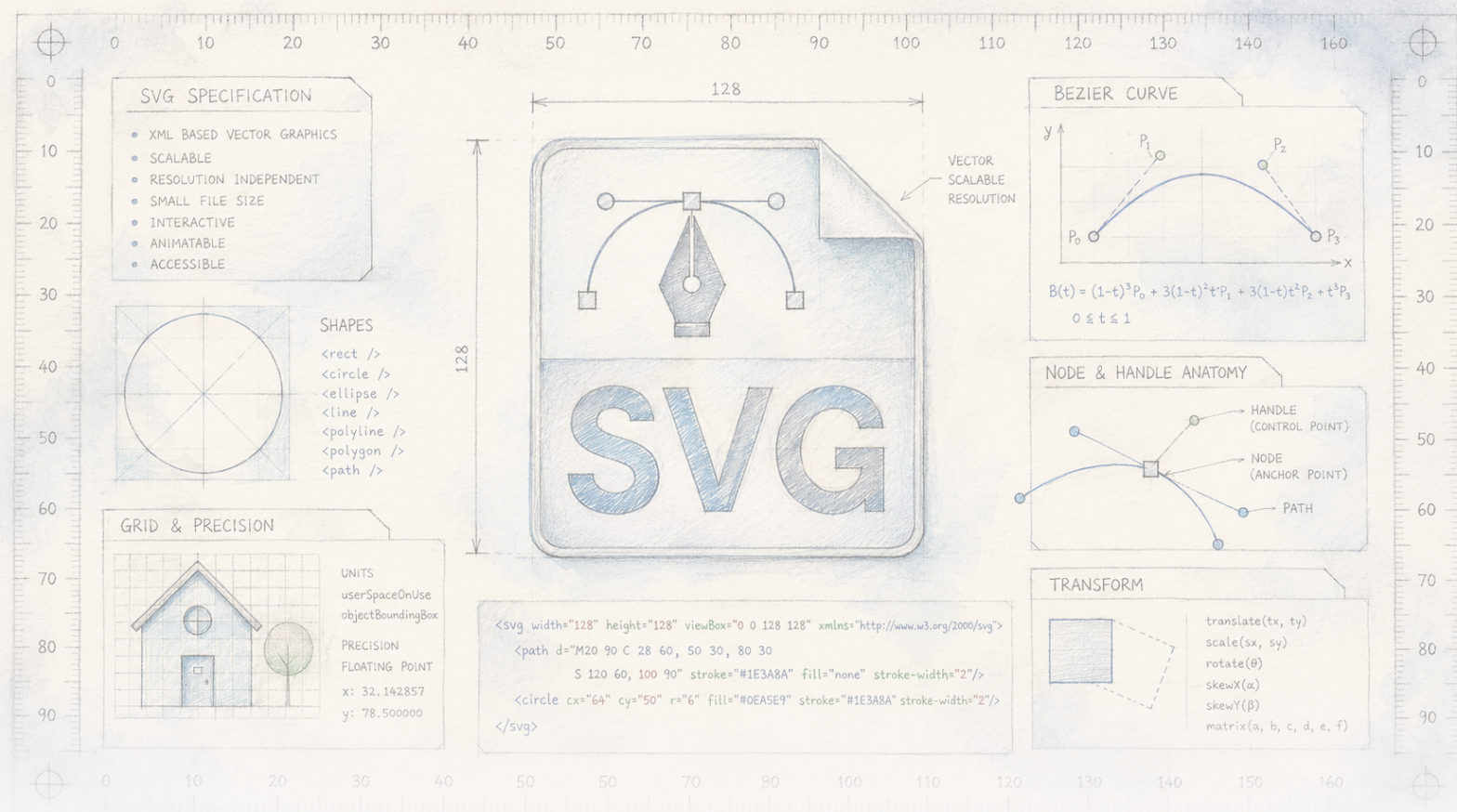




The CLEAN PAPER

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

ስዕል የሚሰል AI ገጹን እንዲመለከት ማስተማር

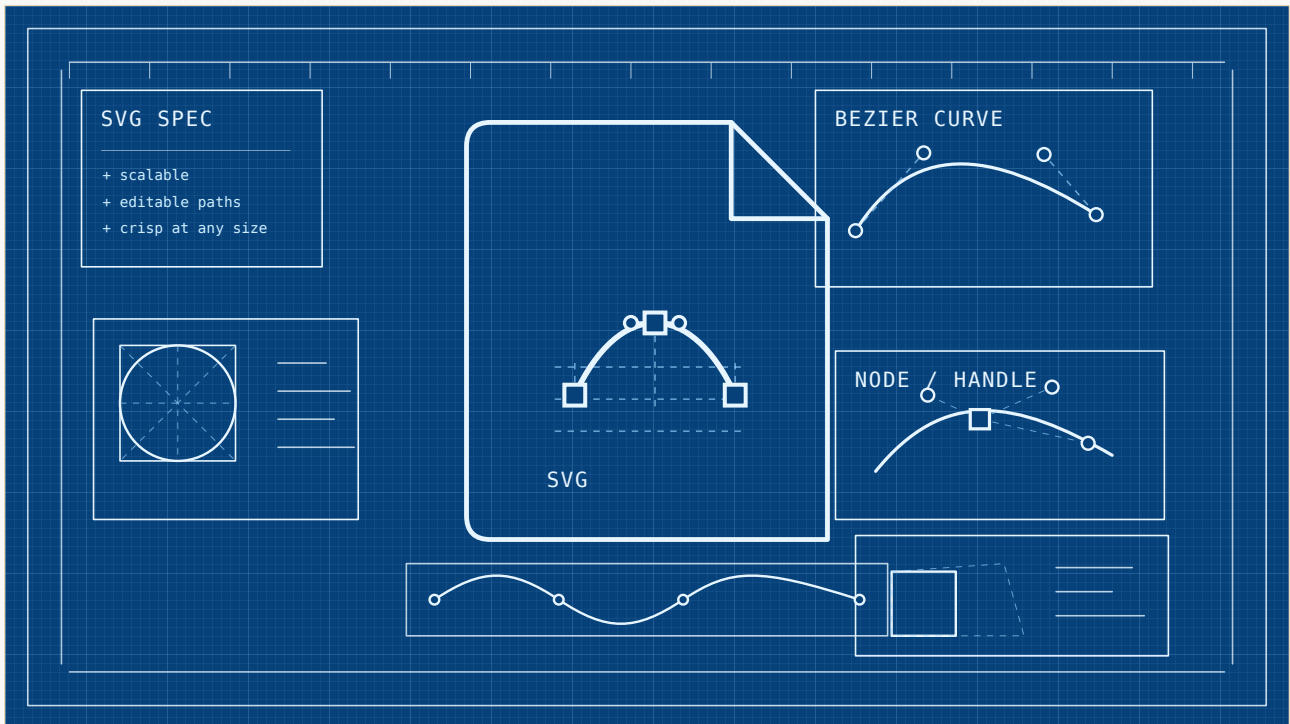
Vector graphics የሚፈጥሩ language models በተለምዶ ውጤታቸውን ሳያዩ የመሳሉ ትዕዛዞችን ይጽፋሉ። አዲስ ዘዴ ሞዴሉ እያንዳንዱን stroke ከጨመረ በኋላ canvas-ውን እንዲያይ ያደርጋል። ግን በግልጽ የቀረበው ያልተጠበቀ ውጤት ይህ ነው፡- ዓይን መስጠት ብቻ ጥራትን ያበላሻል፤ visual feedback-ን ለመጠቀም እንደገና ማሰልጠን ያስፈልጋል። ከዚያ በኋላ ሞዴሉ እስከ 20 እጥፍ በሚበልጥ ውሂብ ከሰለጠነ ተፎካካሪዎች ጋር ይመሳሰላል ወይም ትንሽ ይበልጣል — በአንድ benchmark ላይ ያለ እውነተኛ ግን መጠነኛ ውጤት፤ አብዮት አይደለም።

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)

