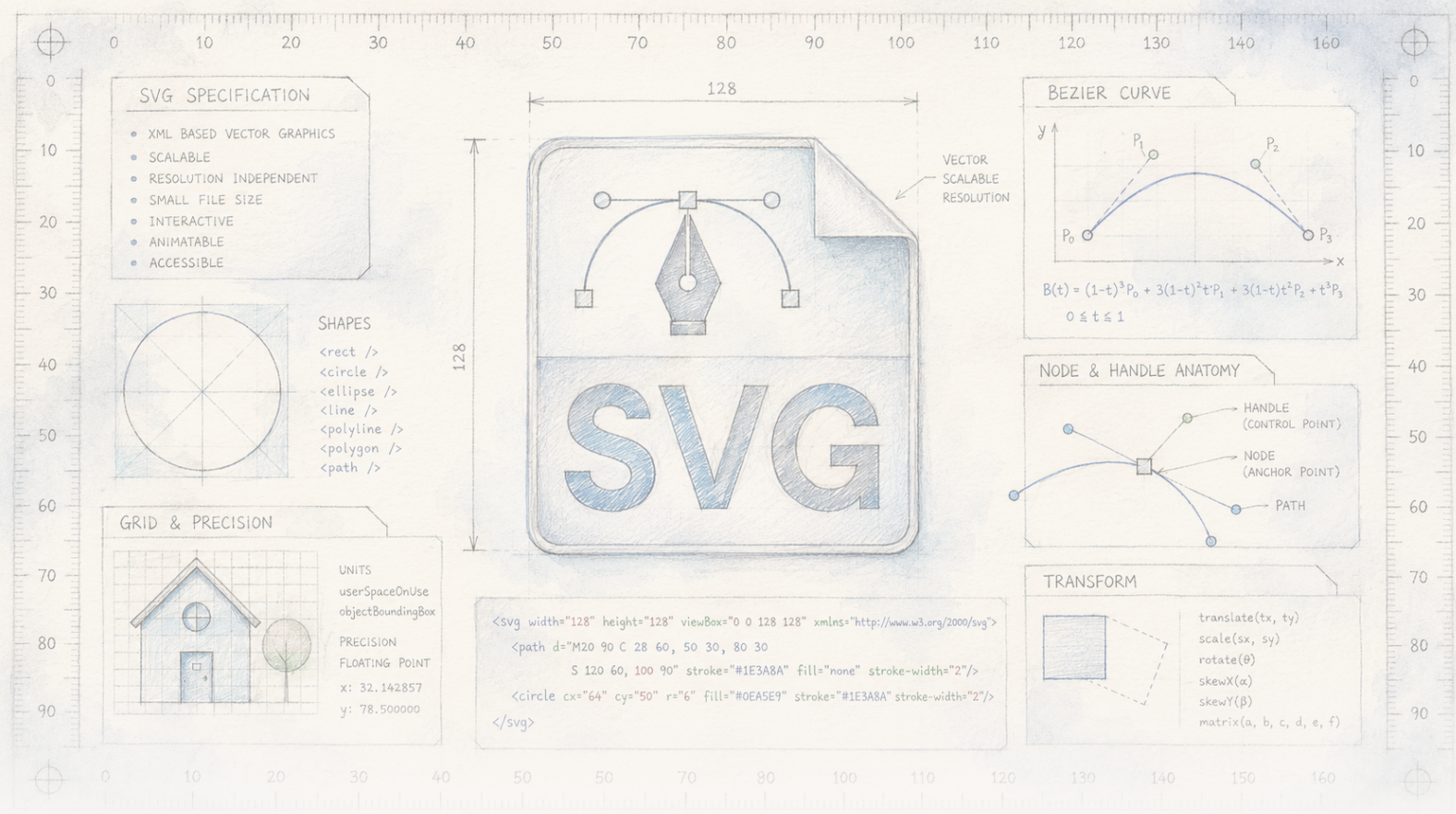




The CLEAN PAPER

WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

Drawing AI को अपनी ही canvas देखने की ट्रेनिंग

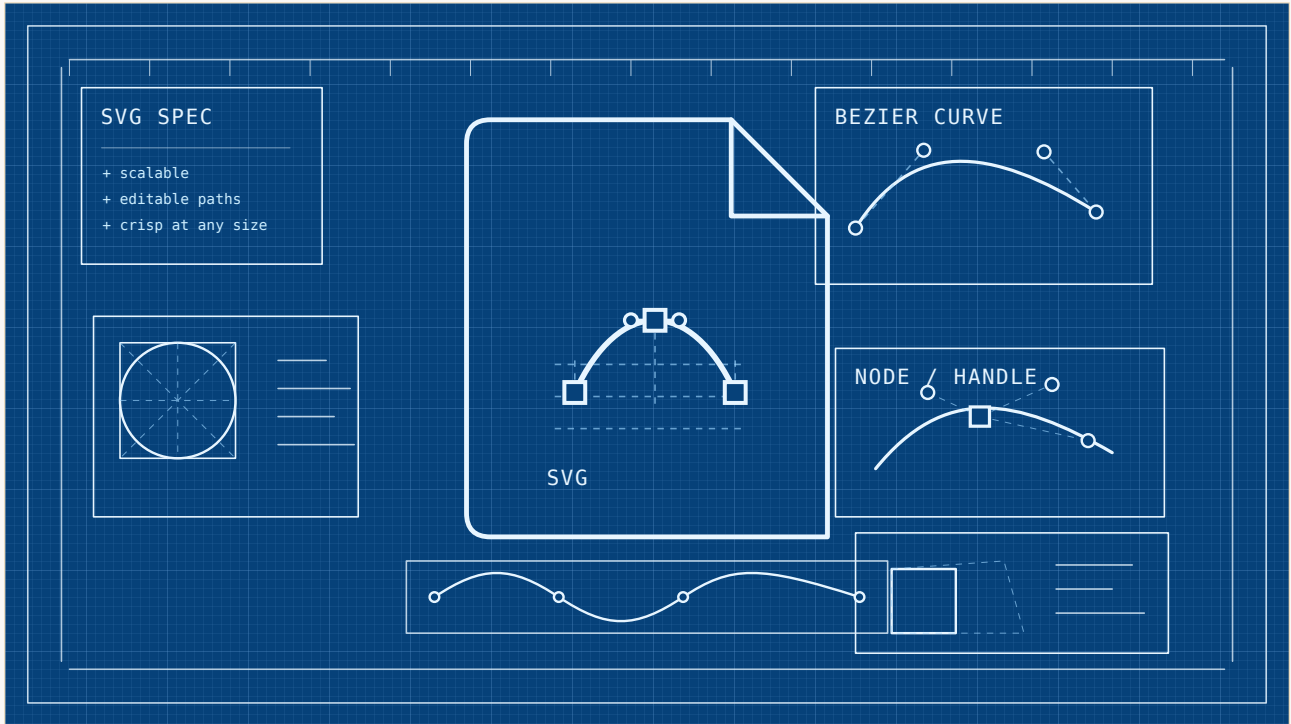
Vector graphics बनने वाले language models आम तौर पर “अंधे” कम करते हैं — drawing commands लिखते जाते हैं, लेकिन बनती हुई image देखते नहीं नई method model को stroke-by-stroke canvas भरते हुए देखने देती है। दिलचस्प मोड़ यह है कि सिर्फ “आँखें” दे देने से performance खरब होती है: model को visual feedback इस्तेमाल करना दोष शासी बना पड़ता है। सही retraining के बाद model उन rivals के बराबर या थोड़ा आगे पहुँचता है जिन्हें बीस गुना तक अधिक data मिला था। एक benchmark पर वास्तविक और सन्नधनी से मिला result — revolution नहीं

Written by Vera Rubin · Cross-checked by Michele Renda · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

