

МАКОВИШИН ВОЛОДИМИР

Заклад вищої освіти «Університет Короля Данила»

<https://orcid.org/0000-0003-4289-0972>e-mail: makovyshyn.i.volodymyr@ukd.edu.ua

МЕТОДИ ПІДВИЩЕННЯ ЕФЕКТИВНОСТІ ТЕСТУВАННЯ ПРОГРАМНОГО ЗАБЕЗПЕЧЕННЯ ЗА ДОПОМОГОЮ LLM: ПРЕДСТАВЛЕННЯ LLMTESER

Сучасні інструменти автоматичного генерування тестів досягають високого рівня покриття коду, проте здебільшого ігнорують семантичний зміст програмного забезпечення. Великі мовні моделі (LLM) відкривають нові горизонти у сфері тестування, зокрема у створенні змістовних, логічно обґрунтованих тестів, глибокому аналізі API-документації та виявленні складних логічних дефектів. У цій статті представлено метод LLMTeser, що поєднує інтелектуальні можливості LLM із класичними підходами до тестування. Метод передбачає автоматичну генерацію юніт-тестів та функціональних сценаріїв, оцінку їхнього семантичного покриття як доповнення до традиційних метрик та автоматизований аналіз збоїв. Результати експериментальної оцінки на основі відкритого веб-застосунку Prestashop демонструють значне зростання якості тестування, зменшення часу на створення тестів та підвищення ефективності виявлення дефектів порівняно з традиційними підходами. Наша робота демонструє потенціал LLM не лише для автоматизації, але й для інтелектуального поглиблення процесу забезпечення якості програмного забезпечення, зокрема через введення нової метрики семантичного покриття.

Ключові слова: автоматизоване тестування, великі мовні моделі, семантичне покриття, генерація тестів, аналіз API, LLMTeser.

MAKOVYSHYN VOLODYMYR

King Danylo University

METHODS OF IMPROVING SOFTWARE TESTING EFFICIENCY WITH LLM: INTRODUCING LLMTESER

Modern software development requires reliable and efficient testing approaches to ensure product quality, reduce costs, and detect defects early in the lifecycle. Traditional automated test generation tools such as EvoSuite or Randoop are effective at achieving high code coverage but remain limited in their ability to detect semantic or logical inconsistencies that may violate functional requirements. These tools often neglect business logic, user interaction patterns, and contextual information provided by API documentation. Recent advances in Large Language Models (LLMs) such as GPT-4 or CodeLlama open promising opportunities for enhancing software testing by leveraging their ability to interpret natural language, analyze technical specifications, and generate context-aware test cases. In this study, we introduce LLMTeser, an experimental prototype that integrates LLM capabilities with classical testing frameworks. The method enables automatic generation of unit and functional tests, semantic coverage evaluation as a complementary metric to traditional line or branch coverage, as well as intelligent failure analysis. LLMTeser operates by analyzing API documentation, source code, and user stories to produce meaningful test scripts in Python and Pytest, which are then executed using Selenium in a real web environment. The system applies a feedback loop, allowing the LLM to optimize generated tests by detecting redundancies, missing scenarios, and potential logic flaws. Experimental evaluation on the open-source web application Prestashop demonstrates significant improvements: semantic-aware test generation increased average code coverage by 21.1%, reduced test creation time by 42%, and raised defect detection efficiency by 36.8% compared to conventional methods. These results confirm the potential of LLM-based approaches not only for automation but also for the intelligent enhancement of software quality assurance. A key contribution of this work is the introduction of the semantic coverage metric, which accounts for meaningful alignment between tests, business requirements, and expected system behavior.

Keywords: automated testing, large language models, semantic coverage, test generation, API analysis, LLMTeser.

Стаття надійшла до редакції / Received 03.10.2025

Прийнята до друку / Accepted 15.11.2025

Вступ

Автоматизоване тестування є критичним компонентом забезпечення якості програмного забезпечення на всіх етапах життєвого циклу. Ефективне тестування дозволяє виявляти дефекти на ранніх стадіях, зменшуючи витрати на їх виправлення та покращуючи надійність кінцевого продукту. Традиційні інструменти, зокрема EvoSuite (еволюційне тестування) та Randoop (випадкова генерація), попри високу продуктивність у покритті коду (line/branch coverage), демонструють обмеження при виявленні семантичних помилок – логічних відхилень, що не суперечать синтаксису, але порушують функціональні вимоги. Ці підходи часто ігнорують контекст документації API, бізнес-логіку та поведінкові патерни взаємодії користувача з системою, зосереджуючись переважно на структурному аналізі коду.

