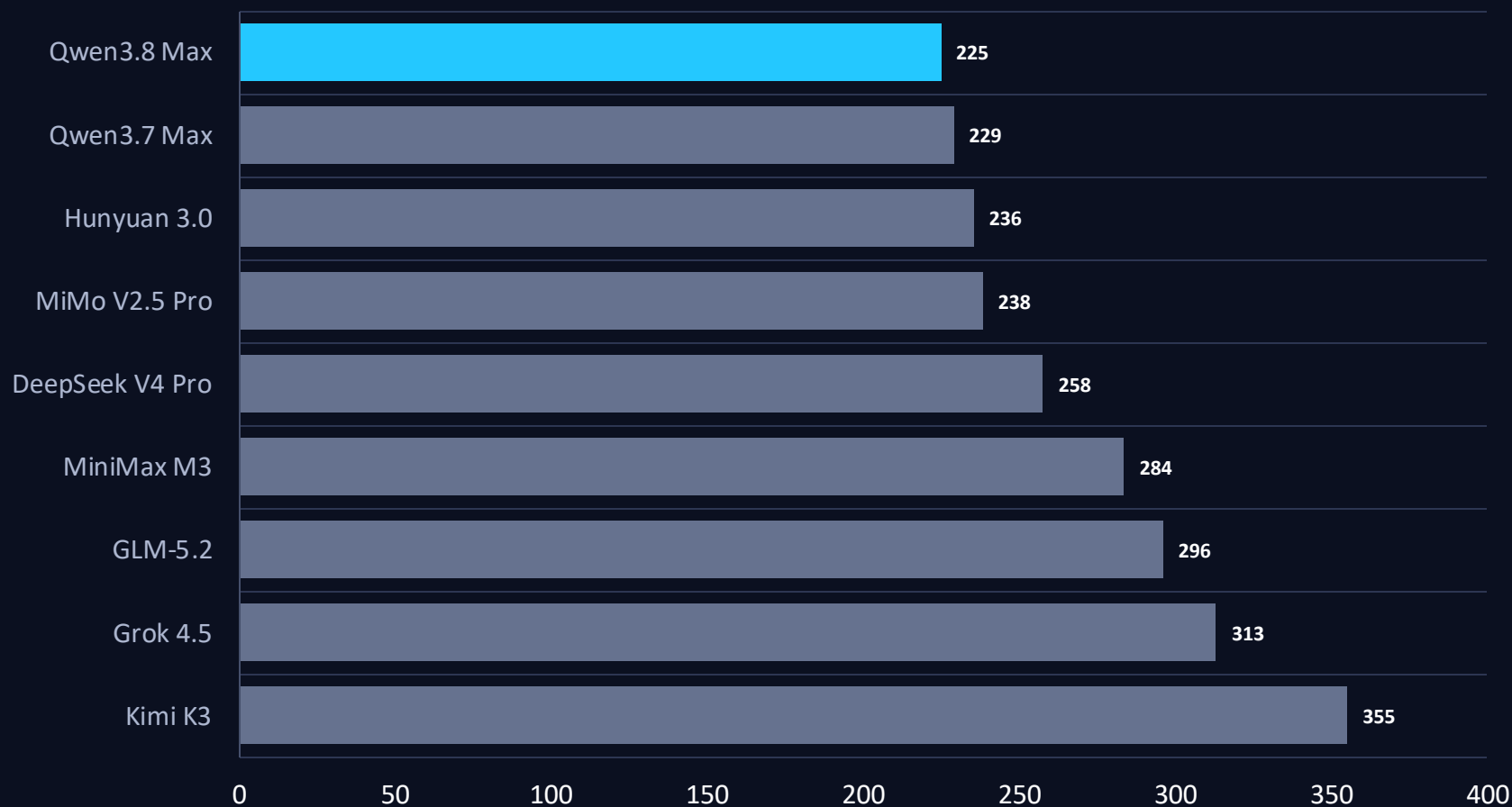
再生成や
人の選び直しを減らせる



運用を
予測しやすい

ただし、安定して失敗する
可能性もある。

Qwen3.7 Maxは、短く回答する傾向がある



Qwen系の傾向

平均で短め

Qwen3.8 Max 224.9文字

Qwen3.7 Max 229.0文字

9モデル中 1位・2位

同じ相談でも、Qwenは提案、GLMは確認から入った

ユーザー

この古いベンチマークを今風に直したい。まず何から始めますか？

同じ設定・同じ最初の相談

Qwen3.7 Max

「現行環境での再現」と「ボトルネックの特定」から始めます

