



AI dla testera manualnego

Jak tester manualny odnajduje się w świecie AI bez utraty jakości i bezpieczeństwa

PO CO TESTEROWI AI

AI odciąża testera manualnego w rutynowych zadaniach: generuje warianty, streszcza logi i podpowiada scenariusze. Dzięki temu masz więcej czasu na **myślenie eksploracyjne**, ocenę ryzyka i rozmowę z zespołem.

Uwaga na proporcje: AI oszczędza czas przy powtarzalnych zadaniach, ale w niejednoznacznych przypadkach weryfikacja może trwać dłużej niż praca ręczna.

Zasada nadrzędna

Weryfikuj każdą odpowiedź przed użyciem. Model podaje wynik prawdopodobny — Ty oceniasz zgodność z wymaganiami.

POJĘCIA W PIGUŁCE

- **LLM (model językowy):** Program przewidujący kolejne słowa. Dobrze tworzy tekst, ale nie zna Twojego systemu.
- **Prompt:** Polecenie i kontekst, który podajesz modelowi. Jakość odpowiedzi = jakość promptu.
- **Kontekst:** Wszystko, co model „widzi” w rozmowie. Bez kontekstu uzupełnia luki założeniami.
- **Okno kontekstu:** Ilość tekstu, którą model obejmuje naraz. Przy zbyt małym oknie może zgubić początek długich wymagań lub logów.
- **Token:** Kawałek tekstu (~4 znaki w angielskim, w polskim mniej). Limit tokenów mówi, ile model „zmieści” w promptcie i odpowiedzi.
- **Model rozumujący:** Wolniejszy i droższy, ale lepszy w analizie wymagań, przypadkach brzegowych i szukaniu przyczyn błędów.
- **Model szybki (lekki):** Tani i szybki, słabszy w złożonej logice. Dobry do streszczeń, porządkowania zgłoszeń, redakcji i prostych wariantów.

- **Halucynacja:** Odpowiedź brzmiąca wiarygodnie, ale nieprawdziwa.
- **RAG:** Dołącza do promptu fragmenty dokumentów, by model odpowiadał na ich podstawie, nie z pamięci.
- **Agent AI:** Model, który sam wykonuje kroki (klika, wysyła zapytania). Wymaga Twojej autoryzacji.

GDZIE AI POMAGA W QA MANUALNYM

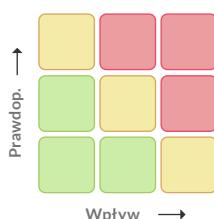
- Przypadki testowe z kryteriów akceptacji
- Pomysły na testy eksploracyjne, brzegowe i negatywne
- Dane testowe, przykładowe payloady JSON, regexy
- Streszczanie długich logów i raportów
- Przegląd wymagań pod kątem testowalności
- Porządkowanie kroków reprodukcji błędu
- Tłumaczenie i redakcja zgłoszeń

Ocena kompletności testów

AI wygeneruje wiele przypadków, ale Ty oceniasz istotność ryzyk. Zestaw może wyglądać kompletnie, a pomijać ścieżki krytyczne.

TESTUJ W OPARCIU O RYZYKO

Nie przetestujesz wszystkiego, więc ustaw priorytety. Liczy się prawdopodobieństwo awarii i jej wpływ na użytkownika oraz biznes. AI zbierze ryzyka, ale ich wagę oceniasz Ty.



Czerwone testujesz najpierw, zielone na końcu.

SYGNAŁY WYSOKIEGO RYZYKA

Obszar jest bardziej ryzykowny, gdy ma:

- Skomplikowaną logikę i wiele warunków
- Świeżo lub często zmieniany kod
- Krytyczne ścieżki i częste użycie
- Historię defektów
- Integracje i zależności zewnętrzne
- Wpływ na dane, pieniądze, bezpieczeństwo lub zgodność

Analiza i priorytetyzacja ryzyk

Wypisz ryzyka funkcji: co może pójść nie tak, kogo dotknie i jak poważne będą skutki. Uporządkuj od najważniejszych. Dodaj założenia.

