

TOWARDS AI-NATIVE SOFTWARE ENGINEERING (SE 3.0): A VISION AND A CHALLENGE ROADMAP

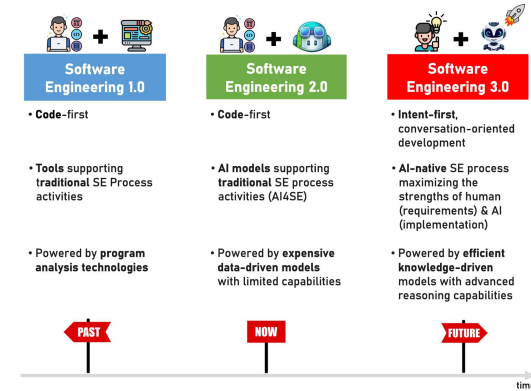
Presenter: Haonan Zhang

January 16, 2025



1

Different stages of SE



2

Significant cognitive overload on the human

```
1 #!/usr/bin/env ts-node
2
3 import { fetch } from "fetch-h2";
4
5 // Determine whether the sentiment of text is positive
6 // Use a web service
7 async function isPositive(text: string): Promise<boolean> {
8   const response = await fetch("http://text-processing.com/api/sentiment/", {
9     method: "POST",
10    body: "text=${text}",
11    headers: {
12      "Content-Type": "application/x-www-form-urlencoded",
13    },
14  });
15  const json = await response.json();
16  return json.label === "pos";
17 }
```

Source: GitHub Copilot

3

Inefficient model training on domain-specific tasks

An example of the performance of FMs used for domain-specific tasks

Model	Cross-Task				Cross-Website				Cross-Domain			
	Ele.	Acc.	Op. F1	Step SR	Ele.	Acc.	Op. F1	Step SR	Ele.	Acc.	Op. F1	Step SR
Supervised Fine-Tuning												
FLAN-T5												
- Base	40.5	74.4	37.3	28.7	69.6	27.9	38.2	69.1	36.2			
- Large	52.2	70.7	48.8	35.3	65.8	32.7	41.9	64.6	39.5			
- XL	56.8	74.6	52.5	42.6	69.9	39.5	43.8	65.2	40.7			
BLIP2-T5												
- Base	39.5	74.9	36.1	34.0	70.8	32.2	38.2	72.8	37.5			
- Large	50.0	72.1	46.0	39.5	71.5	36.3	40.9	70.1	39.4			
- XL	52.9	74.9	50.3	41.7	74.1	38.3	43.8	73.4	39.6			
In-Context Learning												
GPT-3.5	19.4	59.8	16.8	14.9	56.5	14.1	25.5	57.9	24.2			
GPT-4	40.2	63.4	31.7	27.4	61.0	27.0	36.2	61.9	29.7			
SEA-CT												
- Attributes	4.7	39.5	4.7	9.7	37.8	9.7	16.0	41.4	15.3			
- Annotation	15.1	66.5	13.0	11.3	63.4	10.5	16.5	65.1	14.7			
- Choice	48.9	69.1	40.6	48.5	70.6	41.7	44.0	70.9	40.9			
- Oracle	72.9	80.9	65.7	74.4	83.7	70.0	72.8	73.6	62.1			

“The most effective grounding strategies we explored in this paper still exhibit a 20-25% performance gap compared to oracle grounding”

Boyuan Zheng, Boyu Gao, Jihyung KIL, Huan Sun, and Yu Su. 2023. GPT-4V(Vision) is a generalist web-agent, if grounded. In Proceedings of the 41st International Conference on Machine Learning (ICML'24)

4

Unaffordable resource usage

It is very expensive to fine-tune the models for improved domain-specific capability

“We predict that to achieve 95% accuracy for in-domain low-level tasks, 1M episodes would be required, while 2M episodes would be required to obtain 95% episode completion rates for 5-step high-level tasks”

Li, Wei, et al. On the Effects of Data Scale on Computer Control Agents. In Proceedings of the 38th Conference on Neural Information Processing Systems (NeurIPS24)

5

Suboptimal code quality produced by FMs

Task Description	
Write a function to merge two sorted arrays into a single sorted array. Input: Two arrays, array1 and array2, each sorted in non-decreasing order. Output: A single array that combines the elements of array1 and array2, sorted in non-decreasing order. Example: Given array1 = [2, 7, 11, 15], array2 = [5, 8, 15, 17], return [2, 5, 7, 8, 11, 15, 15, 17].	
VS Code Copilot v1.156.683	GPT-4
<pre>def merge_arrays(array1, array2): # Initially, concatenate the two # arrays merged_array = [] for element in array1: merged_array.append(element) for element in array2: merged_array.append(element) # Sort the merged array n = len(merged_array) for i in range(n): for j in range(i+1, n): if merged_array[i] > merged_array[j]: merged_array[i], merged_array[j] = merged_array[j], merged_array[i] return merged_array</pre>	<pre>def merge_arrays(array1, array2): # Merging two arrays efficiently result = [] i, j = 0, 0 while i < len(array1) and j < len(array2): if array1[i] < array2[j]: result.append(array1[i]) i += 1 else: result.append(array2[j]) j += 1 # Append remaining elements of array1 while i < len(array1): result.append(array1[i]) i += 1 # Append remaining elements of array2 while j < len(array2): result.append(array2[j]) j += 1 return result</pre>