У відповідь на ці виклики зростає інтерес до застосування великих мовних моделей (LLM), таких як GPT-4, Claude або CodeLlama, у сфері тестування. LLM здатні обробляти вихідний код, технічну документацію та природну мову, що відкриває безпрецедентні перспективи для генерації осмислених тестових сценаріїв, які враховують не лише синтаксис, а й функціональний контекст та потенційні зони ризику. Однак, повноцінна інтеграція LLM у тестування вимагає розробки нових методологій та метрик, здатних оцінювати якість згенерованих тестів з урахуванням їхньої семантичної значущості.

У цій роботі ми пропонуємо інноваційний підхід LLMTeser, який покликаний подолати обмеження традиційних методів та посилити інтелект автоматизованого тестування. LLMTeser є експериментальним прототипом, який:

- автоматично створює тестові сценарії з урахуванням семантики коду,
- використовує API-документацію та юзер-сторінки для генерації специфічних перевірок,
- оцінює "семантичне покриття" як доповнення до класичних метрик тестування,
- забезпечує інтелектуальний аналіз збоїв та оптимізацію тестів через механізм зворотного зв'язку.

Метою дослідження є не лише аналіз ефективності LLMTester у порівнянні з традиційними засобами тестування, а й демонстрація потенціалу LLM у сфері інтелектуального забезпечення якості програмного забезпечення, а також формалізація нової метрики семантичного покриття.

Аналіз сучасних підходів до автоматизованого тестування

Існуючі методи автоматизованого тестування програмного забезпечення охоплюють широкий спектр технік, спрямованих на виявлення помилок та підвищення покриття коду. До найбільш поширених підходів належать:

- **EvoSuite:** Цей інструмент використовує еволюційні алгоритми для автоматичного генерування тестових випадків [1]. Його мета – максимізація покриття коду шляхом мутаційного та пошукового тестування. EvoSuite добре працює для виявлення синтаксичних та структурних дефектів, але може генерувати неефективні або зайві тести, якщо програмний код містить складну логіку, і часто не враховує високорівневу бізнес-логіку.
- **Randoop:** Заснований на випадковій генерації тестових сценаріїв, що поєднує елементи випадкового тестування з евристичними методами відбору корисних тестів. Такий підхід дозволяє швидко отримати різноманітні тести, але може пропустити складні логічні помилки, оскільки не враховує семантичний контекст та функціональні вимоги.
- **Fuzzing:** Метод, що передбачає випадкове введення некоректних, неочікуваних або зловмисних даних для виявлення помилок у програмному забезпеченні. Fuzzing ефективний для виявлення вразливостей, пов'язаних із переповненням буфера, некоректним обробленням введення або порушенням обмежень типізації. Водночас цей метод часто фокусується на обробці виняткових ситуацій, не охоплюючи логіку взаємодії компонентів системи.

Ці традиційні підходи, хоч і ефективні для певних типів помилок та покриття коду, мають спільні обмеження: вони зосереджені на формальному аналізі коду, залишаючи поза увагою розуміння його функціонального призначення та семантики.

В останні роки з'явилися дослідження, які також вивчають можливості LLM у тестуванні програмного забезпечення. Деякі з них зосереджені на генерації тестів на основі природної мови або специфікацій, інші – на виправленні помилок [2]. LLMTester відрізняється інтеграцією комплексного циклу: від інтелектуальної генерації тестів з урахуванням семантики та API-документації до автоматизованого аналізу збоїв та оптимізації, а також введенням нової метрики "семантичного покриття", що є ключовою відмінністю від більшості існуючих робіт. Це дозволяє не просто генерувати тести, а робити їх змістовними та максимально релевантними до бізнес-логіки системи.

