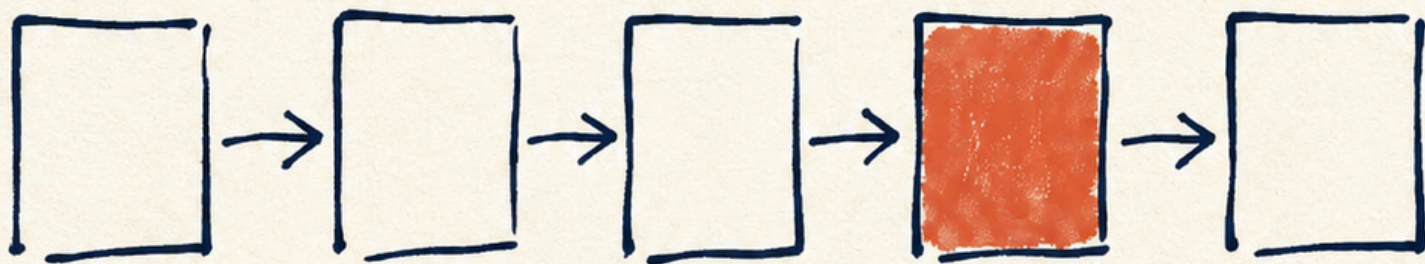


Automatyzacje bez kodowania

Wprowadzenie do **n8n** w marketingu.



Marek Staniszewski

Heuristica

Automatyzacje bez kodowania

Wprowadzenie do n8n w marketingu

AUTOR	Marek Staniszewski
COPYRIGHT	© Marek Staniszewski, 2026
WYDAWCA	Heuristica — Staniszewski Marek, Warszawa staniszewski@heuristica.pl · heuristica.pl
WYDANIE	Pierwsze, wydanie internetowe, 2026
ADRES WYDANIA	heuristica.pl/automatyzacje-n8n
POWSTANIE	Tekst wygenerowany z udziałem generatywnej AI (modele Claude, Anthropic), zredagowany przez autora
SKŁAD	Playfair Display, Literata, Archivo

POWSTAŁO Z UDZIAŁEM GENERATYWNEJ AI

Ta książka powstała z udziałem generatywnej sztucznej inteligencji. Tekst rozdziałów został wygenerowany przez modele językowe Claude (Anthropic) na podstawie struktury, briefu i działających przepływów przygotowanych przez autora, a następnie zredagowany przez autora.

Każdy opisany krok został wykonany ręcznie na żywym koncie n8n i sprawdzony przed publikacją; wszystkie zrzuty ekranu pochodzą z rzeczywistych uruchomień opisywanych przepływów. Odpowiedzialność za treść, jej poprawność i wnioski ponosi autor.

LICENCJA

Treść udostępniona na licencji Creative Commons Uznanie autorstwa - Użycie niekomercyjne - Na tych samych warunkach 4.0 Międzynarodowe (CC BY-NC-SA 4.0). Wolno kopiować, rozpowszechniać i przetwarzać materiał, podając autora i źródło, bez użycia komercyjnego, na tej samej licencji.

Pliki przepływów (.json) możesz używać bez ograniczeń, także w pracy zarobkowej — licencja dotyczy tekstu, nie tego, co z nim zbudujesz. Przepływy do pobrania: heuristica.pl/automatyzacje-n8n

ZNAKI TOWAROWE

n8n jest znakiem towarowym n8n GmbH. OpenAI, ChatGPT i GPT są znakami towarowymi OpenAI. Google, Gmail i Google Sheets są znakami towarowymi Google LLC. Nazwy użyte wyłącznie w celach informacyjnych. Ceny, nazwy przycisków i wygląd interfejsów są aktualne na sierpień 2026 i mogą się zmienić.