DOBRY PROMPT W 5 KROKACH

1. **Rola:** „Jesteś testerem. Zaproponuj testy...”
2. **Kontekst:** wklej kryteria akceptacji, fragment wymagań lub opis funkcji
3. **Zadanie:** napisz, czego oczekujesz, np. „Wypisz ryzyka”, „Przygotuj przypadki testowe”
4. **Ograniczenia:** podaj zakres, format, granice (np. „tylko negatywne”)
5. **Format wyjścia:** wskaż rodzaj oczekiwanej odpowiedzi np. tabela lub lista; poproś też o założenia

ANATOMIA PROMPTU



Kontekst promptu

Im więcej podasz, tym mniej model zgaduje.

Kontekst sesji

Zbyt mocne okrojenie globalnego kontekstu zwiększa ryzyko, że model pominie ważne informacje z wcześniejszej części rozmowy.

GOTOWE PROMPTY

Przypadki testowe z kryteriów akceptacji

Wygeneruj przypadki pozytywne, negatywne i brzegowe dla tych kryteriów. Zwróć tabelę: ID | warunek wstępny | kroki | oczekiwany wynik. Dodaj założenia.

Pomysły na testy eksploracyjne

Zaproponuj 10 testów eksploracyjnych. Uwzględnij dane brzegowe, przerwanie procesu, uprawnienia i i18n.

Bug report z moich notatek

Uporządkuj moje notatki w zgłoszenie: tytuł, kroki, wynik oczekiwany vs. faktyczny, środowisko. Notatki: [WKLEJ].

ZASADY BEZPIECZNEGO KORZYSTANIA Z AI

Korzystanie z zatwierdzonych narzędzi: używaj tylko firmowo dopuszczonych narzędzi AI. Inne mogą wyносить dane poza organizację.

Anonimizacja danych: Przed wysłaniem promptu zamień prawdziwe dane na znaczniki, np. [KLIENT], [NR_UMOWY].

Oddzielenie instrukcji od danych: Jasno oznacz, co jest poleceniem, a co materiałem do analizy. Zmniejsza to ryzyko prompt injection.

Human-in-the-loop: Autoryzuj działania AI. Zatwierdzaj każdy istotny krok wykonywany przez agenta AI.

Procedury awaryjne: Wiedz, gdzie są aktualne instrukcje i kogo powiadomić, gdy coś pójdzie nie tak.

LISTA KONTROLNA PRZED UŻYCIEM AI

Korzystanie z zatwierdzonych narzędzi: używaj tylko firmowo dopuszczonych narzędzi AI. Inne mogą wyносить dane poza organizację.

- › Czy wybrane narzędzie AI jest na **liście dopuszczonych**
- › Czy usunąłem z promptu prawdziwe dane i wstawiłem **placeholder** np. [KLIENT]
- › Czy analizując dokument, **oddzieliłem instrukcję** od wklejanego tekstu
- › Czy traktuję odpowiedź AI jako **roboczy szkic** wymagający mojej akceptacji
- › Czy działania agenta AI podlegają moim regułom autoryzacji **Human-in-the-loop**
- › Czy wiem, **gdzie znaleźć aktualne procedury** na wypadek pomyłki

CO ROBIĆ, A CZEGO UNIKAĆ

RÓB	NIE RÓB
• Sprawdzaj kod, fakty i liczby • Anonimizuj dane • Oznaczaj założenia • Poprawiaj prompt iteracyjnie	• Nie wklejaj haseł, PII ani danych klienta • Nie zakładaj, że podane API istnieje • Nie publikuj wyniku bez weryfikacji • Nie uruchamiaj działań bez autoryzacji

Kontekst promptu

To, co wkleisz do zewnętrznego narzędzia, opuszcza firmę. Trzymaj się polityki bezpieczeństwa i listy dopuszczonych narzędzi.

JAK WYCHWYCIĆ HALUCYNACJĘ

- › Sprawdź u źródła nazwy pól, endpointy i wersje API
- › Policz i zweryfikuj liczby oraz identyfikatory
- › Poproś o źródła i założenia; brak źródeł traktuj jako sygnał ostrzegawczy, podajne źródła weryfikuj

NIM ZGŁOSISZ BŁĄD ZNALEZIONY Z AI

Upewnij się, że:

- Błąd został zreprodukowany ręcznie i nie opierasz zgłoszenia wyłącznie na odpowiedzi modelu
- Z opisu usunięto dane wrażliwe
- Zgłoszenie zawiera rzeczywiste kroki reprodukcji i oczekiwany wynik