ዓይንዎን ጨፍነው መሳል

ዛሬ ካሉት AI ሞዴሎች አንዱን ምስል እንዲስልጥ ቢጠይቁ፣ በውስጡ ከሁለት በጣም የተለያዩ ነገሮች አንዱ ይከሰታል። የተለመዱት የምስል ጀነሬተሮች — ከአንድ ዓረፍተ ነገር ፎቶ የሚፈጥሩት — pixelsን በቀጥታ ይሳሉ፣ እና ሲሠሩም ካንቫሱን ማየት ይችላሉ። ግን ሌላ፣ ዝምተኛ የመሳል መንገድ አለ፤ በዚያ ሞዴሉ ምንም አይቀባም። መመሪያ ይጽፋል፡- እዚህ ክብሩ ሳል፣ እዚያ መስመር አድርግ፣ ይህን ቅርጽ በሰማያዊ ሙሉ። እነዚህ መመሪያዎች ኮድ ናቸው — በድር ላይ ከሚታዩ አብዛኞቹ icons እና logos በስተጀርባ ያለው Scalable Vector Graphics፣ ወይም SVG። ጥቅሙም ግልጽ ነው፣ ውጤቱ ቋሚ የpixel ፍርግርግ ሳይሆን ሳይደበዝዝ ደጋግመው መጠኑን ማስተካከል፣ ቀለሙን መቀየር እና ማርትዕ የሚችሉት የቅርጾች ስብስብ ነው።



ዋናው SVG blueprint ሥዕል፤ ምስሉ ራሱ ሊስተካከል የሚችል vector file ነው፣ ከቋሚ pixels ይልቅ paths፣ ፍርግርግ lines፣ nodes እና handles የተገነባ።

Laura Nesso / The Clean Paper Permission on file

ችግሩ እስከ ቅርብ ጊዜ ድረስ ይህን የመሳል ኮድ የሚጽፍ ሞዴል በዓይነ ስውርነት እንደሚሠራ ነበር። መመሪያዎቹን — ክብ፣ መስመር፣ ቀለም ሙሉ — በአንድ ጊዜ ሙሉ በሙሉ ያወጣ ነበር፤ ምን እንደሠራ ለማየት አንድ ጊዜም ምስል መፍጠር ሳያደርግ። ዓይንዎን ጨፍነው ፊት ለመሳል እንደሚሞክሩ ያስቡ፤ ዓይኖቹን በግምት ትክክለኛ ቦታ ላይ ሊያኖሩ ይችላሉ፣ ግን ሁለተኛው ዓይን ጉንጭ ላይ መውደቁን ወይም ለፀጉር የሳሉት ቅርጽ አፍንጫውን መሸፈኑን ማወቅ አይችሉም። እነዚህ ሞዴሎችም በግምት እንዲህ ይሠሩ ነበር፤ ስለዚህ በኮድ ደረጃ ፍጹም ትክክል ሆኖ በዓይን ግን የተዘበራረቀ ውጤት ብዙ ጊዜ ይፈጥሩ ነበር።

በGuotao Liang የሚመራ ቡድን አሁን በጣም ግልጽ የሚመስለውን ነገር አድርጓል፤ ሞዴሉ ዓይኑን እንዲከፍት አደረጉት። የእነሱ ምስል መፍጠር-in-the-ዑደት ዘዴ እያንዳንዱ ደረጃ ከተጠናቀቀ በኋላ ግማሽ የተሠራውን ሥዕል ምስል መፍጠር ያደርጋል እና ሞዴሉ ቀጣዩን stroke ከመሳሉ በፊት ያንን ምስል ይመልስለታል። በእውነት አስደሳቹ ነገር ይህ መርዳቱ አይደለም — እንዲረዳ ለማድረግ ምን ማድረግ እንዳስፈለጋቸው ነው።

ሥዕልን እንዴት ነው ነጥብ የምንሰጠው?

