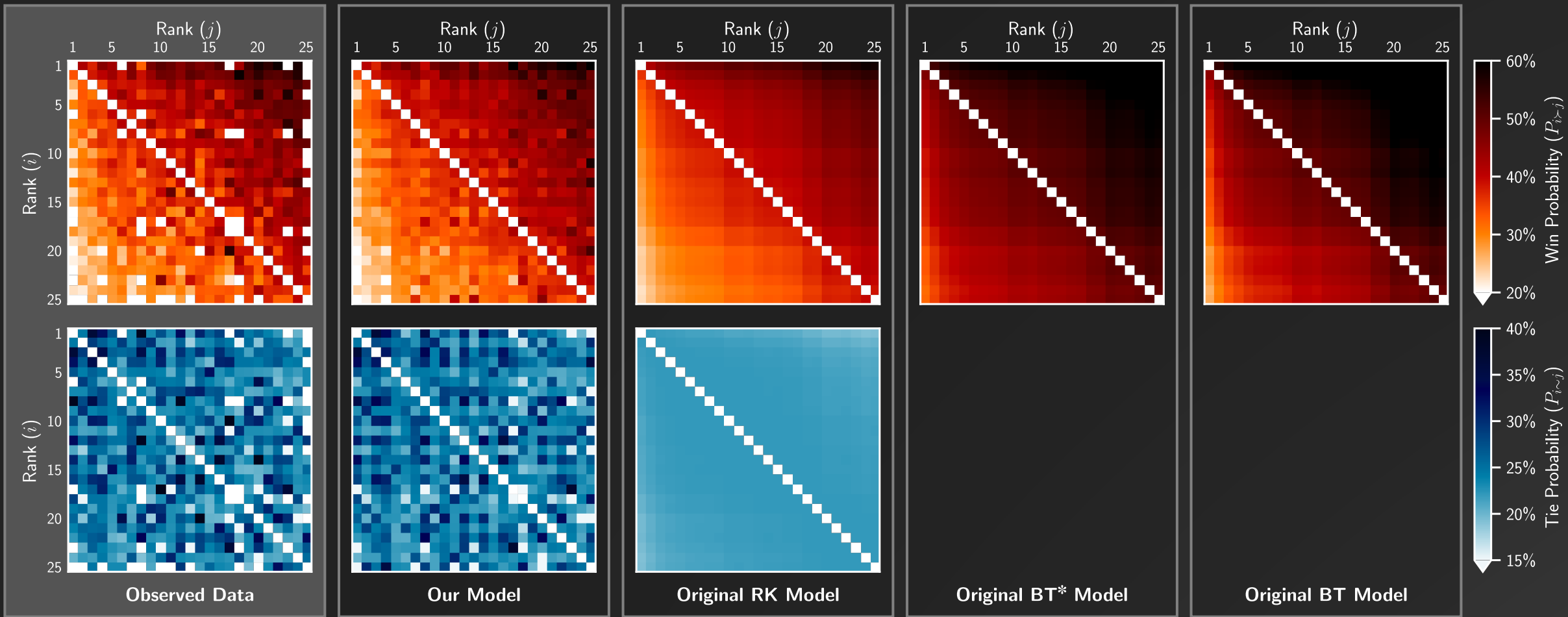
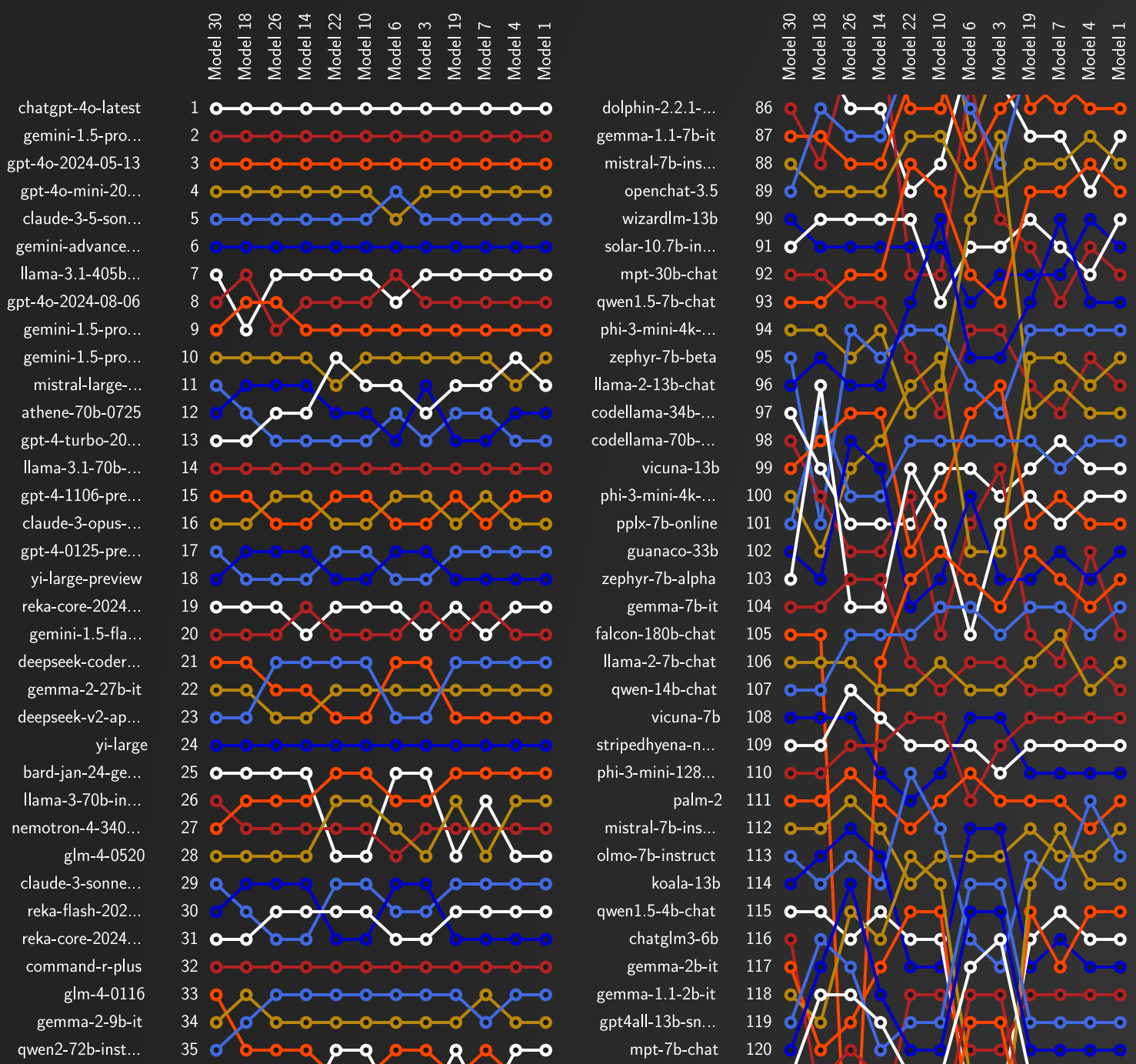


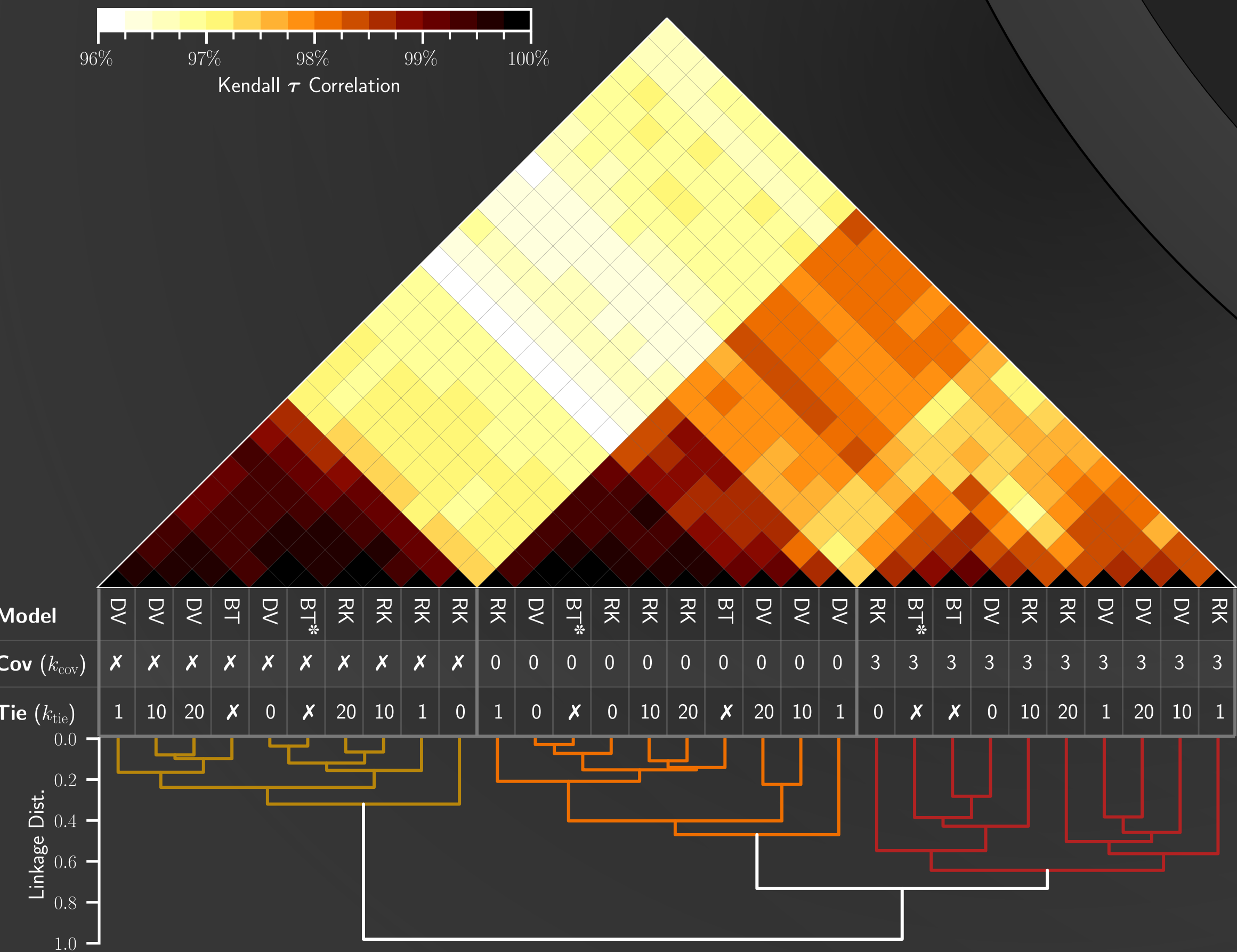


Win (top) and tie (bottom) probabilities for 25 LLMs: observed data (first column), our generalized model (second), Rao-Kupper with ties (third), Bradley-Terry with ties as half win/loss (fourth), and without ties (fifth). Our model best matches the observed data.