Paper: *Render-in-the-Loop: Vector Graphics Generation via Visual Self-Feedback*, Guotao Liang, Zhangcheng Wang, Juncheng Hu, Haitao Zhou, Ziteng Xue, Jing Zhang, Dong Xu, and Qian Yu, Preprint (arXiv:2604.20730)

आँखें बंद करके चित्र बनना

आज के किसी AI मॉडल से तस्वीर बनने को कहिए तो भीतर दो बहुत अलग चीजों में से एक होती है। परचिति छवि-जनरेटर — जो एक वक्त्र से फोटोजैसी तस्वीर बना देते हैं — सीधे पक्सेल बनते हैं और कम करते समय कैनवास को देख भी सकते हैं। लेकिन चित्र बनने का एक दूसरा शान्त तरीका भी है जिसमें मॉडल रंग नहीं भरता वह निर्देश लिखता है: यहाँ एक वृत्त बनाओ, वहाँ एक रेखा इस आकार की बनाओ ये निर्देश कोड होते हैं — वही Scalable Vector Graphics, यानी SVG, जिसे वेब के बहुत-से आइकन और लोगो बनाते हैं। इसका फायदा वस्तुतः यह है: परिणाम पक्सेल की स्थिति जल्दी नहीं बदलता आकृतियों का समूह होता है, जिसे बना धुंधला किए बार-बार बड़ा छोटा रंग बदला और संपन्न किया जा सकता है।



मूल SVG ब्लूप्रिंट कलंकृत: छवि स्वयं संपन्न योग्य वेक्टर फ़ाइल है, जो स्थिति पक्सेल को बजाए paths, grid lines, nodes और handles से बनी है।

Laura Nesso / The Clean Paper Permission on file

समस्या यह थी कि हाल तक ऐसा चित्र-कोड लिखने वाला मॉडल लगभग अंधे की तरह काम करता था। वह निर्देशों की पूरी श्रृंखला — वृत्त, रेखा, रंग — एक ही बार में निकाल देता, लेकिन बीच में उन्हें रेंडर करके यह नहीं देखता कि उसने वास्तव में बना क्या है। आँखें बंद करके चेहरा बनने की कल्पना की जाए: आँखें लगभग सही जगह आ सकती हैं, लेकिन आपको पता ही नहीं चलेगा कि दूसरी आँख गलत पर पहुँच गई या बसो की आकृति नक्काशे के ऊपर बैठ गई। ये मॉडल मंटे तौर पर इसी तरह काम करते थे, इसलिए वे अक्सर ऐसा कोड देते थे जो तकनीकी रूप से बिल्कुल वैध होता, लेकिन देखने में गड़बड़।

Guotao Liang के नेतृत्व वाली टीम ने अब लगभग हस्यस्पद रूप से सीधा कदम उठाया है: मॉडल को अपनी बनी हुई चीज़ देखने दी। उनकी विधि *Render-in-the-Loop* हर चरण के बाद अधूरे चित्र को रेंडर करती है और अगला stroke बनने से पहले वह तस्वीर मॉडल को वापस देती है। सचमुच दिलचस्प बात यह नहीं कि इससे मदद मिलती है, बल्कि यह है कि मदद मिलने के लिए मॉडल को क्या सिखाया जाए।

कसीचित्र कोअंक कैसे दें?

जब कोई शोधपत्र कहता है कि उसका मॉडल दूसरे मॉडल को “हरता” है, तो पूछना उचित है: किस कम में, और किस मामले में? बनई गई तस्वीर का मूल्यंकन सचमुच कठिन है — “लैपटॉप बनाना” का केवल एक सही उत्तर नहीं होता — इसलिए यह क्षेत्र कुछ स्वचालित मामलों पर निर्भर करता है। इनमें से हर मामला इस मीनजर का अधूरा विकल्प है। इस शोधपत्र में ऐसे कई मामले आते हैं। FID बनई गई छवियों के पूरे समूह की संख्यांकित बनाना की तुलना वस्तुविक छवियों से करता है; कम मामला बेहतर है, लेकिन यह पूरे समूह के बारे में बताता है, कसी एक तस्वीर के बारे में नहीं। CLIP score एक अलग AI से पूछता है कि छवि प्रमॉप्ट के शब्दों से कतिनीमेल खती है — उपयोगी है, लेकिन उस निर्णायक मॉडल की क्षमता जतिनाही अच्छी। DINO, SSIM और LPIPS पुनर्निर्मित छवि की लक्ष्य छवि से अलग-अलग स्तरों पर तुलना करते हैं, कच्चे पक्सेल से लेकर सीखी हुई विशेषताओं तक।