አንድ ወረቀት ሞዴሉ ሌላን “አሸንፏል” ሲል፣ በምን ነገር እና በምን መለኪያ እንደሆነ መጠየቅ ተገቢ ነው። የተፈጠረ ምስልን መገምገም በእውነት ከባድ ነው — “laptop ሳል” ለሚለው አንድ ብቻ ትክክለኛ መልስ የለም — ስለዚህ መስኩ ሰው ውጤቱን በዓይኑ እንደሚገመግም ፍጹም ማካካሻ ያልሆኑ ጥቂት automatic scores ላይ ይተማመናል።

በዚህ ወረቀት ውስጥ ጥቂቶቹ ይገኛሉ። FID ብዙ የተፈጠሩ ምስሎችን አጠቃላይ ስታቲስቲካዊ ባህሪ ከእውነተኛ ምስሎች ጋር ያነጻጽራል፤ ዝቅ ያለ እሴት ይሻላል፤ ግን የሚገልጸው ቡድኑን እንጂ አንድን ምስል አይደለም። CLIP ነጥብ ሌላ AI ምስሉ hprompt-ው ቃላት ጋር ይዛመዳል ወይ? ብሎ እንዲፈርድ ያደርጋል — ጠቃሚ ነው፤ ግን የዳኛውን ብቃት አያልፍም። DINO፣ SSIM እና LPIPS ደግሞ እንደገና የተገነባ ምስልን ከዲላማ ጋር፣ ከቀጥታ pixel ደረጃ እስከ ተማሩ ባህሪያት ድረስ ያነጻጽራሉ።

ከእነዚህ አንዱም “እውነት” አይደለም። Proxies ናቸው፣ እና ብዙውን ጊዜ በትንሽ መጠን ይንቀሳቀሳሉ። አንድ ሞዴል 127.6 ሲያገኝ ሌላው 128.8 እንደሚያገኝ ሲያነቡ፣ ይህ በታሰበው አቅጣጫ ያለ እውነተኛ ልዩነት ነው፤ ግን ትንሽ መግፊያ እንጂ ታላቅ ልዩነት አይደለም፤ እና ያ ቁጥር ዓይነት ከሚፈርደው ጋር በቀጥታ አይጓዝም። “20 እጥፍ በሚበልጥ ውሂብ የሰለጠነ ሞዴሉን አሸነፈ” የሚል ርዕስ ሲያዩ ይህን ማስታወስ ጠቃሚ ነው።

ደራሲዎቹ ምን አደረጉ?

ደራሲዎቹ ለመፍታት የፈለጉት ችግር ይህ የዓይነ ስውር መሳል ራሱ ነው። አሁን ያሉ ሞዴሎች SVG መጻፍን እንደ ንጹሕ የጽሑፍ ሥራ ይይዛሉ፤ እስከዚያ የተጻፈውን ኮድ ተመልክተው ቀጣዩን ክፍል ይገምታሉ፤ ነገር ግን ምስል መፍጠር አያደርጉም። ይህም ዘመናዊ multimodal ሞዴሎች ቀድሞውኑ ያላቸውን ኃይለኛ “ዓይኖች” — vision encoder — ሥራ ፈት ያደርጋቸዋል። ምስል መፍጠር-in-the-ቦደት ሥራውን ደረጃ በደረጃ የሚከናወን የእይታ ሂደት አድርጎ ያዋቅረዋል። ከእያንዳንዱ የdrawing ኮድ ክፍል በኋላ፣ ከፊል SVG-ው ወደ ምስል ምስል መፍጠር ይደረጋል እና እንደ ምስል ወደ ሞዴሉ ይመለሳል፤ ስለዚህ ቀጣዩ ክፍል የሚመረጠው እስከዚያ የተሠራውን canvas በእውነት እያየ ነው።

የመጀመሪያ ውጤታቸው ጥንቃቄ የሚያስፈልግ ነው፤ እና ይህን በግልጽ ማቅረባቸው የሚመስገን ነው፤ ይህን ቦደት በቀጥታ በአሁኑ off-the-shelf ሞዴል ላይ መጨመር አይሠራም። ልዩ ሥልጠና ሳይሰጡ መካከለኛ rendersን ለጠንካራ አጠቃላይ ሞዴሎች ሲሰጡ፣ ጥራት አልተሻሻለም — በሁሉም ልኬቶች ተበላሽ። ዓይኑን በዚህ ሥራ ለመጠቀም ተምሮ የማያውቅ ሞዴል፣ ድንገት ብቻውን እንዴት እንደሚጠቀምበት አያውቅም።