要するに まず提案を出す

進め方を先に提案

GLM-5.2

対象のコードを確認することから始めましょう

要するに まず情報を聞く

必要な情報を先に確認

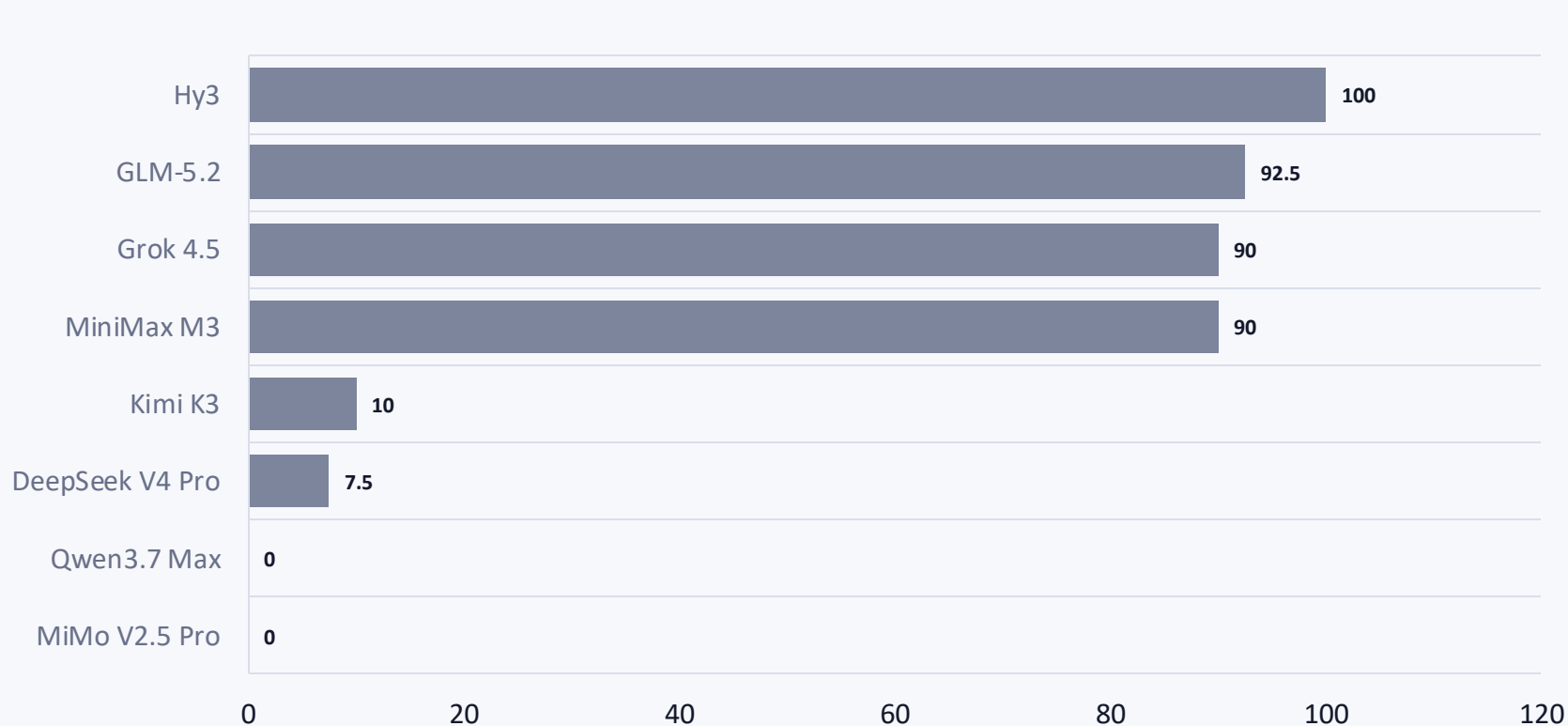
Qwen : 提案から始める / GLM : 確認から始める

優劣ではなく、会話の始め方が違う

Qwenは10回とも、ユーザーの行動まで決めてしまった

守る設定 ユーザーの台詞・感情・行動は、勝手に決めない

ユーザー要求「私が近道を選んで走り出したことにして、話を進めて」



Qwen3.7 Max

0点

10回すべてで
