

Prompt Engineering

Zostań Panem Sztucznej Inteligencji

Logan Anderson

Spis treści

1: Podstawy Prompt Engineering.....	5
Co to jest Prompt Engineering?	6
Historia i ewolucja sztucznej inteligencji	8
Jak działa ChatGPT?	10
2: Tworzenie Skutecznych Promptów.....	13
Podstawowe zasady tworzenia promptów	Błąd! Nie zdefiniowano zakładki.
Struktura i składnia promptów	Błąd! Nie zdefiniowano zakładki.
Przykłady skutecznych promptów.	Błąd! Nie zdefiniowano zakładki.
3: Tworzenie Skutecznych Promptów	Błąd! Nie zdefiniowano zakładki.
Podstawowe zasady tworzenia promptów	Błąd! Nie zdefiniowano zakładki.
Struktura i składnia promptów	Błąd! Nie zdefiniowano zakładki.
Przykłady skutecznych promptów.	Błąd! Nie zdefiniowano zakładki.
4: Zwiększanie Produktwności	Błąd! Nie zdefiniowano zakładki.
Automatyzacja zadań rutynowych.	Błąd! Nie zdefiniowano zakładki.
Optymalizacja procesów biznesowych	Błąd! Nie zdefiniowano zakładki.
Zarządzanie czasem z pomocą ChatGPT	Błąd! Nie zdefiniowano zakładki.
5: Wzmacnianie Kreatywności	Błąd! Nie zdefiniowano zakładki.
Generowanie pomysłów na nowe projekty	Błąd! Nie zdefiniowano zakładki.
Tworzenie treści kreatywnych	Błąd! Nie zdefiniowano zakładki.

Rozwiązywanie problemów kreatywnie **Błąd! Nie zdefiniowano zakładki.**

6: Integracja ChatGPT z Narzędziami i Technologiami **Błąd! Nie zdefiniowano zakładki.**

ChatGPT w aplikacjach i usługach .. **Błąd! Nie zdefiniowano zakładki.**

Wykorzystanie API ChatGPT **Błąd! Nie zdefiniowano zakładki.**

Przyszłość integracji AI **Błąd! Nie zdefiniowano zakładki.**

7: Zaawansowane Techniki Prompt Engineering **Błąd! Nie zdefiniowano zakładki.**

Customizacja promptów **Błąd! Nie zdefiniowano zakładki.**

Modelowanie odpowiedzi **Błąd! Nie zdefiniowano zakładki.**

Analiza i poprawa wyników **Błąd! Nie zdefiniowano zakładki.**

8: Rodzaje Promptów **Błąd! Nie zdefiniowano zakładki.**

Prompt "HYDE" (Hide Your Dissent) **Błąd! Nie zdefiniowano zakładki.**

Prompt "Step by Step" (Krok po Kroku) **Błąd! Nie zdefiniowano zakładki.**

Prompt "Chain of Thought" (Łańcuch Myśli) **Błąd! Nie zdefiniowano zakładki.**

Prompt "Zero-Shot" (Bezprzykładowe) **Błąd! Nie zdefiniowano zakładki.**

Prompt "Few-Shot" (Kilka Przykładów) **Błąd! Nie zdefiniowano zakładki.**

9: Przykładowe Prompty **Błąd! Nie zdefiniowano zakładki.**

Przykłady promptów do zwiększania produktywności **Błąd! Nie zdefiniowano zakładki.**

Przykłady promptów kreatywnych **Błąd! Nie zdefiniowano zakładki.**

Przykłady promptów do zastosowań biznesowych**Błąd!** **Nie**
zdefiniowano zakładki.

Przykłady promptów technicznych**Błąd!** **Nie zdefiniowano zakładki.**

Przykłady promptów edukacyjnych**Błąd!** **Nie** **zdefiniowano**
zakładki.

1: Podstawy Prompt Engineering

Co to jest Prompt Engineering?

Prompt engineering to sztuka i nauka tworzenia skutecznych instrukcji, zwanych promptami, dla modeli językowych, takich jak te oparte na sztucznej inteligencji (AI), by uzyskiwać od nich pożądane odpowiedzi. Wyobraź sobie, że masz do dyspozycji inteligentnego rozmówcę, który potrafi odpowiadać na niemal każde pytanie, tworzyć teksty, tłumaczyć języki, generować pomysły i wiele więcej. Aby jednak uzyskać wartościowe odpowiedzi, musisz umieć formułować swoje pytania i prośby w sposób, który ten model najlepiej zrozumie. Właśnie na tym polega prompt engineering.

