

AI-Augmented Testing solution of LLM-Integrated Mobile Applications: Architecture, Test Oracle Strategies, and Empirical Evaluation

Urvish Gajjar

Sr Test Manager, Health care service corporation Dallas, TX, USA

Uv15gajjar@gmail.com

Abstract

—The integration of large language models (LLMs) into mobile applications introduces adaptive tutoring, conversational question answering, and automated feedback generation, but it also breaks the deterministic input-output assumptions on which conventional mobile test automation relies. This paper proposes an AI-augmented testing framework for LLM-integrated e-learning mobile applications that combines automated test-case generation

Applying only conventional scripted assertions to LLM-driven features leads to brittle tests that either over-reject semantically correct but differently worded answers, or under-detect genuinely incorrect, unsafe, or hallucinated content.

This paper makes three contributions. First, we present a layered reference architecture for LLM-integrated e-learning mobile applications that explicitly positions an AI test harness as a peer component to the prompt orchestrator and retrieval-augmented generation (RAG) knowledge base, rather than as an external, bolted-on test script. Second, we describe an AI-augmented testing methodology that pairs LLM-assisted test-case generation with a hybrid oracle combining embedding-based semantic similarity and LLM-as-a-judge rubric scoring, executed against the mobile client through Appium. Third, we report an empirical comparison of the proposed approach against manual and conventionally scripted testing on authoring effort, regression-cycle duration, defect-escape rate, and oracle false-negative rate, using a reference e-learning application as the evaluation subject.

The remainder of this paper is organized as follows. Section II reviews related work on LLM-based software testing and mobile test automation. Section III presents the system architecture. Section IV describes the AI-augmented testing methodology, including test-case generation and oracle design. Section V presents an implementation with representative code listings. Section VI reports the experimental evaluation. Section VII discusses implications and threats to validity, and Section VIII concludes with directions for future work.

II. Related Work

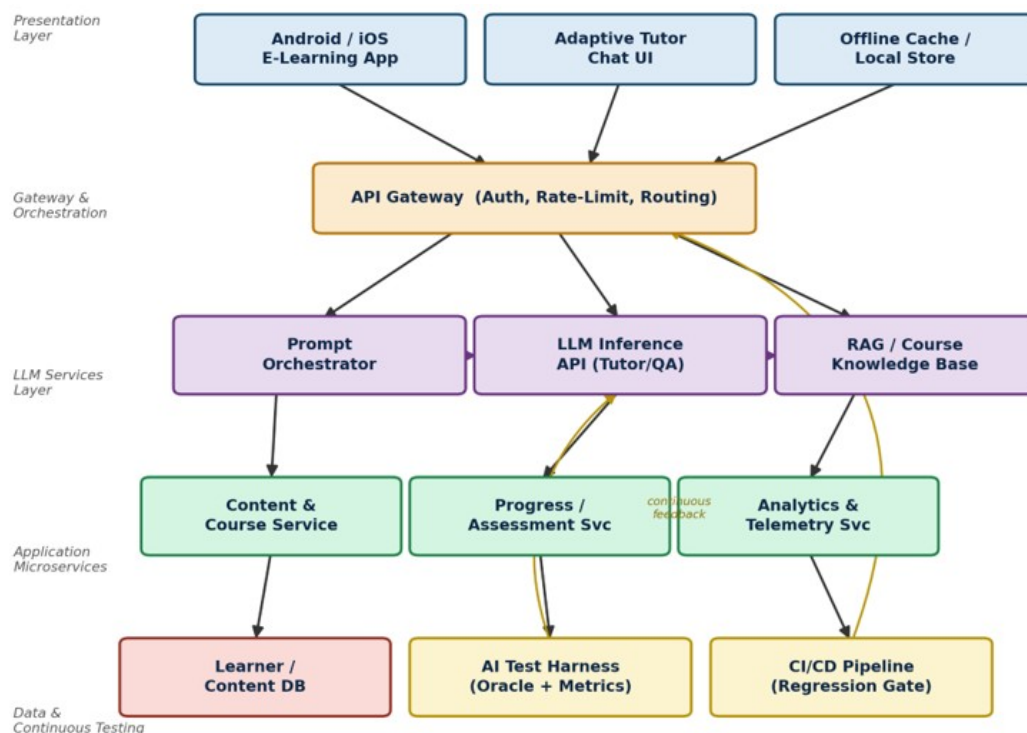
Software testing with LLMs has become an active research area. Wang et

distinct from, the oracle-focused problem this paper addresses: validating the content of AI-generated responses rather than only the reachability of UI states.

Standardized testing terminology in this paper follows ISO/IEC/IEEE 29119-1 [12]. Distinct from the cited work, this paper focuses specifically on the intersection of mobile UI test execution and LLM output verification within the e-learning domain, proposing a hybrid oracle and reporting empirical results from an applied deployment rather than a benchmark study alone.

III. System Architecture

Fig. 1 presents the reference architecture used in this study, organized into five layers: presentation, gateway and orchestration, LLM services, application microservices, and data with continuous testing.