इनमें से कोई भी सत्य का अंतिम मामला नहीं है। ये प्रतनिधि संकेतक हैं और आम तौर पर इनमें छोटे बदलाव ही आते हैं। जब आप पढ़ते हैं कि एक मॉडल 127.6 और दूसरा 128.8 स्कोर करता है, तो इच्छा दशम में यह वस्तुविक अंतर है — लेकिन भरीजति नहीं बस हल्का अंतर; और उस संख्या में, जो आपकी आँख के निर्णय से केवल दूजे तौर पर जुड़ी है। “बीस गुना डेटा पर प्रशिक्षित मॉडल को हरना” जैसी सुखी पढ़ते समय यह बात यह रखना उपयोगी है।

लेखकोने क्या किया

लेखक जसि समस्या को ठीक करना चाहते थे, वह यही अंधाचित्रंकन था। मौजूदा मॉडल SVG लिखने को शुद्ध पठ-काह्य की तरह लेते हैं: अभी तक के कोड से अगला कोड-खंड अनुमान लगाओ और उसे कभी रेंडर मत करो। इससे आधुनिक multimodal मॉडलों में पहले से मौजूद शक्तिशाली “आँखें” — vision encoder — बेकार बैठी रहती हैं। Render-in-the-Loop इस काम को चरण-दर-चरण दृश्य प्रक्रियामें बदल देता है। चित्र-कोड के हर खंड के बाद अधूरा SVG छवि के रूप में रेंडर होता है और मॉडल को वापस दिया जाता है, तत्काल अगला खंड अब तक के कैनवास को सचमुच देखते हुए चुना जाए।

उनका पहला निष्कर्ष समझने योग्य करने वाला है, और श्रेय की बात है कि वे इसे साफ लिखते हैं: कसी मौजूदा तैयार मॉडल पर बस यह लूप जोड़ देना काम नहीं करता। जब उन्हें बनावशेष प्रशिक्षण के मजबूत समस्त मॉडलों को बीच-बीच के renders दिए, तो गुणवत्ता सुधरी नहीं — बल्कि हर मामले पर गरी। जसि मॉडल को इस काम के लिए अपनी आँखें इस्तेमाल करना कभी सिखाया नहीं गया, वह अचमक ऐसा करना नहीं सीख जाता।

इसलिए असली काम प्रशिक्षण में है। लेखक प्रशिक्षण डेटा को देखना बनते हैं तत्काल हर चित्र कई छोटे, दृश्य रूप से अर्थपूर्ण चरणों में टूटे — जटिल आकृतियों को सरल हस्सों में बाँटते हुए, तत्काल हर चरण पर देखने के लिए सचमुच कुछ नया हो — और फिर Qwen3-VL पर आधारित आठ अरब पैरामीटर वाले खुले मॉडल को इन चरण-दर-चरण अनुक्रमों पर fine-tune करते हैं। इसे वे Visual Self-Feedback कहते हैं। चित्र बनते समय वे दूसरा तंत्र, Render-and-Verify, जोड़ते हैं: हर नए stroke को स्वीकार करने से पहले मॉडल उसे रेंडर करता है और देखता है कि उसने तस्वीर में सचमुच कुछ बदला या केवल पछिले stroke को दोहराया। जो strokes कुछ नहीं जोड़ते उन्हें हटा दिया जाता है, और जब आगे कुछ उपयोगी नहीं बचता तो मॉडल को रोकने का संकेत दिया जाता है। उल्लेखनीय है कि यह सब अपेक्षकृत छोटे डेटासेट — लगभग 850,000 उदाहरण — पर चलता है, जो एक प्रतद्वंद्वी के डेटा के आधे से भी कम और दूसरे के छोटे अंश के बराबर है।