ስለዚህ የሥራው አብዛኛው ክፍል ማስተማር ላይ ነው። የስልጠና መረጃውን እንደገና ያዋቅራሉ፤ እያንዳንዱን ሥዕል በዓይን ትርጉም ያላቸው ብዙ ትናንሽ ደረጃዎች ይከፍላሉ — ውስብስብ ቅርጾችን ወደ ቀላል ክፍሎች በመከፋፈል በእያንዳንዱ ደረጃ አዲስ የሚታይ ነገር እንዲኖር — ከዚያም በQwen3-VL ላይ የተገነባ 8-billion-መለኪያ ያልተፈታ ሞዴል በእነዚህ ደረጃ-በደረጃ sequences ላይ fine-tune ያደርጋሉ። ይህን የእይታ Self-Feedback ብለው ይጠሩታል። በመሳል ጊዜ ሁለተኛ መንገድ፣ ምስል መፍጠር-and-Verify፣ ይጨምራሉ፤ ሞዴሉ አዲስ stroke ከመቀበሉ በፊት ምስል መፍጠር ያደርገውና በእውነት ምስሉን ቀይሮታል ወይስ ያለፈውን ብቻ ደግሟል ብሎ ይመረምራል። ምንም የማይጨምሩ strokes ይጣላሉ፤ ከዚያ በላይ ምንም ሲያሻሽል ሞዴሉ እንዲቆም ይነገረዋል። ይህ ሁሉ በአንጻራዊነት ትንሽ የመረጃ ስብስብ — ወደ 850,000 ምሳሌዎች፣ አንድ ተፎካካሪ ከተጠቀመበት ከግማሽ በታች እና ከሌላው በጣም ትንሽ ክፍል — ላይ መሠራቱ የሚጠቀስ ነው።

ምን አገኙ?

- በዚህ መንገድ ከሰለጠነ በኋላ ሞዴሉ ከዓይነ ስውሩ ስሪት ይልቅ የተሻለ ይሳላል — በተለይ በውድቀት ጉዳዮች ላይ፤ ዓይነ ስውሩ ሞዴል ዓይነን ጉንጭ ላይ ሲያስቀምጥ ወይም የተጠየቀውን bar chart ትቶ generic monitor ሲያወጣ ልዩነቱ ግልጽ ነው።
- በመደበኛው መመዘኛ (MMSVGBench) ላይ፤ ዘዴው ከጠንካራ ተፎካካሪዎች ጋር ተወዳዳሪ ነው፤ በብዙ ልኬቶችም ትንሽ ይበልጣል — OmniSVG ከሁለት እጥፍ በላይ ውሂብ ላይ የሰለጠነ ሲሆን፤ InternSVG ወደ 20 እጥፍ በሚጠጋ ውሂብ ላይ ስልጥፏል።
- ልዩነቶቹ ትንሽ ናቸው። ለምሳሌ nicon ስብስብ ላይ፤ ዋናው Fimage-quality ነጥብ 127.6 ነው፤ የምርጥ ተፎካካሪው 128.8፤ prompt-matching ነጥብ ደግሞ 0.293 ከ0.291 ጋር ነው — እውነተኛና የተደጋጋሚ ልዩነቶች፤ ግን ቀጭን።
- የተጨመሩት ሁለቱም ክፍሎች ጠቃሚ ናቸው፤ ልዩ ስልጠናውን ወይም በመሳል ጊዜ verification-ውን ካጠፉ፤ ቁጥሮቹ በሚለካ ሁኔታ ይወርዳሉ። በተለይ verification-ው ሞዴሉ አንድን ነገር ደጋግሞ በመሳል እንዳይጣበቅ ያደርጋል።
- ደራሲዎቹ ራሳቸው የሚያጎሉት ውጤት efficiency ነው — መሪ ሞዴሎች ከተጠቀሙበት በጣም ያነሰ ስልጠና መረጃ ተጠቅመው ይህን ደረጃ መድረሳቸው።

ይህ ምን አያረጋግጥም?