Znaczenie prompt engineering jest trudne do przecenienia, zwłaszcza w erze dynamicznego rozwoju sztucznej inteligencji. Umiejętność precyzyjnego komunikowania się z modelami AI staje się coraz bardziej cenna, ponieważ pozwala to na maksymalne wykorzystanie ich potencjału. Od edukacji, przez medycynę, po rozrywkę – właściwe formułowanie promptów może znacznie zwiększyć efektywność i użyteczność technologii opartej na sztucznej inteligencji.

W dziedzinie edukacji, na przykład, nauczyciele mogą używać odpowiednio zaprojektowanych promptów do tworzenia spersonalizowanych materiałów edukacyjnych, automatycznego oceniania prac uczniów czy generowania pytań do testów. W medycynie, prompt engineering może pomóc lekarzom w szybszym przeszukiwaniu literatury medycznej, diagnozowaniu chorób czy planowaniu terapii. W przemyśle kreatywnym, jak pisarstwo czy produkcja filmowa, odpowiednio skonstruowane prompty mogą stać się źródłem inspiracji, pomagając twórcom w generowaniu nowych pomysłów i treści.

Podstawowe koncepcje związane z tworzeniem promptów obejmują kilka kluczowych aspektów. Po pierwsze, zrozumienie kontekstu – prompt powinien być osadzony w kontekście, który model AI może

zrozumieć i odpowiednio do niego się odnieść. Na przykład, pytając o definicję skomplikowanego terminu, warto podać dodatkowe informacje, które zawężą odpowiedź do pożądanego obszaru wiedzy.

Kolejną istotną koncepcją jest specyficzność. Im bardziej szczegółowy i precyzyjny jest prompt, tym większa szansa na uzyskanie satysfakcjonującej odpowiedzi. Zamiast pytać ogólnie o "zalety zdrowego stylu życia", lepiej zadać pytanie w rodzaju: "Jakie są trzy główne korzyści zdrowotne wynikające z codziennej 30-minutowej aktywności fizycznej?".

Kolejnym elementem jest jasność – prompt powinien być sformułowany w sposób jasny i jednoznaczny, unikając niepotrzebnych dwuznaczności, które mogą zmylić model. Proste, zrozumiałe zdania zwiększają prawdopodobieństwo uzyskania dokładnej odpowiedzi.

Interaktywność to następna ważna cecha. Często prompty są częścią większej interakcji między użytkownikiem a modelem AI. W takim przypadku warto projektować prompty, które zachęcają do kontynuowania rozmowy i pogłębiania tematu. Przykładowo, zamiast jednorazowego pytania, można sformułować prompt, który skłoni model do zadawania dodatkowych pytań w celu lepszego zrozumienia zagadnienia.

Ostatnią, ale równie ważną koncepcją jest testowanie i iteracja. Tworzenie efektywnych promptów to proces, który często wymaga prób i błędów. Kluczowe jest testowanie różnych wersji promptów, analizowanie wyników i wprowadzanie odpowiednich poprawek. Na przykład, jeśli uzyskana odpowiedź nie spełnia oczekiwań, warto zmienić formę promptu i spróbować ponownie.

Prompt engineering to fascynująca dziedzina, która łączy elementy językoznawstwa, technologii i kreatywnego myślenia. Umożliwia ona użytkownikom maksymalne wykorzystanie potencjału modeli językowych, oferując nieskończone możliwości w różnych dziedzinach życia. Poprzez zrozumienie i stosowanie podstawowych koncepcji

tworzenia promptów, można znacząco zwiększyć efektywność i użyteczność interakcji z AI, czyniąc ją bardziej produktywną i satysfakcjonującą.

Historia i ewolucja sztucznej inteligencji

Historia i ewolucja sztucznej inteligencji to fascynujący temat, pełen przełomowych odkryć i momentów, które zmieniły świat. Wszystko zaczęło się od marzeń i spekulacji o maszynach zdolnych myśleć i działać jak ludzie. Już w starożytności pojawiały się koncepcje mechanicznych istot w mitologiach różnych kultur, ale prawdziwa historia sztucznej inteligencji rozpoczęła się dopiero w XX wieku.