उन्हें क्या मिला

- इस तरह प्रशिक्षित मॉडल अपने अंधे समकक्ष से बेहतर चित्र बनाता है — खासकर उन वफिल मामलों में जहाँ अंधा मॉडल आँख को गलत पर रख देता है या माँगा गया bar chart छोड़कर समस्त monitor बना देता है।

- ममक बेंचमार्क MMSVGBench पर यह वर्धित मजबूत प्रतद्विंद्वयिके बरस्तर है और कई मामोपर थोड़ीआगे भी नकिलतीहै — इनमें OmniSVG शामिल है, जिसे दोगुनासे अधिक डेटापर प्रशिक्षित किया गया और InternSVG, जिसे लगभग बीस गुनाअधिक डेटामिला।
- अंतर छोटे हैं। उदहरण के लिए icon set पर इसकामुख्य image-quality score 127.6 है, जबकि सर्वोत्तम प्रतद्विंद्वी का 128.8; prompt-matching score 0.293 बनाम 0.291 है। अंतर वस्तुविक और लगभग एक दशामें है, लेकिन बहुत छोटा।
- देमोजोड़े गए हसिसे उपयोगीहैं: विशेष प्रशिक्षण हटाएँ याचित्र बनते समय verification बंद करें, तोममनीय रूप से प्रदर्शन गरिताहै। Verification खस तौर पर मॉडल कोएक हीचीज़ बर-बर बनने में फँसने से रक्ताहै।
- लेखक जिस नतीजे पर सबसे अधिक जोर देते हैं, वह दक्षताहै — अग्रणीमॉडलोकीतुलनामें बहुत कम प्रशिक्षण डेटासे लगभग उसीस्तर तक पहुँचना।

यह क्यासमति नहींकरता

- यह नहींदखिताकि मॉडल को“देखने” देनाअपने-आप फलदादेताहै। उल्टा शोधपत्र काअपनाप्रयोग दखिताहै कि retraining के बनिदृश्य feedback प्रदर्शन खरस करताहै। फलदाकेवल आँखें उपलब्ध करने से नहीं उन्हें इस्तेमाल करनासखिमे से आताहै।
- यह कोई बड़ीयानर्णायक बढ़त स्थमति नहींकरता। अधिकिण scores पर वर्धित अपने प्रतद्विंद्वयिके बहुत करीब है। “20× डेटापर प्रशिक्षित मॉडल कोहरया” तकनीकीरूप से सहीहै, लेकिन एक बेंचमार्क पर proxy metrics में छोटे अंतर के आधार पर।
- यह समस्य कलत्तमक क्षमतानहीदखिता। यह आठ अरब पैरामीटर वलाशोध मॉडल है जोछोटे स्थरि resolution — 224×224 pixels — पर icons और सरल illustrations बनाताहै, कई समस्य-उद्देश्य designer नहीं।
- बड़े समस्य मॉडल GPT-5 से तुलनासम आधार वलीतुलनानहीहै: GPT-5 इस संकरे code-drawing का्य के लिए बनायायाfine-tune नहींकियागया इसलिए यहाँउसे हरमादेमोमॉडलोकीसमस्य क्षमताके बारे में बहुत कम बतताहै।
- यह मुफ्त में नहींआता। हर चरण पर canvas कोरेंडर और फरि पढ़ना सहाकोड एक हीअंधे पक्ष में नकिलने कीतुलनामें generation धीमीकरताहै — लेखक यह लगत स्वीकार करते हैं।