- ሞዴል “እንዲያይ” መፍቀድ ነፃ የሆነ ድል ነው ብሎ አያሳይም። በተቃራኒው፤ የወረቀቱ ሙከራ ያለ retraining የእይታ feedback ጥራትን እንደሚያበላሽ ያሳያል። ጥቅሙ የሚመጣው ከretraining ነው፤ ከ“ዓይኖቹ” ብቻ አይደለም።
- ትልቅ ወይም ወሳኝ መሪነት እንዳለው አያረጋግጥም። በአብዛኛዎቹ scores ዘዴው ከተፎካካሪዎች ጋር እጅግ ቅርብ ነው፤ “20× ውሂብ ላይ የሰለጠነ ሞዴሉን ያሸንፋል” የሚለው እውነት ቢሆንም፤ በአንድ መመዘኛ ላይ proxy metrics ያሳዩት ትንሽ ልዩነት ነው።
- አጠቃላይ የስነ-ጥበብ ችሎታን አያሳይም። ይህ 8-billion-መለኪያ የምርምር ሞዴል በትንሽ ቋሚ resolution (224×224 pixels) icons እና ቀላል ምሳሌዎችን የሚስል ነው፤ አጠቃላይ-purpose designer አይደለም።
- ከትልቅ አጠቃላይ ሞዴል (GPT-5) ጋር ያለው ንጽጽር ተመሳሳይ ሁኔታ ያለው አይደለም፤ ያ ሞዴል ለዚህ ጠባብ ኮድ-drawing ተግባር አልተገነባም እና አልተስተካከለም፤ ስለዚህ እዚህ መሸነፉ ስለ ሁለቱም ሞዴሎች በአጠቃላይ ብዙ አይናገርም።
- ነፃ ወጪ የሌለው ዘዴ አይደለም። በእያንዳንዱ ደረጃ ካንቫሱን ምስል መፍጠር ማድረግ እና እንደገና ማንበብ፤ ኮድ-ውን በአንድ ጊዜ በዓይነ ስውርነት ከማውጣት ይልቅ generation-ን ያዘገያል፤ ደራሲዎቹም ይህን ይቀበላሉ።

ማስረጃው ምን ያህል ጠንካራ ነው?

- ጠንካራ ነው፤ በራሱ የሙከራ ወሰን ውስጥ በቁጥጥር የተደረገ ንጽጽር ሲሆን። Ablations ግልጽ ናቸው፤ ስልጠናውን ወይም verification-ውን ሲያስወግዱ ቁጥሮቹ ይወርዳሉ፤ ስለዚህ ሁለቱ ክፍሎች ደራሲዎቹ እንደሚሉት እውነት ሥራ እየሠሩ ናቸው።
- ያልተጠበቀውን ውጤት በግልጽ ያቀርባል። Naive የእይታ feedback ይጎዳል የሚለው ውጤት አልተደበቀም፤ በወረቀቱ ውስጥ ከሚያስደስቱት ነገሮች አንዱ ነው፤ እና “ብዙ input ሁልጊዜ የተሻለ ነው” የሚለውን ስሜት

ያስተካክላል።

- **Leaderboard አባባል ላይ ግን ቀጭን ነው።** ከብዙ ውሂብ በሰለጠኑ ተፎካካሪዎች ላይ ያሉት ድሎች ትንሽ ናቸው፤ እና በአንድ መመዘኛ ላይ ብቻ ይኖራሉ፤ ያ መመዘኛ ደግሞ በአንዱ ተፎካካሪ ደራሲዎች ተገንብቷል። በአንድ መከራ ስብስብ ላይ proxy metrics ያሳዩት ትንሽ margins ጠቋሚ ናቸው፤ የመጨረሻ ፍርድ አይደሉም።
- **በትልቅ ሚዛን እና በእውነተኛ ዓለም አልተፈተነም።** እዚህ ያሉት ሁሉ icons እና ቀላል ምሳሌዎች በዝቅተኛ resolution ናቸው። ተመሳሳይ ሐሳብ ለውስብስብ፤ high-resolution ወይም የእውነተኛ ዓለም ንድፍ ሥራ ይሠራ ወይ? የሚለው ለወደፊት ሥራ ተትቷል፤ ደራሲዎቹም ይህን በግልጽ ይላሉ።

ለምን አስፈላጊ ነው?