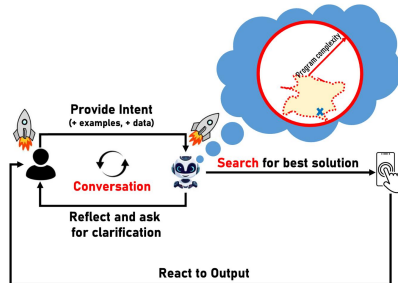
Figure 1: Example codes with distinct time complexity generated by Copilot and GPT-4, respectively. Code accessed on January 15, 2024.

“While functionally correct, this approach suffers from sub-optimal time complexity and space complexity”

Dong HUANG, Yuhao QING, Weiyl Shang, Heming Cai, and Jie Zhang, EFBench: Benchmarking the Efficiency of Automatically Generated Code. In Proceedings of the 38th Conference on Neural Information Processing Systems (NeurIPS24)

6

Vision of Software Engineering 3.0

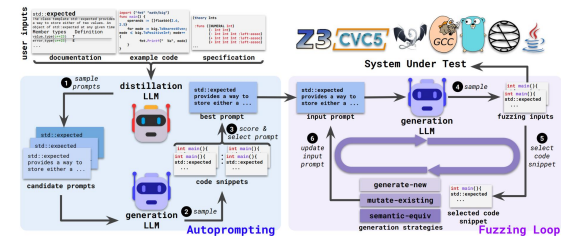


Development is no longer driven by code but by intents expressed through back-and-forth conversations between human developers and their AI teammates.

7

What fuzzing looks like in SE 3.0

Fuzz4ALL: Universal Fuzzing with Large Language Models		
Chunqiu Steven Xia University of Illinois Urbana-Champaign, USA xiaqc@illinois.edu	Matteo Paltenghi University of Stuttgart, Germany matt.paltenghi@stuttgart.de	Jia Le Tian University of Illinois Urbana-Champaign, USA jia@illinois.edu
Michael Pydel University of Stuttgart, Germany michael.pydel@stuttgart.de	Lingming Zhang University of Illinois Urbana-Champaign, USA lingmingz@illinois.edu	

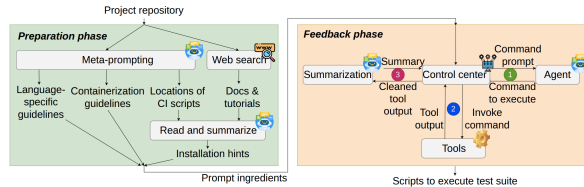


8

What test execution looks like in SE 3.0

You Name It, I Run It: An LLM Agent to Execute Tests of Arbitrary Projects

ISLEM BOUZENIA¹, University of Stuttgart, Germany
MICHAEL PRADEL¹, University of Stuttgart, Germany

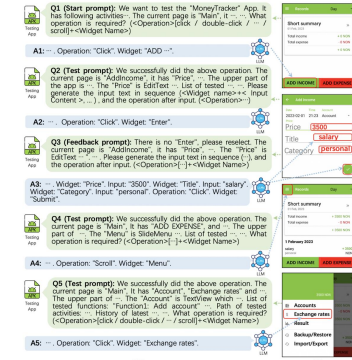


9

What GUI testing looks like in SE 3.0

Make LLM a Testing Expert: Bringing Human-like Interaction to Mobile GUI Testing via Functionality-aware Decisions

Zhe Liu^{1,2}, Chongyang Chen¹, Junjie Wang^{1,2}, Mengshao Chen^{1,2}, Boyu Wu^{1,4},
Xing Chen¹, Dandan Wang^{1,2}, Qing Wang^{1,2,5},
¹State Key Laboratory of Intelligent Games, Beijing, China
²Institute of Software Chinese Academy of Sciences, Beijing, China
³University of Chinese Academy of Sciences, Beijing, China
⁴Corresponding author: liuzhe@pku.edu.cn
⁵Science & Technology on Integrated Information System Laboratory
liuzhe@pku.edu.cn, chencc@pku.edu.cn, chencc@pku.edu.cn, chencc@pku.edu.cn



10

What GUI testing looks like in SE 3.0

Video Demo

Source: <https://osai-nlp-group.github.io/Seckit/>

11

Balance between ask too many and not asking enough

Optimizing Instructions and Demonstrations for Multi-Stage Language Model Programs