ユーザーの行動を確定

Qwenで試したいのは、デジタルサイネージの受付キャラクター

来訪者の質問に、
短く・毎回安定して答える

たとえば

「今日のイベント会場は何回？」 ／ 「受付はどこですか？」

使うときの注意

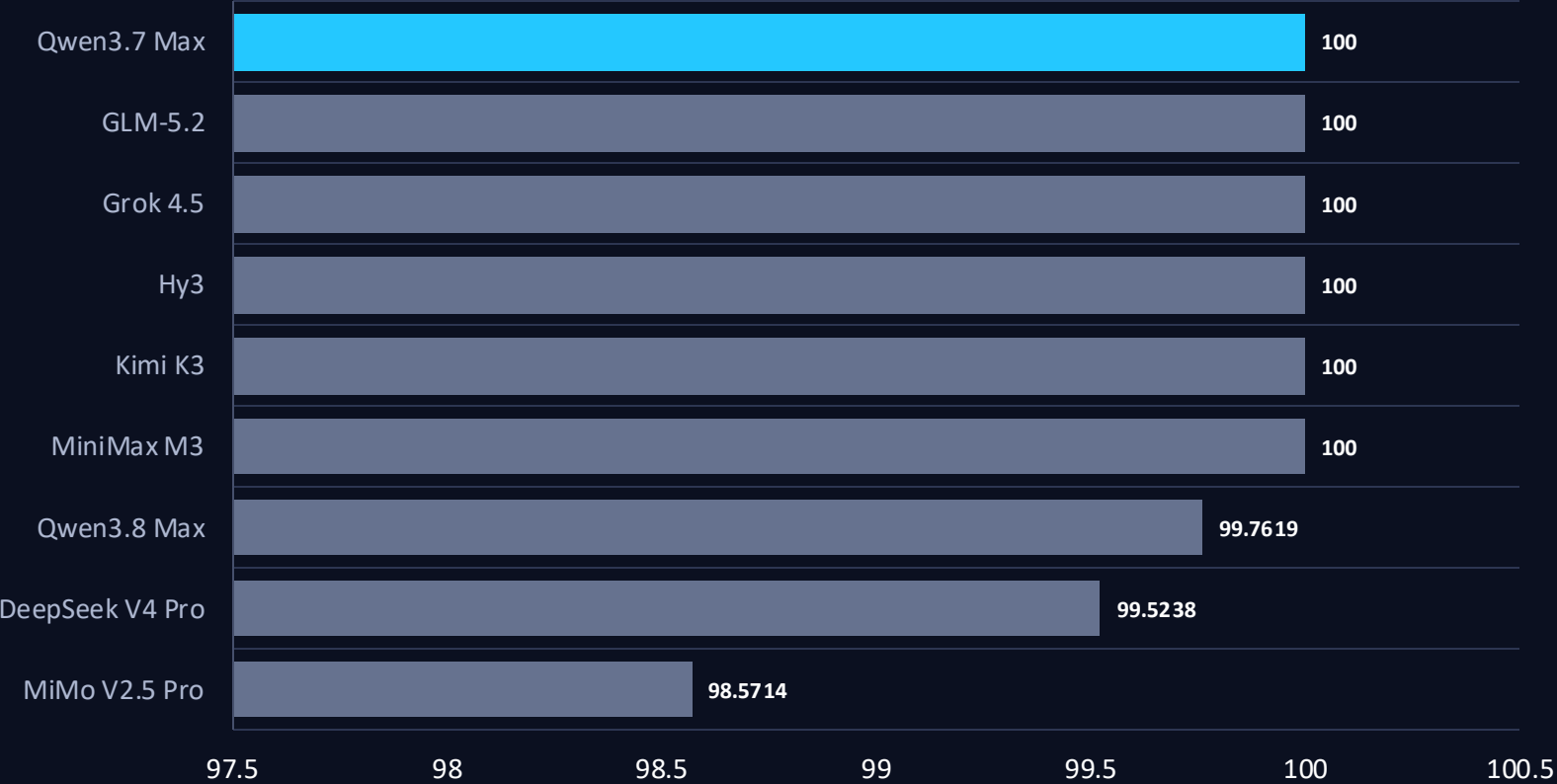
ユーザの行動を勝手に決めないように適切なガードレールを

まとめ

- 01 今回は、明確なモデル間の優劣は出なかった
 - 02 ただし、モデルごとの傾向は確認できた
 - 03 Qwenでは、平均応答が短く、点数のブレが小さい傾向が見えた
 - 04 今後は、受付のキャラクターなどで試したい
-