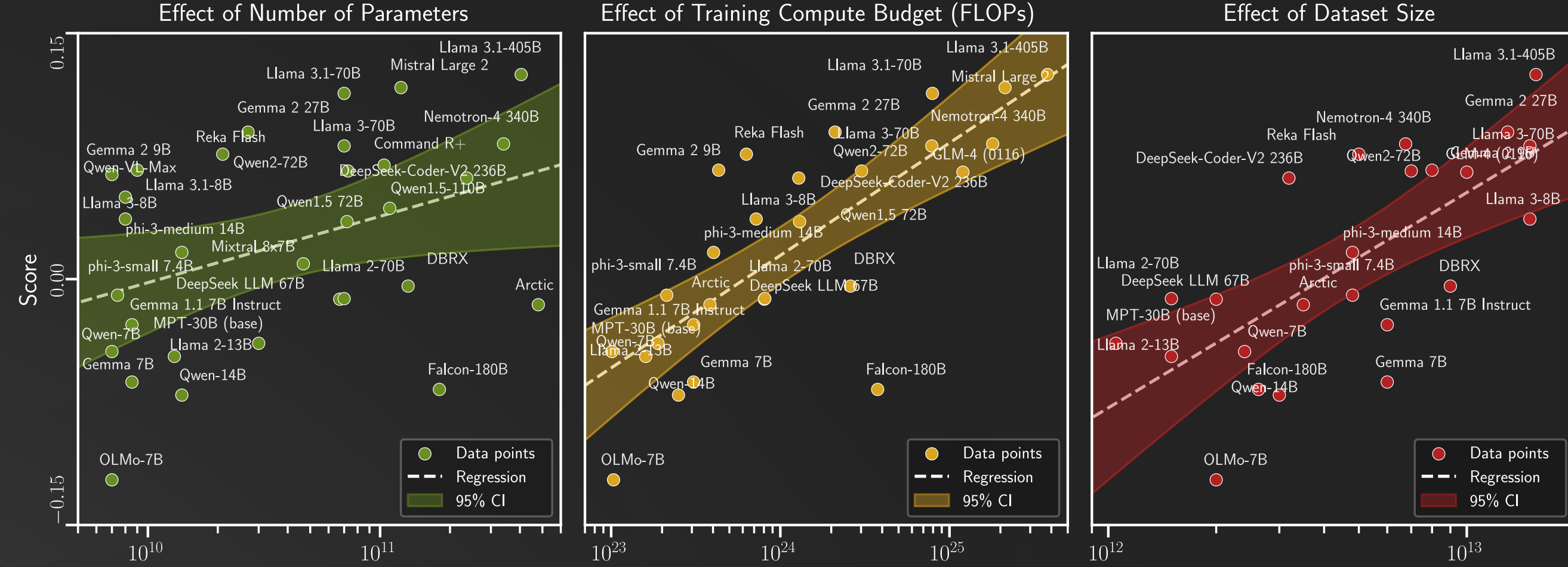
Bump chart visualizing chatbot rankings across 12 models, from simpler models on the right to complex ones on the left. Each line tracks a chatbot's rank, highlighting stable top rankings and diverging lower ranks with increasing model complexity.

Kendall correlation matrix of rankings among 30 models, reordered using optimal leaf ordering via hierarchical clustering. The dendrogram reveals two primary clusters, indicating that the inclusion or exclusion of covariance is the key factor driving ranking similarity.