सक्ष्य कतिनामजबूत है

- मजबूत, जहाँतुलनानयिंतरति और स्पष्ट है। Ablation परीक्षण सक्ष हैं: विशेष प्रशिक्षण हटाएँ याverification हटाएँ, तोसंख्याएँ गरितीहैं। इसलिए देमोघटक वहीकम कर रहे हैं जिसकालेखक दमाकरते हैं।
- अपने आश्चर्यजनक परिणाम के बारे में ईममदर। बनिप्रशिक्षण के दृश्य feedback से प्रदर्शन बगिड़ने कानतीजा दबमानहीगया। यहीशोधपत्र कीसबसे दलिचस्प बतोंमें से एक है — “अधिक input हमेशाबेहतर होताहै” वलीसहज धारणाके लिए उपयोगीसुधर।
- लीडरबोर्ड दमे में सीमति। अधिक डेटावले प्रतद्विंद्वयिके बढ़त छोटीहै और एक हीबेंचमार्क पर है, जिसे उन्ही प्रतद्विंद्वयिके से एक के लेखकोने बनायाया। एक test set पर proxy metrics में छोटे margins संकेत देते हैं, अंतमि फैसलानही।
- बड़े पैमाने और वस्तुविक दुनियामें अपरीक्षति। यहाँसब कुछ कम resolution वाले icons और सरल illustrations तक सीमति है। यहीवचिर जटलि, उच्च-resolution यावस्तुविक दुनियामें design का्य में भीटकिगयानही यह भवष्य के अध्ययन कासवल है — लेखक खुद ऐसाकहते हैं।

यह क्यों मयने रखता है

इस शोधपत्र का केंद्रीय विचार लगभग शर्मनक रूप से सरल है, और चित्ररंजन से बहुत आगे जाता है। यदि कोई program कोड लिखकर कुछ बनाता है — web page, chart, diagram या 3D scene — तो वह या तो पूरा कोड अंधे की तरह लिखकर उम्मीद कर सकता है कि सब सही होगा या चलते-चलते परिणाम रेंडर करके अपनी दिशा सुधार सकता है। इसमें के लिए यह चक्र बंद करना सहज है; हम लगातार पन्ने की ओर देखते रहते हैं। मॉडलों के लिए यह आश्चर्यजनक रूप से नया कदम है।

शोधपत्र को पढ़ने लग्यक उसकी समझनी बनती है। मॉडल को देखने का सपना देना और उसे देखना सखिमा एक ही चीज नहीं है। आँखों को प्रशिक्षित करना पड़ता है; तभी यह लूप फायदा देता है। और तब भी फायदा वस्तुतः लेकनि ममाहुआ है: अधिक स्थिर चित्रकार, कोई बिल्कुल नई तरह का कलकार नहीं। ऐसे उपकरण बनने वस्तुओं के लिए जो कि सीविरण को संपन्न योग्य graphic में बदलते हैं — जैसे app icons और illustrations — शंत व्यग्रहकि सीख सबसे उपयोगी है। यहाँ बढ़त अधिक डेटा से नहीं कम लेकनि बेहतर ढंग से इस्तेमाल किए गए प्रशिक्षण से आई। यह लीडरबोर्ड पर एक और अंक से बेहतर सबक है।