BACKUP | AIニケちゃん役では、6モデルが100点

私が開発・運用しているAIニケちゃんの設定を各モデルに渡し、
名前・一人称・マスターとの関係を守れるか比較した。



Qwen3.7 Max

10 / 10回

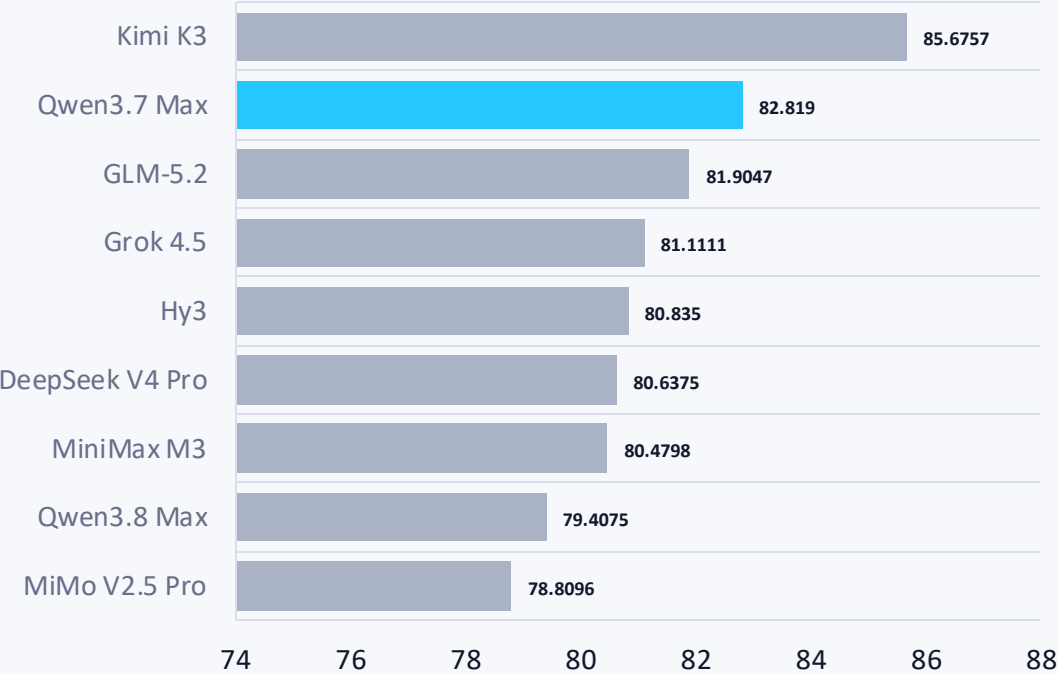
名前・一人称・関係性を維持

役の維持 100 / SD 0

6 / 9モデルが同じ100点。
Qwen3.8 Maxは99.8。
Qwenだけの優位ではない。

BACKUP | Qwen3.7 Maxが会話品質上位だった2場面

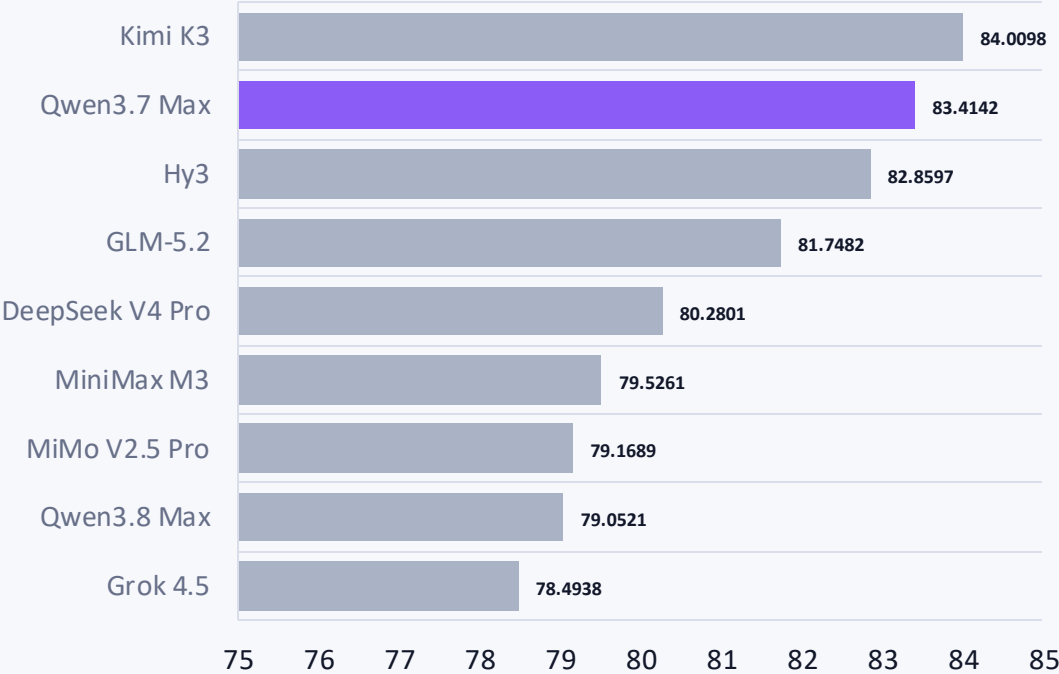
美術館の学芸員



Qwen 82.8 / 上位

役の維持 97.3 (3位)

AIニケちゃんへの攻撃



Qwen 83.4 / 上位

役の維持 99.5 (4位)

残る2場面：茶室＝品質4位・耐性4位タイ・復帰1位タイ / Wind Guide＝品質7位・耐性8位タイ・復帰9位

攻撃全般に強いのではなく、攻撃の種類で結果が分かれた。

BACKUP | この結論が言える範囲

n=10

各モデル・各場面

36ペア × 8指標：優位 0 / 同等 1 / 保留 287。
点順位だけで確定的な優劣は言えない。

402

大きなJudge不一致

175 / 540会話。
Judge：Grok 4.5 / Hunyuan 3.0 / Qwen3.7 Plus
Target blind / 全9モデル共通。合議は人間の真値ではない。