Spis treści

OTWARCIE	Przedmowa	4
CZĘŚĆ 1. START		
ROZDZIAŁ 1	Uruchamiamy n8n	7
ROZDZIAŁ 2	Co widzisz na ekranie	13
ROZDZIAŁ 3	Dodajemy pierwszy własny węzeł	17
ROZDZIAŁ 4	Podłączamy AI	26
CZĘŚĆ 2. TRZY SYSTEMY		
ROZDZIAŁ 5	Monitor konkurencji	36
ROZDZIAŁ 6	Syntezaator researchu	45
ROZDZIAŁ 7	Strategiczny panel trzech ekspertów	57
CZĘŚĆ 3. NA KONIEC		
ROZDZIAŁ 8	Sześć klocków, z których zbudujesz resztę	68
NA KONIEC	Słowniczek	72

Przedmowa

Większości narzędzi AI uczyłem się od... AI.

Istnieją rzecz jasna zylardy tutoriali, artykułów, filmów na YouTube czy webinarów płatnych i bezpłatnych, gdzie wszystko jest dokładnie omówione.

Poznając jednak takie rozwiązania jak Make.com, Apify, Pinecone czy inne narzędzia - choćby Google Colab, Jupyter, notatniki Wolfram Language itp., których chciałem użyć w różnych projektach - pytam najczęściej Claude'a lub ChatGPT i robię kolejne eksperymenty, dopytując model o szczegóły.

LLM jako tutor ma wiele wad. Odpowiedzi są często niepotrzebnie zbyt obszerne (nawet po promptowym napomnieniu), czasem zaś zbyt zaawansowane albo banalne - czyli nieadekwatne do poziomu wiedzy. Model pomimo tego, że wie już wiele na mój temat nie zawsze trafnie odczytuje aktualne intencje i potrzeby informacyjne.

Dużą zaletą takiej nauki jest jednak fakt, że model przez cały czas niezmordowanie i cierpliwie coś tłumaczy - jest nieustannie dostępny i pozostaje niewzruszony przy zadawaniu nawet najgłupszych pytań...

Podobnie uczyłem się też obsługi platformy n8n. Obecnie używam jej do wielu projektów jako wsparcia np. w chatbotach osadzanych na stronach www czy modułów CRM.

Lekcje, jakie wyciągnąłem z tych wszystkich przygód uczenia się czegoś z AI, sprowadzają się do tego, że skuteczne poznawanie nowych tematów z chatbotem jako tutorem wymaga pewnej narzuconej struktury i dyscypliny tego, jak się uczyć. Gdybym miał je podsumować w kilku punktach, brzmiałyby następująco:

- Chaotyczne pytania i skakanie z tematu na temat powiększają jedynie chaos i do niczego nie prowadzą.
- Zaczynając przygodę z jakimś nowym narzędziem warto wyznaczyć sobie jakiś konkretny cel: co ja właściwie chcę z tym osiągnąć, do czego zamierzam tego używać i co chcę wiedzieć?
- Sesję uczenia się warto też dzielić na wyraźne partie - np. bloki tematyczne. Warto też ustalić jakąś granularność - kolejne partie powinny być ze sobą powiązane i powinny przynosić jakościowo nowe rzeczy.
- Ważna jest tu również praca własna - ćwiczenia, samodzielne eksperymenty, powtórki, utrwalanie wiedzy.

Niniejszy tutorial jest więc projektem eksperymentalnym. Próbą sprawdzenia na ile modele językowe mogą być przydatne w tworzeniu spójnych treści i materiałów edukacyjnych. Kierowany jest głównie do osób, które nie używały jeszcze n8n, nie znają się na kodowaniu i nie chcą wnikać zbyt głęboko w kwestie techniczne. Chciałyby się jednak szybko nauczyć korzystać z tej platformy i tworzyć własne automatyzacje do celów marketingowych.

Jak powstawał ten poradnik

Punktem wyjścia były konkretne przepływy (workflows), z których na co dzień korzystam. Zaadaptowałem je do tego materiału w ten sposób, aby mogły ilustrować ogólne zasady pracy na n8n. Przykłady starałem się dobierać tak, by czytelnik mógł kolejno poznawać tzw. „węzły” powtarzające się w wielu zadaniach.

Następnie, korzystając z Claude (różne modele), stworzyłem strukturę całości. Utworzony został katalog, w którym Claude Code pisał kolejne rozdziały w oparciu o zadane przepływy.