C. LLM Services Layer

A prompt orchestrator assembles context, retrieves relevant course material from a retrieval-augmented generation (RAG) knowledge base, and issues calls to the LLM inference API. Grounding responses in retrieved course content reduces, but does not eliminate, hallucination, which motivates the safety and factuality checks in the oracle described in Section IV-B.

D. Application Microservices

Conventional services handle course content management, learner progress and assessment, and analytics/telemetry. These services remain deterministic and are validated using conventional test techniques, which this paper

the regression suite, mitigating the risk, noted by Jalil et al. [3], of LLM-authored tests asserting incidental rather than intended behavior.

B. Test Oracle Design for LLM Outputs

Because exact-match assertions are unsuitable for open-ended natural-language output, the harness applies a two-stage hybrid oracle. The first stage computes cosine similarity between sentence-transformer embeddings of the actual tutor response and a curated reference (golden) answer; responses below a calibrated similarity threshold are flagged as semantic drift. The second stage submits the response, reference answer, and a rubric (for example, scientific accuracy, age-appropriateness, absence of hallucinated facts) to a separate judge LLM, following the LLM-as-a-judge pattern of Zheng et al. [7]. Using a distinct model as judge, rather than the model under test, reduces self-preference bias. A response must pass both the similarity threshold and the rubric verdict, and must additionally score below a hallucination threshold, to be marked as passing.

C. Mobile UI Test Execution

Functional and integration assertions, including the rendering of the AI response inside the chat UI, offline fallback banners, and navigation flows, are executed against real devices and emulators through Appium, consistent with prior cross-platform mobile automation practice [8], [9]. This keeps oracle validation of LLM content and conventional UI verification within a single CI-gated pipeline rather than two disconnected toolchains.

V. Implementation

</

```

1  """
2  test_tutor_oracle.py
3  AI-augmented test oracle for validating LLM tutor responses
4  in the e-learning mobile application.
5  """
6  import pytest
7  from sentence_transformers import SentenceTransformer, util
8  from llm_client import LLMClient
9
10 embedder = SentenceTransformer("all-MiniLM-L6-v2")
11 judge = LLMClient(model="tutor-eval-judge")
12
13 SIMILARITY_THRESHOLD = 0.78
14
15 test_cases = [
16     {
17         "prompt": "Explain photosynthesis to a 10th grade student.",
18         "reference": "Photosynthesis converts light energy into chemical "
19         "energy, producing glucose and oxygen from CO2 and water.",
20         "rubric": ["scientifically_accur
```

```

1  """
2  test_lesson_flow_ui.py
3  Cross-platform Appium regression test for the e-learning app's
4  AI tutor chat screen (Android + iOS).
5  """
6  from appium import webdriver
7  from appium.webdriver.common.appiumby import AppiumBy
8  from selenium.webdriver.support.ui import WebDriverWait
9  from selenium.webdriver.support import expected_conditions as EC
10
11 CAPS = {
12     "platformName": "Android",
13     "automationName": "UiAutomator2",
14     "deviceName": "Pixel_7_API_34",
15     "app": "/builds/elearning-app-debug.apk",
16     "autoGrantPermissions": True,
17 }
18
19
20 class TestTutorChatScreen:
21     def setup_method(self):
22         self.driver = webdriver.Remote(

```

VI. Experimental Evaluation

A. Setup

The proposed framework was evaluated against a reference e-learning mobile application covering secondary-school science and mathematics content, deployed on Android (UiAutomator2) and iOS (XCTest) test devices. The evaluation compared two conditions over five regression cycles: (1) a manual/scripted baseline in which test cases were hand-authored and validated with exact or substring assertions, and (2) the proposed AI-augmented approach described in Sections IV and V. A held-out set of 240 learner questions, spanning in-scope curriculum topics

noted by Zheng et al. [7]. External validity is constrained by evaluating a single reference application; generalization to production-scale, multi-language e-learning platforms requires further study.

VII. Discussion

A. Implications for Testing Practice

The results suggest that treating LLM output verification as a first-class testing concern, rather than deferring to manual spot-checking, yields measurable efficiency and quality benefits. However, the hybrid oracle does not eliminate the underlying test oracle problem; it reframes it as a calibration exercise between an embedding-based similarity metric and a second model's judgment, both of which are themselves imperfect and subject to drift as the

VIII. Conclusion and Future Work

This paper presented a layered architecture and AI-augmented testing methodology for LLM-integrated e-learning mobile applications, combining LLM-assisted test-case generation, a hybrid semantic-similarity and LLM-as-judge test oracle, and Appium-based mobile UI execution within a single CI/CD-gated pipeline. An empirical evaluation against manual/scripted testing showed reductions in authoring effort and regression-cycle duration, along with improvements in defect-escape and oracle false-negative rates. Future work includes extending the oracle to multi-turn conversational context, evaluating oracle calibration across multiple LLM providers and languages, and investigating adaptive thresholding techniques that adjust similarity and rubric sensitivity based on