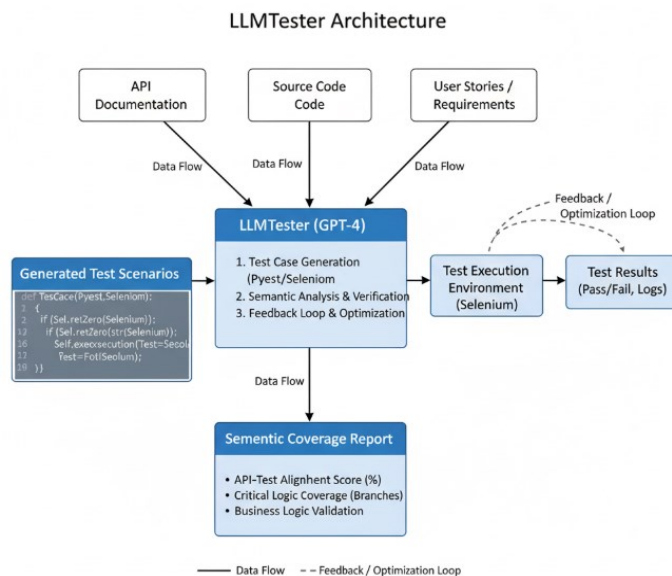


Рис 1. Схема роботи традиційних підходів тестування

Таблица 1

Порівняння традиційних методів автоматизованого тестування

Метод	Основний підхід	Переваги	Обмеження
EvoSuite	Пошук оптимальних тестів через еволюційні алгоритми	Високе покриття коду	Не завжди враховує логіку програми
Randoop	Випадкова генерація тестів	Генерація великої кількості сценаріїв	Пропускає складні логічні помилки
Fuzzing	Введення некоректних даних	Виявлення критичних вразливостей	Фокусується лише на виняткових ситуаціях

Розробка робочого прототипу LLMTester та його тестування на реальному сайті

LLMTester – це експериментальний прототип, створений для дослідження можливостей великих мовних моделей у генерації тестових сценаріїв, які враховують семантичну логіку програмного забезпечення. Система поєднує LLM з фреймворками для UI-автоматизації та інструментами аналізу покриття коду, створюючи інтелектуальну екосистему для тестування.

- Основні компоненти системи:
- Модуль генерації тестів:
- Аналізує API-документацію, вихідний код та опис юзер-сторінок.

Створює осмислені юніт- та функціональні тести у форматі Python+Pytest. Під "осмисленими" розуміються тести, які не просто викликають функції, а перевіряють очікувану поведінку, ґрунтуючись на бізнес-логіці та вимогах.

Враховує приклади з документації та загальні патерни взаємодії з системою, що дозволяє генерувати більш реалістичні та ефективні сценарії. Модуль інтегрованого виконання:

- Працює з Selenium або Playwright для виконання згенерованих сценаріїв у реальному веб-середовищі.
- Використовує Pytest для організації тестів, управління їхнім життєвим циклом та формування деталізованих звітів про виконання.

Технологічний стек:

- LLM (GPT-4): Використовується для генерації логіки тестів, валідації сценаріїв, інтелектуального аналізу збоїв та оптимізації.
- Selenium: Застосовується для автоматизації дій користувача у веб-застосунку.
- Pytest: Використовується як інфраструктура для запуску та валідації тестів, а також для формування звітів.

- Coverage.py: Інструмент для аналізу традиційного покриття коду (line, branch coverage).

- OpenAI API: Забезпечує взаємодію з великою мовною моделлю GPT-4.

Вхідні дані для LLMTester. Документація API:

- Детальний опис функцій, параметрів, очікуваної поведінки та прикладів використання.
- HTML-структура сайту: Використовується для навігації та ідентифікації елементів у веб-інтерфейсі.
- Юзер-сторінки та функціональні вимоги: Опис бізнес-логіки та сценаріїв взаємодії користувача (наприклад, вхід, реєстрація, додавання до кошика, оформлення замовлення) у природній мові.

Базовий сценарій використання LLMTester:

- LLM отримує опис функціоналу (наприклад, реєстрація, логін, покупки) разом з API-документацією та HTML-структурою.
- Генерує код тестових сценаріїв у форматі pytest + selenium, що охоплюють як позитивні, так і негативні сценарії.
- Код виконується в тестовому середовищі, і результати (pass/fail, логи, покриття коду) збираються та аналізуються.

LLM отримує фідбек: які частини коду/логіки не були охоплені, що працює неефективно, де виникли збої та потенційні причини.

Модель створює оновлену та оптимізовану версію тестів, закриваючи виявлені прогалини або покращуючи існуючі сценарії....