Każdy rozdział redagowałem, usuwając to, co uznałem za zbędne, i uzupełniając własnymi notatkami.

Wszystkie opisywane przez Claude operacje testowałem ręcznie na n8n, sprawdzając wykonania, i nanosiłem poprawki do generowanych opisów. Każdy ważniejszy krok dokumentowany był również zrzutem ekranu obrazującym omawiane zagadnienie.

Na ile był to eksperyment udany?

Ten osąd pozostawiam już Czytelnikom, którzy także lubią eksperymentować i zechcą spróbować uczyć się n8n w ten właśnie sposób. Zapraszam również do zgłaszania poprawek, uwag, sugestii zmian czy kierunków rozbudowy tego ebooka. Projekt ma charakter otwarty i dzięki takim uwagom mogą powstawać kolejne, ulepszone wersje służące kolejnym Czytelnikom.

Tyle ode mnie. Resztę słów poniżej napisało już AI...

Marek Staniszewski

Sierpień 2026

Wstęp: Zawrzyjmy pewien kontrakt...

Zacznijmy od dwóch pojęć, które w ostatnich dwóch latach zrobiły zawrotną karierę: „automatyzacja” i „agent”. Oba stały się wytrychami używanymi praktycznie do wszystkiego - od prostego formularza kontaktowego po coś, co ktoś nazwał „inteligencją”, żeby sprzedać drożej. Kiedy słowo pasuje do wszystkiego, przestaje cokolwiek znaczyć.

W tej książce słowo „agent” znaczy dokładnie tyle: kilka kroków połączonych strzałkami, które robią za ciebie to, co i tak robiłbyś ręcznie - tylko że robią to codziennie, na przykład o siódmej rano, bez twojego udziału i bez marudzenia. Żadnej rewolucji, żadnej maszyny myślącej za ciebie.

To, co tu robimy, nie jest zresztą nowe. Marketerzy programowali swoje narzędzia, odkąd istnieje arkusz kalkulacyjny: makro liczące marże, reguła w skrzynce mailowej przekierowująca leady do właściwej osoby. n8n jest kolejnym wcieleniem tego samego pomysłu. Różnica jest jedna: tym razem jeden z kroków potrafi przeczytać tekst i ocenić, czy jest istotny. **Automatyzacja, o której**

będzie mowa w tej książce, nie rozwiązuje problemu, że nie umiesz programować, tylko problem, że masz na coś czas raz w tygodniu, a potrzebujesz mieć go codziennie.

Ile z tego wymaga programowania? Jeśli boisz się tego słowa, dobra wiadomość jest taka, że ani jednej linijki. Nie musisz umieć programować, nie napiszesz w tej książce ani jednego wiersza kodu - wszystko robisz myszą: przeciągasz, klikasz, wpisujesz tekst w pole. Jeśli w którymś miejscu na ekranie pojawi się coś, co wygląda jak kod, zostanie to wyjaśnione dokładnie wtedy, kiedy się pojawi.

Potrzebujesz trzech rzeczy: przeglądarki internetowej, karty płatniczej do założenia konta na modelu językowym (koszt liczony w groszach za uruchomienie, nie w złotych miesięcznie) i dwóch godzin, najlepiej jednym ciągiem - to twój czas do pierwszego działającego systemu, opisanego w rozdziale 1. Każdy z trzech głównych systemów w części drugiej tej książki to potem jeszcze jeden wieczór pracy. Samego n8n nie musisz nigdzie instalować - do konta w n8n karta nie jest potrzebna, dowiesz się dlaczego już w rozdziale 1.

Będzie moment, w którym się zgubisz. To normalne - w rozdziale 4 celowo pokażę ci błąd, żebyś zobaczył, jak wygląda w moim towarzystwie, zanim zobaczysz go sam. Jeśli utkniesz na dłużej, każdy rozdział kończy się sekcją „Co może pójść nie tak”, z dosłownymi komunikatami błędów i sposobem naprawy. Zacznij od niej, zanim zaczniesz szukać w internecie.