Krista Opsahl-Ong^{1*}, Michael J Ryan^{1*}, Josh Purtell²,
David Broman³, Christopher Potts⁴, Matei Zaharia¹, Omar Khattab¹
¹Stanford University, ²Basis, ³KTH Royal Institute of Technology ⁴UC Berkeley

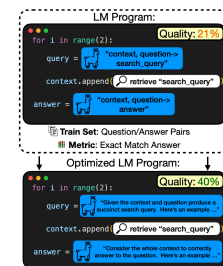


Figure 1: An example of the optimization problem we explore, shown for a multi-hop retrieval LM program. Given some question-answer pairs and a metric, the optimizer proposes new instructions and bootstraps new demonstrations (not pictured) for each stage.

12

Improving the efficiency of code synthesis

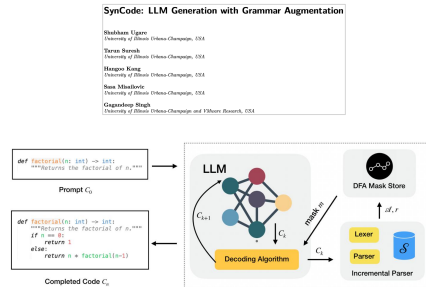


Figure 1: In the SynCode workflow, the LLM takes partial output C_k and generates a distribution for the next token t_{k+1} . The parser processes C_k to produce accept sequences $\mathcal{A}$ and remainder r . These values are used by the DFA mask store to create a token mask, eliminating syntactically invalid tokens. The LLM iteratively generates a token t_{k+1} using the distribution and the mask, appending it to C_k to create the updated code C_{k+1} . The process continues until the LLM returns the final code C_n based on the defined stop condition.

13

Improving runtime performance

Video

Source: <https://osu-nlp-group.github.io/SynCode/>

14

Positive points

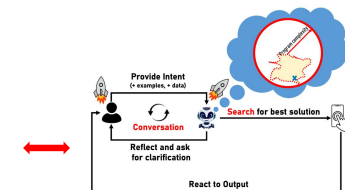
- Redefine software engineering in the era of LLM
- Reveal the challenges we need to address to shift to SE 3.0.

15

Negative points

- Mostly from a theoretical perspective...
- Blurry boundary between SE 2.0 and SE 3.0?

the human developer. In our experience, a typical programming session looks as follows: create a class, write the constructor's signature and have the copilot autocomplete the implementation, create a new method, write a code comment inside the body of that method (e.g., "## Clone a GitHub repository") to induce copilot to generate code, rewrite the comment to force copilot to use some specific package, run tests, discover that tests failed because a key requirement was missed, reason about three alternative copilot suggestions while coding the fix, finish writing the fix, rerun the tests, and so on. As the example highlights, the process is inherently task-driven and puts a significant cognitive overload on the human.



16

Rating

Theoretical foundation: 5/5

Practical experience: 4/5

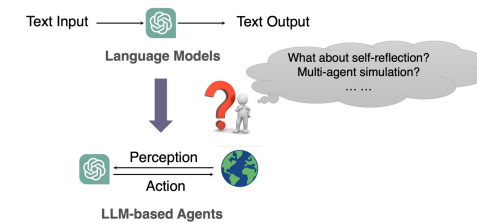
Overall: 4.5/5

17

Discussion

➤ Is SE 3.0 = SE based on agent?

'Modern' agent = LLM + external environment?

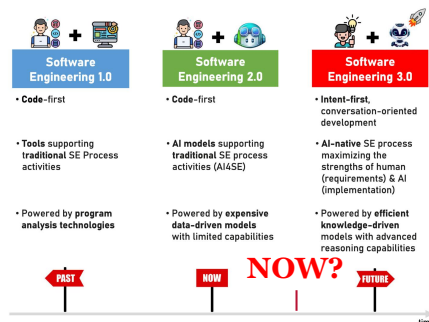


Source: https://youngho.github.io/resources/language_agents_Youngho_2024.pdf

18

Discussion

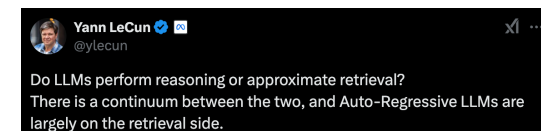
➤ We are now at SE 2.0 or 2.5?



19

Discussion

➤ Do FMs really think or just content retrieval?



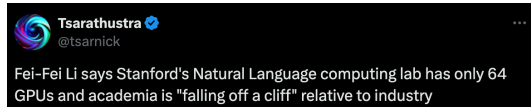
➤ If just content retrieval, then can we ever shift to SE 3.0?

Source: <https://x.com/>

20

Discussion

- Is Prompt Engineering the only thing we can do as a SE researcher?



Source: <https://x.com/>

21

UNIVERSITY OF
WATERLOO



22