

STRICT: Stress-Test of Rendering Image Containing Text

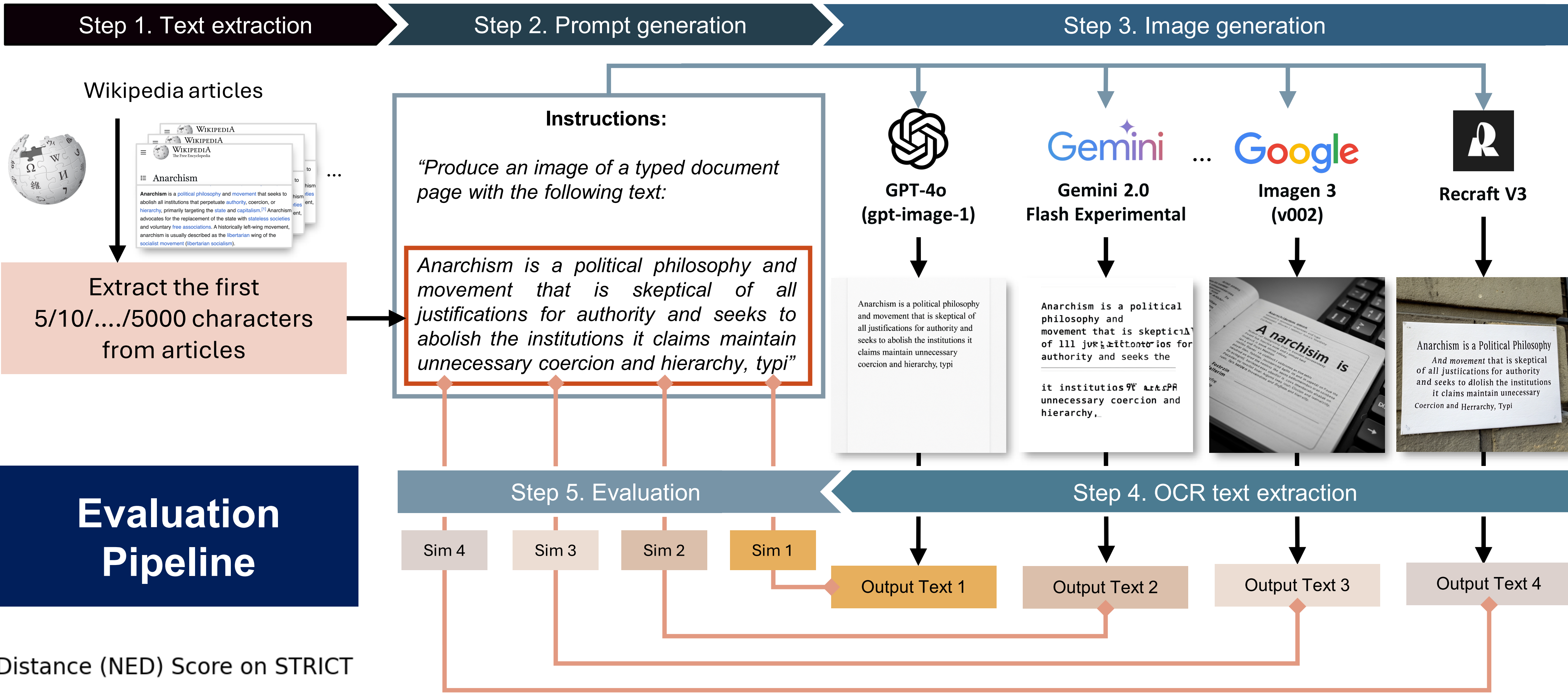
Tianyu Zhang^{1*}, Xinyu Wang^{2*}, Lu Li^{3*}, Zhenghan Tai⁴, Jijun Chi⁴, Jingrui Tian⁵, Hailin He⁶, Suyuchen Wang¹

¹ Mila, University of Montreal, ² McGill University, ³ University of Pennsylvania, ⁴ University of Toronto, ⁵ University of California, Los Angeles, ⁶ Southwestern University of Finance and Economics *Equal contribution