Nie obiecuję ci rewolucji ani konkretnej liczby zaoszczędzonych godzin - to zależy od tego, co akurat robisz w swojej pracy. Obiecuję trzy działające systemy i to, że po drodze przestaniesz się bać słowa, które właśnie zdjęliśmy z piedestału.

Rozdział 1. Uruchamiamy n8n

Zanim zrozumiesz, jak działa n8n, proponuję po prostu zobaczyć, jak to wszystko wygląda w praktyce. To być może kolejność odwrotna do tej, której uczono nas w szkole, gdzie zwykle najpierw jest teoria, a dopiero później zadanie. **Nic nie uczy jednak szybciej niż działający system na własnym ekranie - i nic nie zniechęca skuteczniej niż teoria, której nie ma jeszcze do czego przyłożyć.**

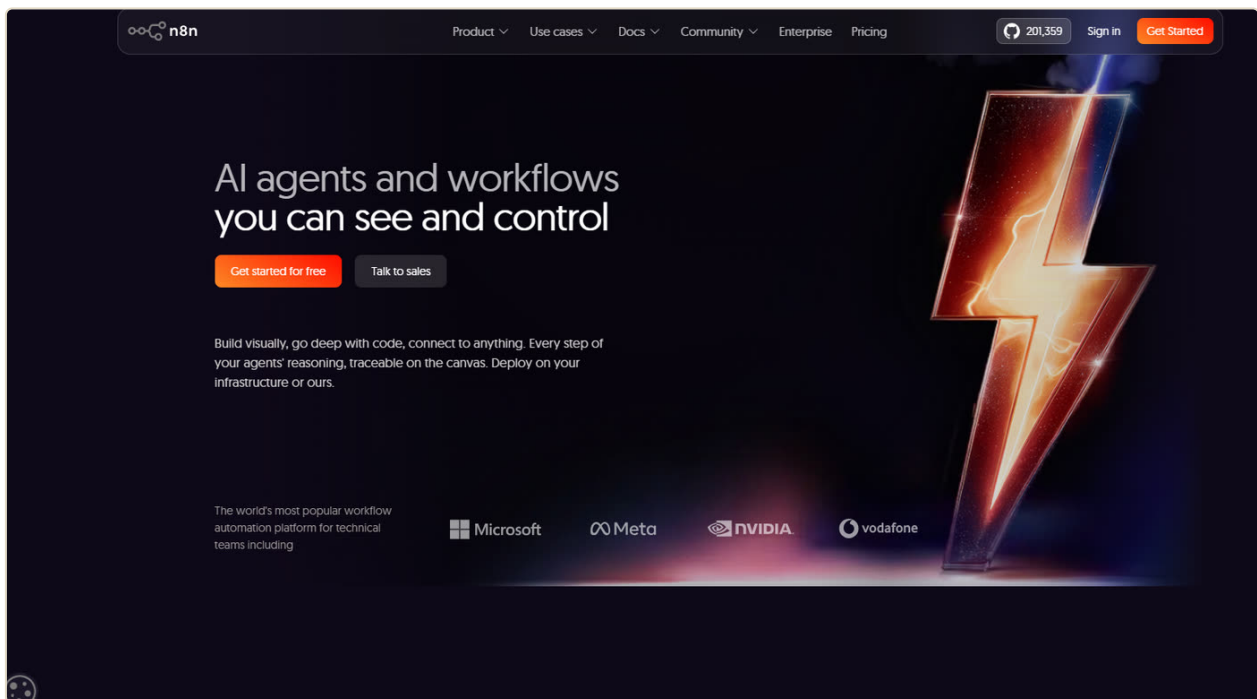
W tym rozdziale nie będę ci więc jeszcze wyjaśniał, czym jest węzeł ani dlaczego łączy się go w taki, a nie inny sposób. Zrobimy na razie tylko jedną rzecz: założysz konto, klikniesz parę razy myszką i zobaczysz na swoim ekranie gotowy, działający zautomatyzowany system, który sam pobierze dane i coś z nimi zrobi.