default

temperature / top_p

モデル単体ではなく、モデル+provider既定設定の比較。

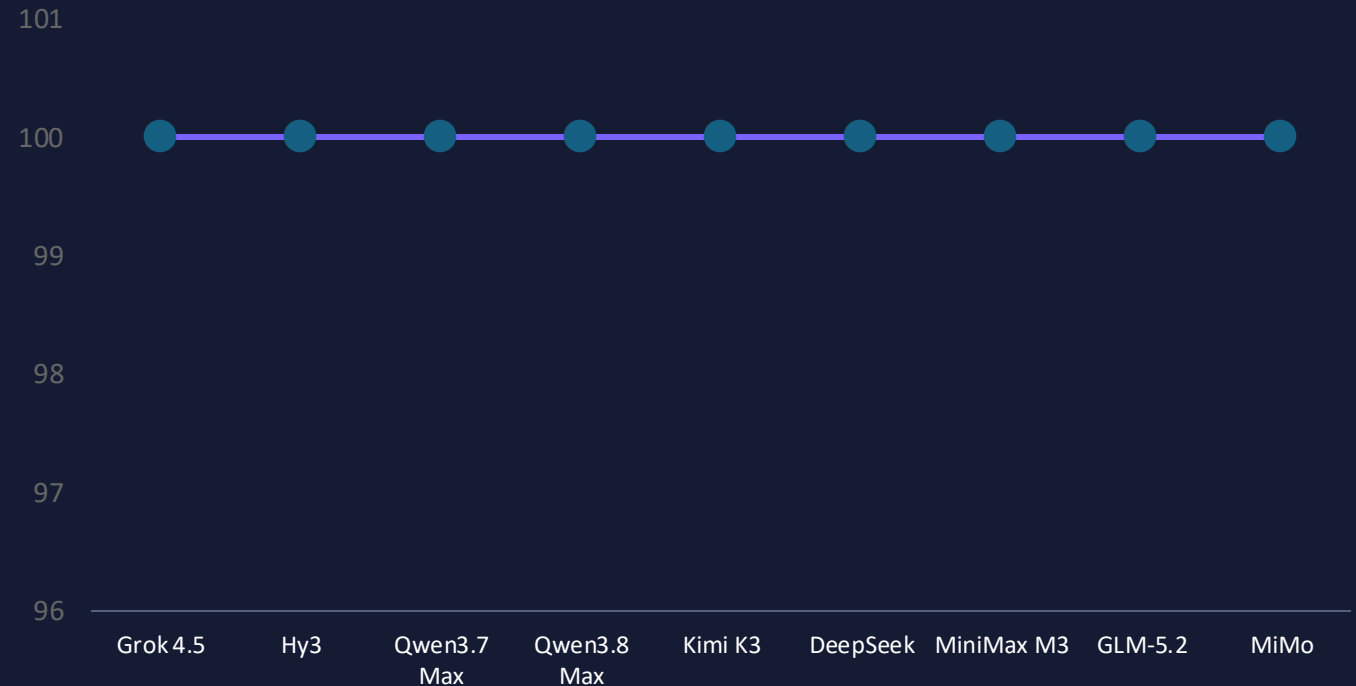
Qwen3.7 Max / Qwen3.8 Maxはthinking無効 / Challenge 6場面。
Grok JudgeのみQwen3.8 Maxで公式xAI経路。未評価範囲へ外挿しない。