የዚህ ወረቀት ማዕከላዊ ሐሳብ እጅግ ቀላል ነው፤ ግን ከመሳል በላይ ይሄዳል። አንድ program ኮድ በመጻፍ ነገር ሊፈጥር ከሆነ — web page፣ chart፣ diagram፣ 3D scene — ሁሉን በዓይነ ስውርነት ጽፎ ተስፋ ሊያደርግ ይችላል፤ ወይም እየሠራ ምስል መፍጠር አድርጎ መንገዱን ሊያስተካክል ይችላል። ይህን ዑደት መዝጋት ለሰው ተፈጥሯዊ ነው፤ ገጹን ደጋግመን እንመለከታለን። ለእነዚህ ሞዴሎች ግን አስገራሚ በሆነ መጠን አዲስ እርምጃ ነው።

ወረቀቱን ለማንበብ የሚያስቆጥረው ያለው ማስታወሻ ነው። ሞዴሉን ማየት እንዲችል ማድረግ፤ ማየትን ማስተማር ጋር አንድ አይደለም። “ዓይኖቼ” ሊሰለጥኑ ይገባል፤ ከዚያ ብቻ ዑደት-ው ጥቅም ያመጣል። እና እንኳን ከዚያ በኋላ ጥቅሙ እውነተኛ ቢሆንም የተመጠነ ነው፤ የተረጋጋ ሰዓሊ እንጂ ፍጹም አዲስ ዓይነት አርቲስት አይደለም። ከdescription ወደ editable graphic የሚቀይሩ መሣሪያዎችን — ለapp icons እና illustrations የሚያገለግሉትን — ለሚገነባ ሰው፤ ጸጥ ያለው ነገር ግን ተግባራዊው ትምህርት ይህ ነው። እዚህ ያለው ድል ከብዙ መረጃ አልመጣም፤ ከርካሽና ብልህ ስልጠና መጣ። ይህ ከleaderboard ላይ ተጨማሪ አንድ ነጥብ ይልቅ የተሻለ ትምህርት ነው።

ንጹህ ማጠቃለያ

Vector graphics — በአብዛኛዎቹ web icons እና logos በስተጀርባ ያለው፤ ሊስተካከልና መጠኑ ሊቀየር የሚችል ኮድ — የሚፈጥሩ ቋንቋ ሞዴሎች በባህላዊ መንገድ “በዓይነ ስውርነት” ይሠሩ ነበር፤ ሁሉንም የመሳል መመሪያዎች ይጽፉ ነበር፤ ግን ውጤቱን ምስል መፍጠር አድርገው አያዩም። በGuotao Liang የሚመራ ቡድን ምስል መፍጠር-in-the-ዑደት የተባለ ዘዴ አቀረበ፤ ከእያንዳንዱ ደረጃ በኋላ ግማሽ የተሠራውን ሥዕል ምስል መፍጠር አድርጎ ወደ ሞዴሉ ይመልሳል፤ ስለዚህ ቀጣዩን stroke የሚሰለው ካንቫስን እያየ ነው። ዋናውና በግልጽ የቀረበው ውጤት ይህን በቀጥታ ለነባር ሞዴል መጨመር ጥራትን እንደሚያበላሽ ነው፤ ጥቅሙ የሚታየው ሞዴሉ የእይታ feedback-ን ለመጠቀም እንደገና ከሰለጠነ በኋላ፤ ምንም የማይቀይሩ strokesን በሚጥል drawing-ጊዜ check በመታገዝ ነው። የተሰለጠነው 8-billion-መለኪያ ሞዴል በአንድ መደበኛ መመዘኛ ላይ እስከ 20 እጥፍ በሚበልጥ ውሂብ ከሰለጠኑ ተፎካካሪዎች ጋር ይመሳሰላል ወይም ትንሽ ይበልጣል — እውነተኛ ውጤት ነው። ግን በproxy scores ላይ ያለ ትንሽ margin፤ ለicons እና ቀላል ዝቅተኛ-resolution ምሳሌዎች ነው። ደራሲዎቹ የሚያነሱት፤ እኛም ማስታወስ ያለብን፤ efficiency ነው፤ ሥራውን በደንብ ማየት አንዳንድ ጊዜ ብዙ መረጃ ለመጨመር ማካካሻ ሊሆን ይችላል።