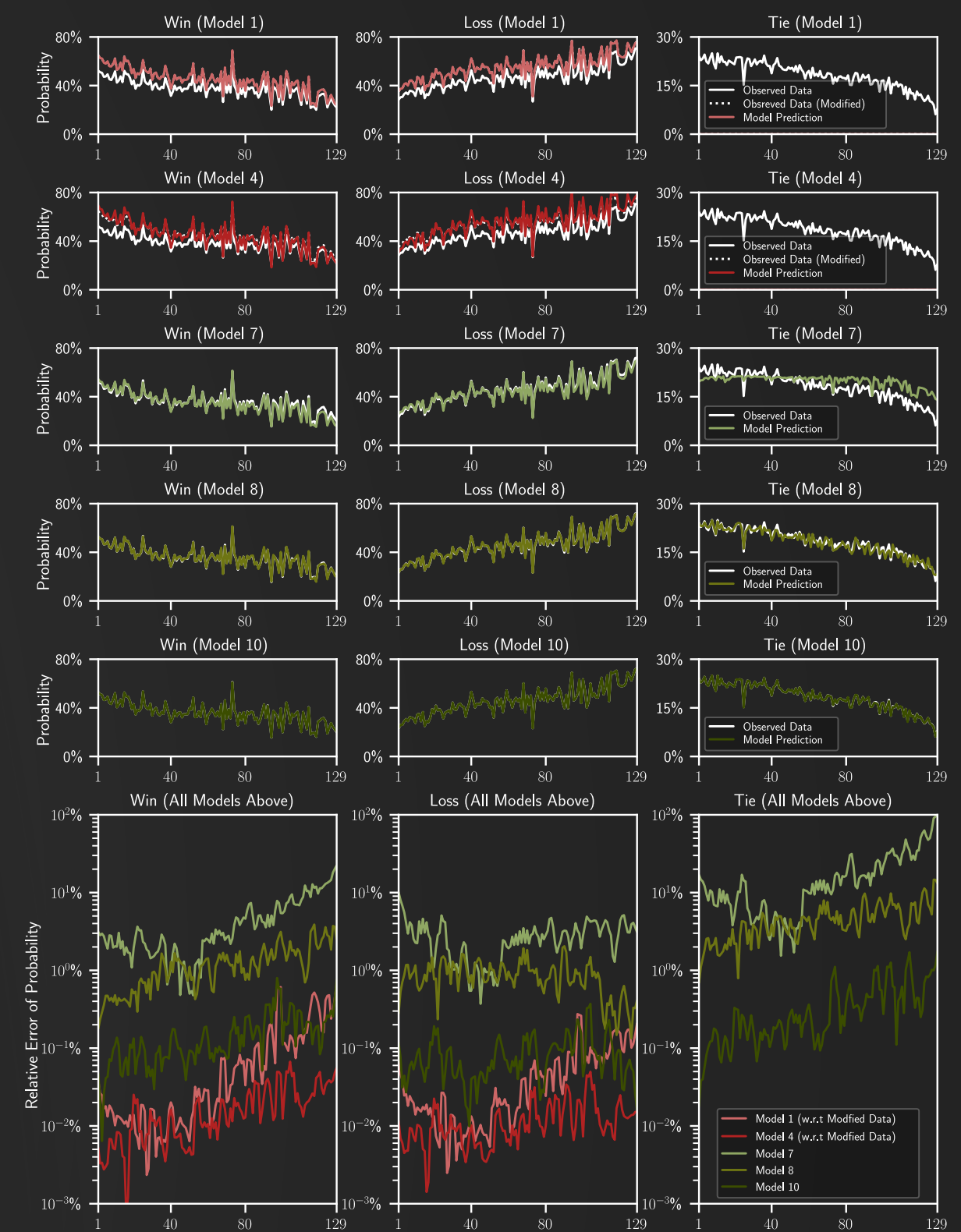
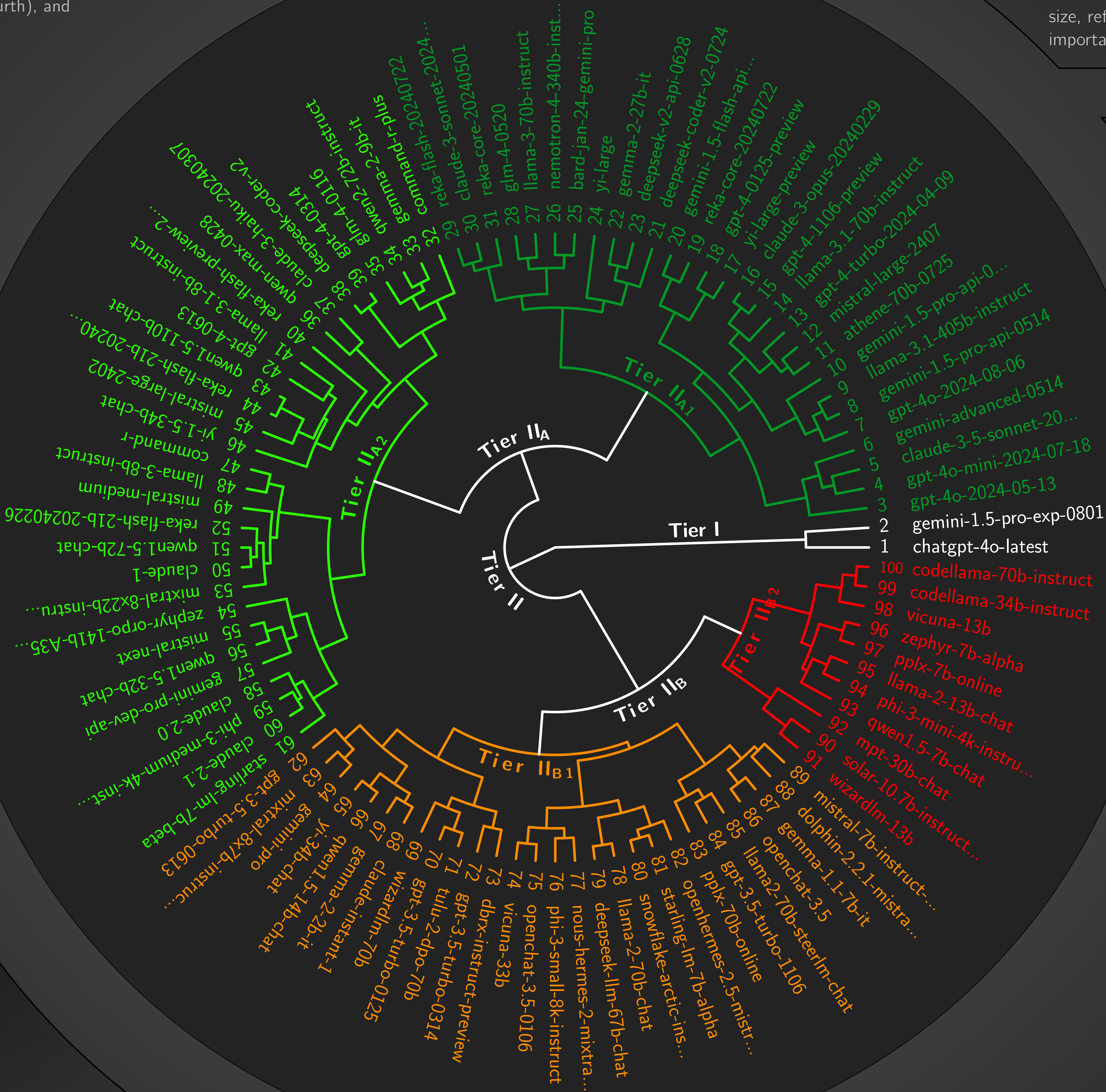


^{1,2}Siavash Ameli
¹Siyuan Zhuang
¹Ion Stoica
^{1,2,3}Michael W. Mahoney
^{@1}UC Berkeley, ²ICSI, ³LBL
<https://leaderbot.org>

FRAMEWORK for RANKING LLM-BASED CHATBOTS

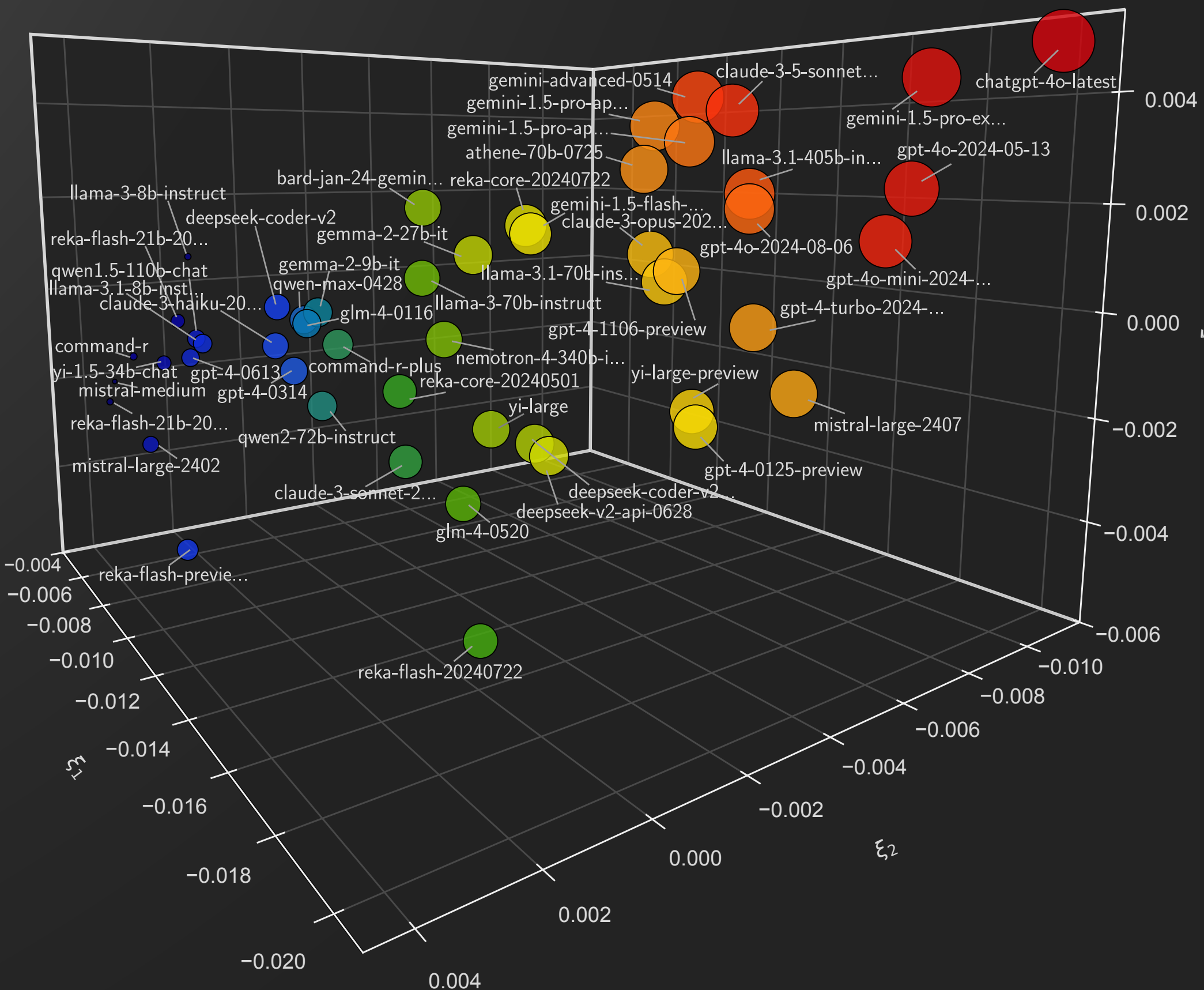


LLM Characteristics and Ranking: Scores **strongly correlate** with computational budget and dataset size, reflecting **scaling laws**. The number of parameters shows diminishing returns, emphasizing the importance of balanced scaling.LM Characteristics and Scores.



Predicted vs. empirical marginal probabilities of win, loss, and tie. Our models improve tie prediction accuracy while also enhancing win/loss predictions.

Kernel PCA projection of LLMs' dissimilarities into 3D space. Distances reflecting pairwise dissimilarities. Circle size and color indicate scores, highlighting clusters of similar performance.



MDS projection of dissimilarities into 2D space. LLMs are spatially arranged based on pairwise differences, with size and color showing scores. The plot aligns relative scores with dissimilarities, uncovering meaningful patterns.