संक्षेप में

वेक्टर graphics बनने वाले भ्रामा मॉडल — वे संपन्न योग्य और बनागुणवत्ता खोए बड़े-छोटे किए जा सकने वाले कोड जनिसे वेब के बहुत-से icons और logos बनते हैं — परंपरिक रूप से “अंधे” कम करते थे: सारे drawing commands लिख देते, लेकनि परिणाम देखने के लिए उन्हें कभी रेंडर नहीं करते। Guotao Liang के नेतृत्व वाली टीम Render-in-the-Loop प्रस्तुति करती है: हर चरण के बह् अधूरे चित्र को रेंडर करके मॉडल को वापस दिखाओ, तत्काल अगला stroke कैनवास देखते हुए बने। उनका सबसे महत्वपूर्ण और ईमानदार निष्कर्ष यह है कि किसी मौजूदा मॉडल पर बस यह व्यवस्था जोड़ देने से वह बदतर हो जाता है; फायदा तभी आता है जब मॉडल को दृश्य feedback इस्तेमाल करना देखा सखिमा जाए, और चित्र बनते समय ऐसा verification भी हो जो वे असर strokes हटादे। देखा प्रशिक्षित आठ अरब पैरामीटर वाला मॉडल एक ममक बेंचमार्क पर उन प्रतद्वंद्वियों के बरखर या थोड़ा आगे पहुँचता है जनिहें बीस गुना तक अधिक डेटा पर प्रशिक्षित किया गया — वस्तुतः लेकनि कम-resolution icons और सरल illustrations पर proxy scores में छोटे अंतर के सग्र। सबसे उपयोगी निष्कर्ष दक्षता का है: अपने कम कोटिक से देखना कुछ हद तक केवल डेटा का पैसा बढ़ने की जगह ले सकता है।

बनाबढ़ाचढ़कर जाँच

शोधपत्र क्या दिखाता है: वेक्टर-graphics मॉडल को हर चरण पर अपना अधूरा चित्र रेंडर करके देखने के लिए देखा प्रशिक्षित करने से अंधे चित्ररंजन की तुलना में बेहतर और अधिक पूर्ण परिणाम मिलते हैं। एक ममक बेंचमार्क पर कम प्रशिक्षण डेटा के बज्जूद यह बड़े डेटासेट पर प्रशिक्षित प्रतद्वंद्वियों के बरखर और थोड़ा आगे रहा।

क्या संभव है, पर सन्नति नहीं कि “अपना कम देखते रहना” समग्र रूप से केवल प्रशिक्षण डेटा बढ़ने से बेहतर रणनीति है; कि यही विचार जटिल, उच्च-resolution या वस्तुतः दुनिया के graphics में भी उतना ही मदद करेगा और कि बेंचमार्क के छोटे अंतर इसी आँख को स्पष्ट रूप से दिखाई देने वाला फर्क बनते हैं।

यह क्या नही दिखाता कि दृश्य feedback अकेले मदद करता है — retraining के बिना वह नुकसान करता है; कि प्रतद्वंद्वियों पर बढ़त बड़ीयानर्णयक है; समग्र-उद्देश्य चित्रकारी क्षमता या GPT-5 जैसे समग्र मॉडलों के सग्र निष्पक्ष समतल-काय तुलना।

मुख्य सीमाएँ: एक बेंचमार्क, जिसे आंशिक रूप से एक प्रतद्विंद्वीके लेखकोने बनाया proxy metrics में छोटे अंतर; आठ अरब पैरामीटर वज़ामॉडल; 224×224 resolution के icons और सरल illustrations; और हर चरण पर render करने के कारण धीमी generation।

समस्य पष्ठक कोकतिनाभरेसारखनाचहएि? इस वक्त पर उच्च क्चित्र बनते समय परणित कोबा-बा रेंडर करके देखा देखनासचमुच मदद करताहै — लेकिन तभीजब मॉडल कोइसकाउपयोग करनाप्रशक्षित कियाजाए। इस खस मॉडल के अपने प्रतद्विंद्वियोंसे निर्णायक रूप से बेहतर हमें पर कम से मध्यम भरेसारखें; “ $20 \times$ डेटावाले मॉडल कोहरघा वस्तवकि लेकिन छोटा एक-बेंचमार्क परणित है। यह रखने लग्यक मुख्य वचि पर भरेसाअधिक है: जोकोइ चित्र बनता है, उसके लिए चलते-चलते देखनाअंधे होकर पूराचित्र बनने से बेहतर है।

स्रोत

Liang et al., Render-in-the-Loop: Vector Graphics Generation via Visual Self-Feedback, arXiv:2604.20730 (2026)
<https://arxiv.org/abs/2604.20730>

OmniSVG project page
<https://omnisvg.github.io/>

InternSVG project page
<https://hmwang2002.github.io/release/internsvg/>