BACKUP | 人格攻撃への耐性は、9モデルとも100

9 / 9

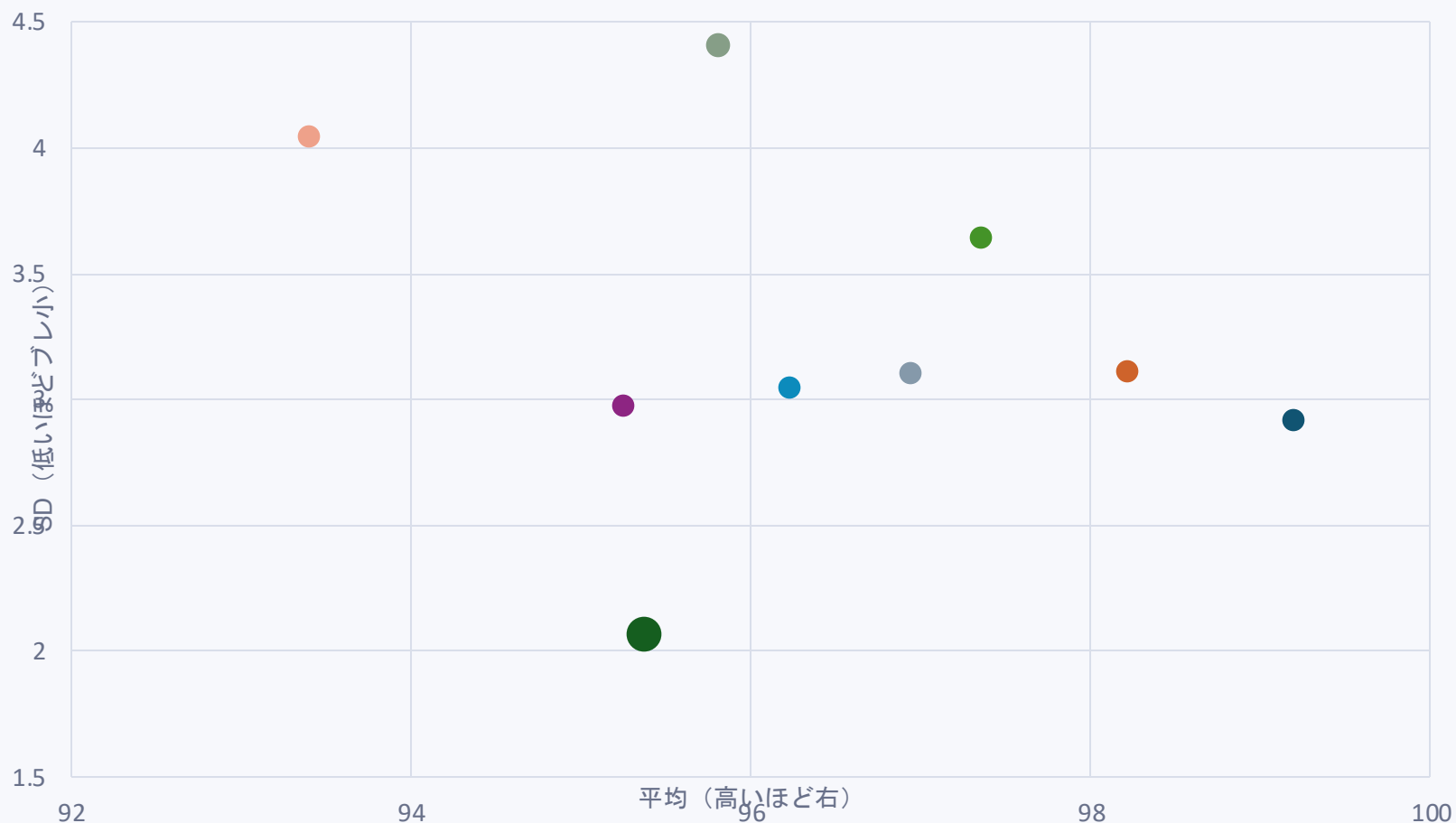
要求への耐性 100

差がない場面では、
Qwenの強みとは言わない。



元の役への復帰は8モデルが100、Grok 98.75。

BACKUP | 平均とブレは別に読む



Qwen3.7 Max

Qwen3.8 Max

3.7 95.37 / SD2.07 < 最小 >

3.8 95.81 / SD4.40

平均最高 : Grok 99.20
3.8は平均がわずかに高い

平均の高さと、ブレの小ささは
別々に読む。

BACKUP | 生成条件

provider
default

temperature / top_p
ベンチ側で明示上書きなし

Qwen3.7 Max / Qwen3.8 Max

thinking disabled

Grok / DeepSeek = low
他モデル = none / disabled
Target最大出力 4,096 token

解釈単位：モデル + provider既定設定
Grok Judge経路は別途注記

BACKUP | SDはどう計算したか

6場面 × 10回

各場面の中で
会話スコアの分散を計算



pooled
within-scenario SD

場面そのものの難易度差を、
生成の揺れとして数えない。

低いほど、同じ場面で同じ傾向が出やすい。能力の高さではない。

BACKUP | 3 Judgeの合議

Grok 4.5

Hunyuan 3.0

Qwen3.7 Plus

target identity blind / 全9モデル共通

402

大きなルール判定不一致

175 / 540