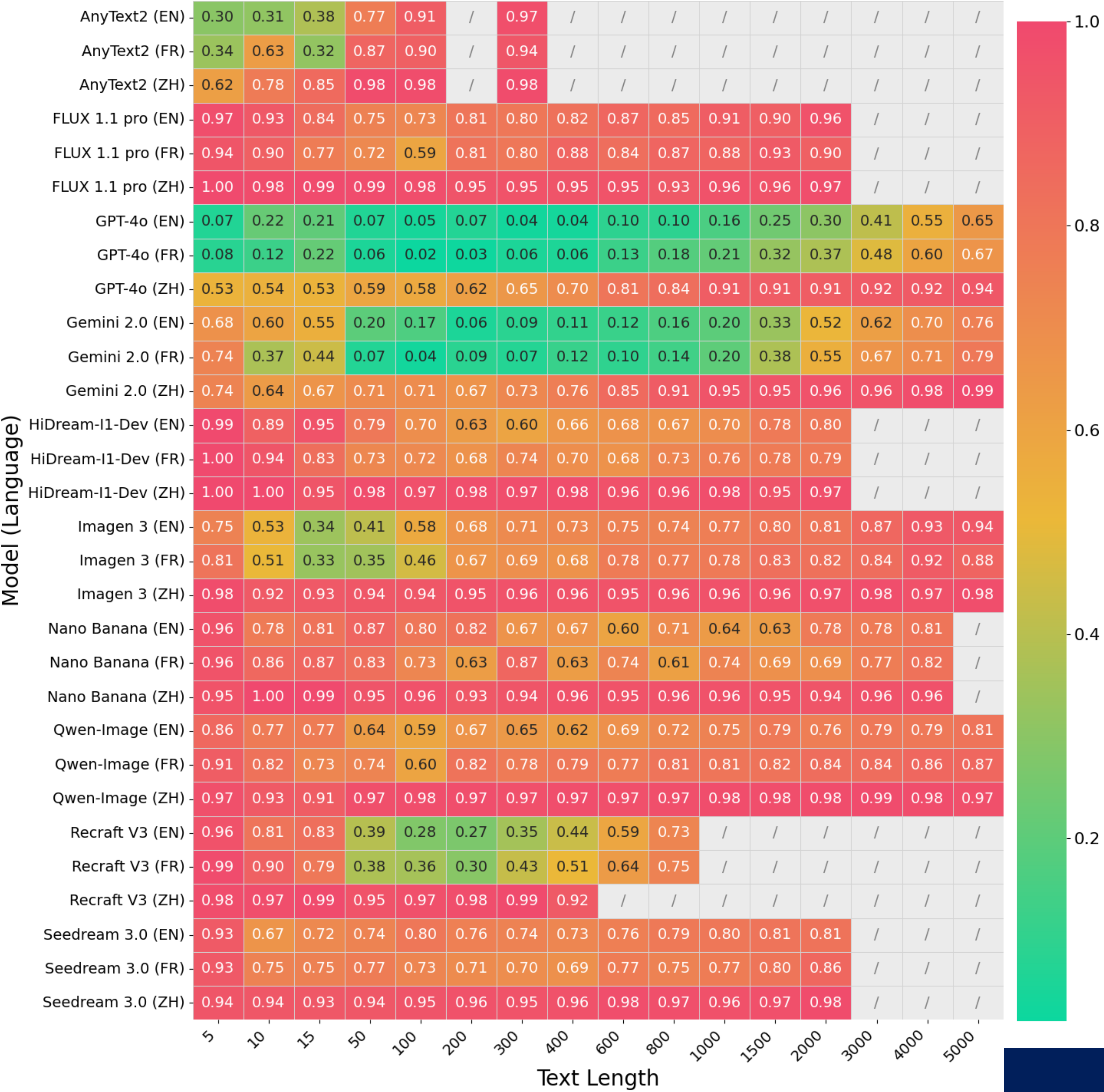


Motivation

- Diffusion models excel at realistic image synthesis
- Still fail to generate **legible, consistent, and instruction-aligned text**
- Root cause: **locality bias** → weak long-range spatial consistency
- Need a systematic test for text fidelity and instruction following