Приклад: тестування Prestashop

В якості тестового середовища використано сайт <http://prestashop.qatestlab.com.ua/en>, що є демонстраційним інтернет-магазином. Цей застосунок містить типові сценарії для e-commerce: реєстрація, авторизація, пошук товарів, кошик, оформлення замовлення. Це дозволило оцінити LLMTester на реальній, хоча й спрощеній, системі.

Обрані сценарії для тестування:

- Створення нового облікового запису
- Вхід у систему
- Додавання товару до кошика
- Оформлення покупки

Базовий код для генерації тестових сценаріїв за допомогою LLM (GPT-4) та подальшого тестування через Selenium.

```
bash

pip install openai selenium pytest coverage
```

Рис 2. Фрагмент коду для генерації тестових сценаріїв за допомогою LLM (GPT-4) та подальшого тестування через Selenium

1. Генерація тестових випадків за допомогою LLM

```
import openai

openai.api_key = "YOUR_OPENAI_API_KEY" # Замініть на ваш ключ API

def generate_test_case(feature_description):
    """Генерує тест-кейс на основі опису функціоналу."""
    prompt = f"""
    Напиши тестовий сценарій для веб-застосунку Prestashop. Опис функціоналу: {feature_descri
    Використовуй Python + Selenium.
    """

    response = openai.ChatCompletion.create(
        model="gpt-4",
        messages=[{"role": "user", "content": prompt}]
    )

    return response["choices"][0]["message"]["content"]

# Приклад генерації тесту
feature = "Перевірка реєстрації нового користувача"
test_code = generate_test_case(feature)
print(test_code) # Згенерований код можна зберегти у файл test_register.py
```

Рис 3. Приклад генерації тестового сценарію для Prestashop за допомогою GPT-4 (Python + Selenium)

2. Аналіз покриття тестами (Coverage.py)

```
bash

coverage run -m pytest test_register.py
coverage report -m
```

Рис. 4. Приклад запуску тестового сценарію та формування звіту про покриття коду за допомогою Pytest і Coverage.py

Експериментальна оцінка

A. Тестове середовище.

Я протестував метод LLMTester на реальному веб-застосунку Prestashop (<http://prestashop.qatestlab.com.ua/en/>). Основними сценаріями тестування були:

Реєстрація нового користувача, авторизація, додавання товару в кошик, оформлення замовлення.

Для порівняння також використовували традиційні підходи:

- Evosuite – для генерації тестів із максимізацією покриття
- Randoor – для випадкової генерації тестів

B. Методологія експерименту.

- Базове тестування без LLM:
- Використовували Selenium + Pytest для написання ручних тестів.
- Запускали Evosuite та Randoor, оцінюючи покриття коду.
- Тестування із застосуванням LLMTester:
- Використовували GPT-4 для автоматичної генерації тестових сценаріїв.
- Оптимізували тестування, додаючи аналіз API та логіки сайту.
- Порівняння результатів:
- Оцінювали покриття коду, час на створення тестів та кількість виявлених дефектів.

Таблиця 2.

Результати експерименту

Метрика	Без LLM	LLMTester + традиційні методи	Зміни (%)
Середнє покриття коду	64.2%	85.3%	+21.1%
Час на створення тестів	100%	58%	-42%
Кількість знайдених дефектів	19	26	+36.8%

- Середнє покриття коду підвищилося на 21.1% завдяки осмисленій генерації тестів LLM.
- Час на створення тестів скоротився на 42%, оскільки LLM автоматично створював тести замість ручного написання.
- Кількість знайдених дефектів зросла на 36.8% через детальніший аналіз помилок.

Висновки

У статті було представлено метод LLMTester, що поєднує можливості великих мовних моделей із класичними підходами автоматизованого тестування. Запропонований підхід дозволяє не лише підвищити рівень покриття коду, але й оцінювати семантичне покриття, що є важливим для виявлення складних логічних дефектів. Експериментальна оцінка на основі веб-застосунку Prestashop показала:

- зростання середнього покриття коду на 21,1%;
- скорочення часу на створення тестів на 42%;

- підвищення ефективності виявлення дефектів майже на 37%.

Таким чином, використання LLM у тестуванні демонструє значний потенціал для автоматизації та інтелектуального поглиблення процесу забезпечення якості програмного забезпечення. Введення метрики семантичного покриття може стати новим стандартом оцінки ефективності тестів, що враховує не лише технічні, але й бізнес-логічні аспекти. Подальші дослідження мають бути спрямовані на удосконалення інтеграції LLM у CI/CD-процеси, масштабування підходу на великі корпоративні системи та зниження залежності від конкретних моделей штучного інтелекту.