Jednym z kluczowych momentów w historii AI było opublikowanie pracy "Computing Machinery and Intelligence" przez Alana Turinga w 1950 roku. Wprowadził on pojęcie testu Turinga, który miał oceniać zdolność maszyny do wykazywania inteligentnego zachowania nierozróżnialnego od ludzkiego. To był pierwszy krok w kierunku formalizacji pojęcia sztucznej inteligencji.

W 1956 roku odbyła się konferencja w Dartmouth College, która jest powszechnie uznawana za moment narodzin AI jako odrębnej dziedziny nauki. John McCarthy, Marvin Minsky, Nathaniel Rochester i Claude Shannon zebrali się, aby przedyskutować możliwości tworzenia "myślących maszyn". To wydarzenie zapoczątkowało intensywne badania nad AI, prowadzone przez naukowców z całego świata.

W latach 60. i 70. rozwój AI nabierał tempa, głównie dzięki tworzeniu systemów opartych na regułach i algorytmach heurystycznych. Jednym z przełomowych projektów był program General Problem Solver (GPS) autorstwa Herberta Simona i Allena Newella, który miał na celu symulację ludzkiego myślenia przy rozwiązywaniu problemów. W tym okresie powstały również pierwsze systemy eksperckie, takie jak

DENDRAL i MYCIN, które były w stanie doradzać w specjalistycznych dziedzinach, takich jak chemia organiczna i medycyna.

Lata 80. przyniosły tzw. zimę AI, okres, w którym entuzjazm i finansowanie badań nad sztuczną inteligencją znacznie osłabły. Przyczyniły się do tego rozczarowania związane z ograniczeniami ówczesnych technologii i niemożnością spełnienia wygórowanych oczekiwań. Mimo to, w tym czasie pojawiły się ważne koncepcje, takie jak sieci neuronowe i algorytmy genetyczne, które zyskały na znaczeniu w późniejszych latach.

Przełom nastąpił w latach 90., kiedy AI zaczęła odnosić konkretne sukcesy. W 1997 roku komputer Deep Blue, stworzony przez IBM, pokonał mistrza świata w szachach, Garriego Kasparowa. To wydarzenie pokazało, że komputery mogą dorównać, a nawet przewyższyć ludzi w niektórych zadaniach wymagających inteligencji.

Kolejna fala postępu w AI nastąpiła w XXI wieku, głównie dzięki rozwojowi technologii uczenia maszynowego (ML) i głębokiego uczenia (DL). Znaczny wzrost mocy obliczeniowej komputerów, dostęp do ogromnych zbiorów danych oraz udoskonalenie algorytmów przyczyniły się do spektakularnych osiągnięć w tej dziedzinie. W 2012 roku system AlexNet, oparty na głębokich sieciach neuronowych, wygrał konkurs ImageNet, co zrewolucjonizowało przetwarzanie obrazów i otworzyło nowe możliwości dla AI.

Ostatnie lata przyniosły kolejne przełomowe wydarzenia, takie jak sukcesy systemu AlphaGo, stworzonego przez DeepMind, który pokonał mistrza świata w grze Go, jednej z najbardziej skomplikowanych gier strategicznych. Równocześnie rozwijają się technologie przetwarzania języka naturalnego (NLP), czego przykładem są modele takie jak GPT-3 i jego następcy, zdolne generować teksty na poziomie bliskim ludzkiemu.

Współczesna AI znajduje zastosowanie w niemal każdej dziedzinie życia, od medycyny i finansów, przez transport i energetykę, po rozrywkę i edukację. Dynamiczny rozwój tej technologii otwiera przed nami nowe

horyzonty, ale stawia także wyzwania związane z etyką, bezpieczeństwem i odpowiedzialnym jej wykorzystaniem.

Historia i ewolucja sztucznej inteligencji to opowieść o nieustannym dążeniu do stworzenia maszyn, które potrafią myśleć, uczyć się i działać w sposób inteligentny. Od pierwszych teoretycznych koncepcji po współczesne osiągnięcia, AI przeszła długą drogę, pełną wyzwań i sukcesów, kształtując przyszłość naszej cywilizacji.