ያለ ማጋነን ፍተሻ

ወረቀቱ የሚያሳየው፦ vector-graphics ሞዴልን የግማሽ የተሠራውን ሥዕል ደረጃ በደረጃ ምስል መፍጠር አድርጎ እንዲመለከት እንደገና ማሰልጠን፤ በዓይነ ስውርነት ከመሳል የተሻለና የተሟላ ውጤት ያመጣል። በአንድ መደበኛ መመዘኛ ላይ ከብዙ መረጃ የሰለጠኑ ተፎካካሪዎች ጋር ይወዳደራል ወይም ትንሽ ይበልጣል፤ እና ይህን ከያነሰ ስልጠና መረጃ ጋር ያደርጋል።

ምክንያታዊ ነገር ግን ያልተረጋገጠው፦ “ሥራህን እያየህ መሥራት” በአጠቃላይ መረጃ ሚዛን ከመጨመር የተሻለ መንገድ መሆኑ፤ ተመሳሳይ ሐሳብ ለውስብስብ፤ high-resolution ወይም የእውነተኛ ዓለም graphics መርዳቱ፤ ትንሹ መመዘኛ margins ሰው በዓይኑ የሚያስተውለውን ልዩነት መወከላቸው።

የማያሳየው፦ የእይታ feedback ብቻውን እንደሚረዳ — ያለ retraining ይጎዳል፤ ከተፎካካሪዎች ላይ ትልቅና ወሳኝ መሪነት፤ general-purpose drawing ability፤ እንደ GPT-5 ከሚሉ አጠቃላይ ሞዴሎች ጋር ፍትሃዊ head-to-head ገጽጽር፤ ምክንያቱም እነሱ ለዚህ ሥራ አልተገነቡም።

ዋና ገደቦች፦ አንድ መመዘኛ፤ ከፊሉ በአንድ ተፎካካሪ ደረሲዎች የተገነባ፤ በproxy metrics ላይ ትንሽ margins፤ 8B ሞዴል፤ icons እና ቀላል illustrations በ224×224፤ እያንዳንዱን ደረጃ ምስል መፍጠር ስለሚያደርግ የዘገየ generation።

አጠቃላይ አንባቢ ምን ያህል እምነት ሊኖረው ይገባል? ዑደት-ውን መዝጋት — እየሳሉ ምስል መፍጠር ማድረግና እንደገና ማየት — በእውነት እንደሚረዳ፤ እና ይህ ነገር ማሰልጠን እንጂ switch ብቻ መክፈት እንዳልሆነ ከፍተኛ እምነት። ይህ የተወሰነ ሞዴል ተፎካካሪዎቹን በወሳኝ ሁኔታ ያሸንፋል በሚለው ላይ ዝቅተኛ እስከ መጠነኛ እምነት፤ “20× መረጃ የሰለጠነውን አሸነፈ” የሚለውን እውነተኛ ግን መጠነኛ፤ single-መመዘኛ ውጤት አድርገው ይዩት። ማስታወስ ያለብዎት አንድ ሐሳብ ላይ ግን ከፍተኛ እምነት፦ ለሚስል ኮድ፤ እየሠራ ማየት በዓይነ ስውርነት ከመሳል ይሻላል።

ምንጮች

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/>