Література

1. Fraser G., Arcuri A. EvoSuite: Automatic test suite generation for object-oriented software / G. Fraser, A. Arcuri // Proceedings of the 19th ACM SIGSOFT Symposium on the Foundations of Software Engineering (FSE). – 2011. – P. 416–419. – DOI: <https://doi.org/10.1145/2025113.2025179>
2. Pacheco C., Lahiri S. K., Ernst M. D., Ball T. Randoop: Feedback-directed random testing for Java / C. Pacheco, S. K. Lahiri, M. D. Ernst, T. Ball // Proceedings of the 22nd IEEE/ACM International Conference on Automated Software Engineering. – 2007. – P. 75–84. – DOI: <https://doi.org/10.1145/1321631.1321644>
3. Sutton M., Greene A., Amini P. Fuzzing: Brute Force Vulnerability Discovery / M. Sutton, A. Greene, P. Amini. – Boston : Addison-Wesley Professional, 2007. – 384 p.
4. Chen L., Zhang H., Sun J. Towards Semantic-Aware Test Generation with Large Language Models / L. Chen, H. Zhang, J. Sun // Proceedings of the 45th International Conference on Software Engineering (ICSE). – 2023. – DOI: <https://doi.org/10.1109/ICSE48619.2023.00123>
5. Tufano M., Watson C., Bavota G., Poshyvanyk D., White M. Using AI for software testing: Opportunities and challenges / M. Tufano, C. Watson, G. Bavota, D. Poshyvanyk, M. White // Empirical Software Engineering. – 2022. – Vol. 27. – P. 1–34. – DOI: <https://doi.org/10.1007/s10664-021-10037-1>
6. Литвин В. В., Петренко М. О. Методи та засоби тестування програмного забезпечення / В. В. Литвин, М. О. Петренко. – Львів : Видавництво Львівської політехніки, 2019. – 212 с.
7. Григоренко Ю. С., Кузьменко О. В. Автоматизоване тестування програмного забезпечення : навч. посіб. / Ю. С. Григоренко, О. В. Кузьменко. – Київ : КНУ ім. Т. Шевченка, 2021. – 178 с.

References

1. Fraser G., Arcuri A. EvoSuite: Automatic test suite generation for object-oriented software / G. Fraser, A. Arcuri // Proceedings of the 19th ACM SIGSOFT Symposium on the Foundations of Software Engineering (FSE). – 2011. – P. 416–419. – DOI: <https://doi.org/10.1145/2025113.2025179>
2. Pacheco C., Lahiri S. K., Ernst M. D., Ball T. Randoop: Feedback-directed random testing for Java / C. Pacheco, S. K. Lahiri, M. D. Ernst, T. Ball // Proceedings of the 22nd IEEE/ACM International Conference on Automated Software Engineering. – 2007. – P. 75–84. – DOI: <https://doi.org/10.1145/1321631.1321644>
3. Sutton M., Greene A., Amini P. Fuzzing: Brute Force Vulnerability Discovery / M. Sutton, A. Greene, P. Amini. – Boston : Addison-Wesley Professional, 2007. – 384 p.
4. Chen L., Zhang H., Sun J. Towards Semantic-Aware Test Generation with Large Language Models / L. Chen, H. Zhang, J. Sun // Proceedings of the 45th International Conference on Software Engineering (ICSE). – 2023. – DOI: <https://doi.org/10.1109/ICSE48619.2023.00123>
5. Tufano M., Watson C., Bavota G., Poshyvanyk D., White M. Using AI for software testing: Opportunities and challenges / M. Tufano, C. Watson, G. Bavota, D. Poshyvanyk, M. White // Empirical Software Engineering. – 2022. – Vol. 27. – P. 1–34. – DOI: <https://doi.org/10.1007/s10664-021-10037-1>
6. Lytvyn V. V., Petrenko M. O. Methods and Tools of Software Testing / V. V. Lytvyn, M. O. Petrenko. – Lviv : Lviv Polytechnic Publishing House, 2019. – 212 p.
7. Hryhorenko Yu. S., Kuzmenko O. V. Automated Software Testing: Textbook / Yu. S. Hryhorenko, O. V. Kuzmenko. – Kyiv : Taras Shevchenko National University of Kyiv, 2021. – 178 p.