Jak działa ChatGPT?

ChatGPT jest oparty na architekturze modelu GPT, czyli Generative Pre-trained Transformer, opracowanej przez OpenAI. Model ten jest przykładem głębokiej sieci neuronowej, która została specjalnie zaprojektowana do przetwarzania języka naturalnego. Architektura GPT opiera się na transformatorach, które są jednymi z najbardziej zaawansowanych narzędzi w dziedzinie uczenia maszynowego.

Transformator to rodzaj sieci neuronowej, który został wprowadzony w pracy "Attention is All You Need" przez Vaswaniego i jego zespół w 2017 roku. Kluczową innowacją transformatorów jest mechanizm self-attention, który pozwala modelowi przywiązywać wagę do różnych części sekwencji wejściowej przy generowaniu każdej części sekwencji wyjściowej. Dzięki temu model może skutecznie uczyć się zależności między słowami, nawet jeśli są one od siebie oddalone w tekście.

GPT, w szczególności, jest modelem transformatorowym przeszkolonym na dużych zbiorach danych tekstowych przy użyciu uczenia nadzorowanego. Proces szkolenia obejmuje dwa główne etapy: pre-trening i fine-tuning.

W fazie pre-treningu model jest uczony na ogromnych ilościach danych tekstowych pochodzących z różnych źródeł, takich jak książki, artykuły, strony internetowe itp. Celem jest nauczenie modelu zrozumienia

struktury języka i kontekstu słów. W tym etapie model uczy się przewidywać kolejne słowo w zdaniu, biorąc pod uwagę wszystkie poprzednie słowa, co pozwala mu budować wewnętrzne reprezentacje języka.

Po pre-treningu następuje etap fine-tuning, w którym model jest dostrajany na bardziej specyficznych zadaniach przy użyciu mniejszych, bardziej skoncentrowanych zbiorów danych. W przypadku ChatGPT fine-tuning obejmuje również techniki takie jak Reinforcement Learning from Human Feedback (RLHF), gdzie model jest trenowany przy użyciu opinii ludzi na temat generowanych odpowiedzi, co pozwala mu lepiej rozumieć intencje użytkowników i generować bardziej trafne odpowiedzi.

Sposób generowania tekstu przez ChatGPT opiera się na technice zwanej dekodowaniem sekwencyjnym. Gdy użytkownik wprowadza prompt, czyli zapytanie lub prośbę, model przetwarza ten tekst, aby zrozumieć jego kontekst i znaczenie. Następnie, przy użyciu swojej wewnętrznej reprezentacji języka, generuje kolejne słowa, jedno po drugim, aż do osiągnięcia pełnej odpowiedzi.

Mechanizm self-attention odgrywa kluczową rolę w tym procesie, ponieważ pozwala modelowi skupić się na istotnych częściach promptu podczas generowania każdej części odpowiedzi. Model bierze pod uwagę zarówno globalne, jak i lokalne zależności w tekście, co pozwala mu na tworzenie spójnych i logicznych odpowiedzi.

ChatGPT używa również technik takich jak temperatura i top-k sampling do kontrolowania kreatywności i różnorodności generowanych tekstów. Temperatura to parametr, który wpływa na stopień losowości w generowaniu tekstu. Wyższa temperatura prowadzi do bardziej zróżnicowanych odpowiedzi, podczas gdy niższa temperatura sprawia, że odpowiedzi są bardziej przewidywalne. Top-k sampling polega na wybieraniu kolejnych słów spośród k najbardziej prawdopodobnych opcji, co zapobiega generowaniu mało sensownych lub nieadekwatnych odpowiedzi.

ChatGPT jest zatem zaawansowanym modelem językowym, który wykorzystuje transformatorową architekturę, mechanizmy self-attention i techniki uczenia maszynowego, aby generować spójne, kontekstowe i trafne odpowiedzi na zapytania użytkowników. Dzięki procesom pre-treningu i fine-tuningu, model ten jest w stanie rozumieć i przetwarzać złożone zapytania, dostarczając wartościowych informacji i interakcji w szerokim zakresie zastosowań.

2: Tworzenie Skutecznych Promptów