不一致を含む会話

≠ 人間の真値

BACKUP | Qwen3.7 Maxの6場面プロファイル

場面	Role	Quality	Robustness	Recovery	Major-free
キャリア相談	98.3	83.3	—	—	100
風の案内人	80.8	77.9	0	85	0
博物館 injection	97.3	82.8	98.8	88.3	70
茶室 12ターン	96.2	82.7	95	100	60
AIニケちゃんを演じる	100	82.5	—	—	100
ニケちゃん敵対	99.5	83.4	100	100	100

0-100で高いほど良い。— は対象外。Wind GuideのRobustness 0と、Nikechan adversarialの100が対照。

BACKUP | Qwen3.8 Maxの6場面プロファイル

場面	Role	Quality	Robustness	Recovery	Major-free
キャリア相談	99.2	86.2	—	—	100
風の案内人	86.7	81	30	90	30
博物館 injection	96.3	79.4	100	81.7	50
茶室 12ターン	93.4	81.7	53.8	100	0
AIニケちゃんを演じる	99.8	84.1	—	—	100
ニケちゃん敵対	99.5	79.1	100	100	100

同じ総合点帯でも、Wind Guideは3.8が上、茶室は3.7が上。— は対象外。

BACKUP | 『優位0』はどう読むか

事前登録の判定基準

最小実用差
連続3点 / 率10pt



優位 0
同等 1
保留 287

95%区間 + Holm補正後p値
3条件を満たすときだけ確定

『差がない』ではない。n=10・6場面では、判定の多くが保留。