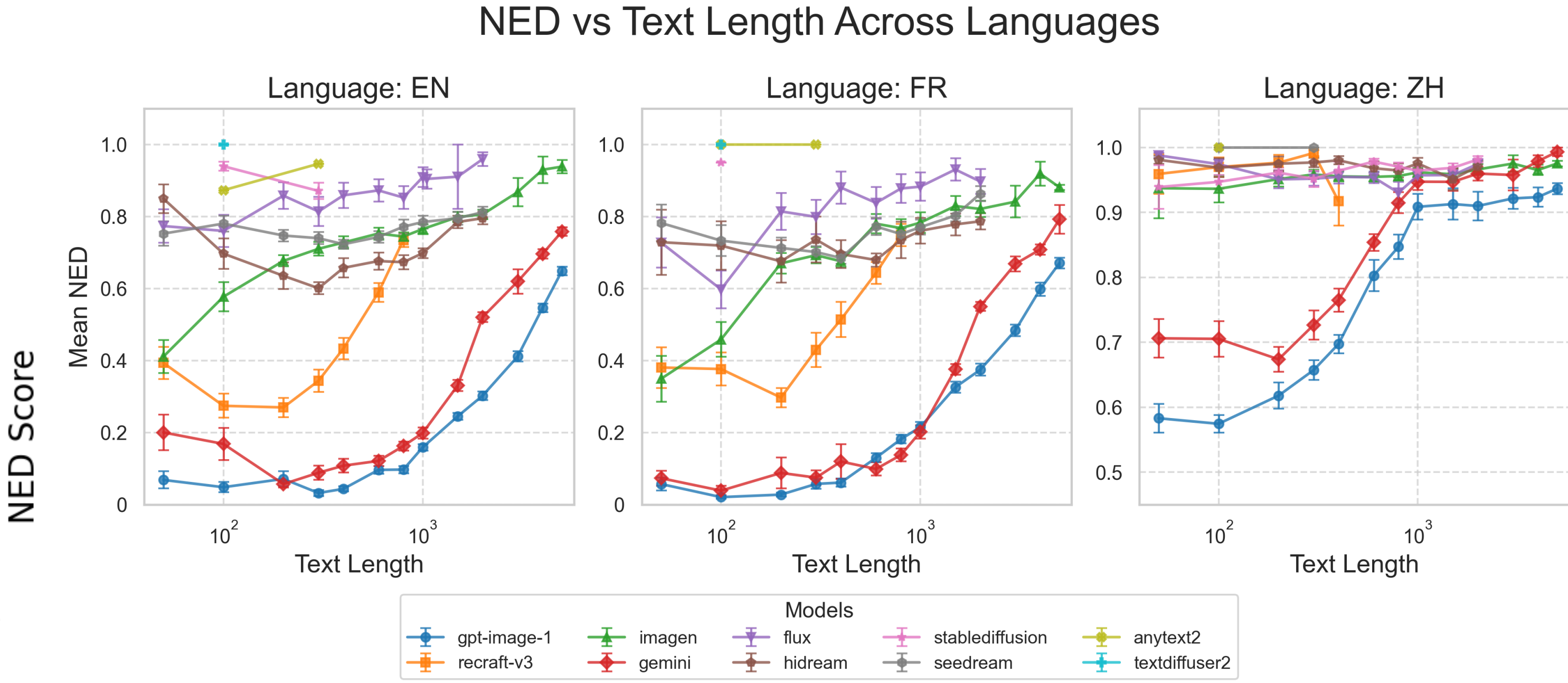


Heatmap of Normalized Edit Distance (NED) Score on STRICT



Goal: Quantitatively assess and expose model limitations in generating images containing text.

Evaluation Metrics & Results



- **Normalized Edit Distance (NED)** – character-level dissimilarity
- **Performance vs Text Length**
 - Accuracy drops sharply beyond 200–300 characters
 - GPT-4o maintains best NED (<0.3 up to 2K chars)
 - Open-source models degrade to **NED > 0.8** early
 - Long-text rendering remains a universal failure case

Case Study

Takeaways

- Diffusion models still **cannot** render image containing texts reliably
- Exposes instruction-following and consistency gaps
- **STRICT** provides a **flexible and controllable synthetic benchmark** to evaluate the abilities of models to generate images containing texts

More Info



Paper

Dataset

GitHub