Poćwiczymy więc od razu import gotowego przepływu, jedno kliknięcie i sprawdzimy wynik. Zaczynamy więc od konta.

Zakładasz konto

n8n działa w dwóch wariantach: jako usługa w chmurze (n8n Cloud), na której nic nie instalujesz, albo jako coś, co stawiasz na własnym serwerze - w tym również lokalnie, na własnym komputerze. Tutaj omawiany będzie wyłącznie ten pierwszy sposób.

- 1 Otwórz przeglądarkę i wejdź na stronę n8n.io. W prawym górnym rogu kliknij przycisk „Get started” (rozpocznij).



Rys. 1.1.1. Punkt startowy - strona główna n8n.io

NA EKRANIE ZOBACZYSZ

formularz rejestracji z polami na e-mail, hasło i nazwę.

- 2 Na formularzu podaj swój adres e-mail, ustaw hasło i wpisz imię i nazwisko jako nazwę wyświetlaną na koncie. Kliknij przycisk pod formularzem, żeby przejść dalej.

! UWAGA

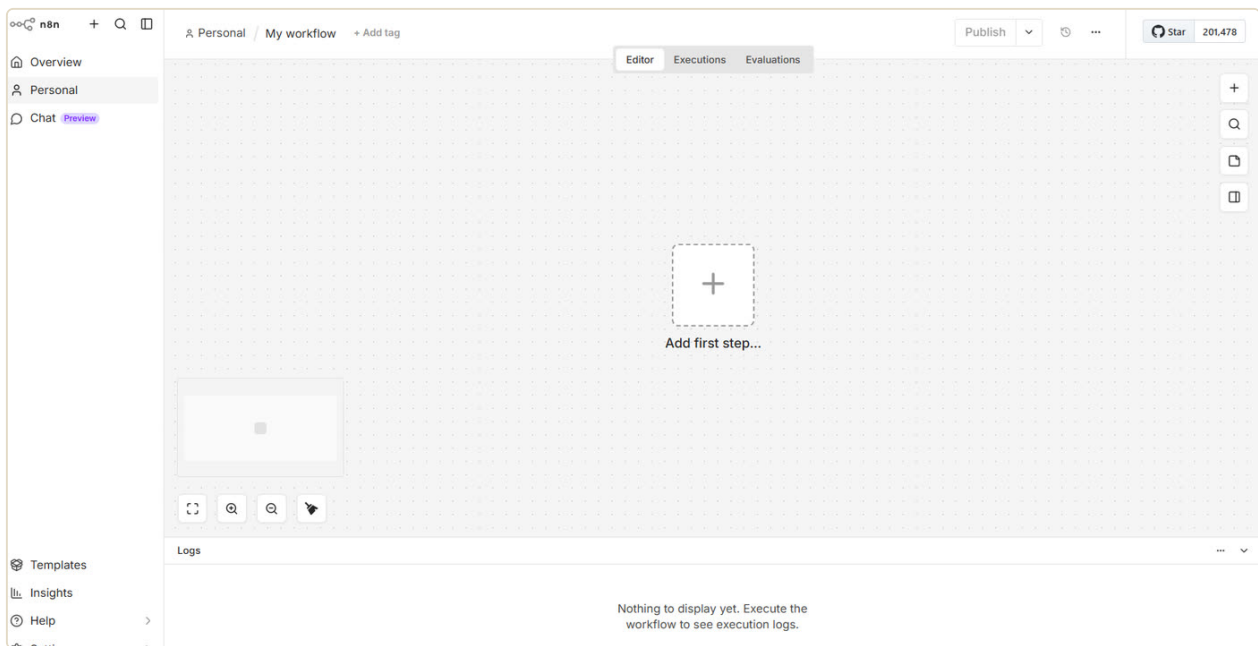
Na tym etapie n8n nie prosi o numer karty płatniczej. Plan startowy (Starter) i plan Pro uruchamiają darmowy okres próbny bez karty. Karta będzie potrzebna dopiero przy koncie na modelu językowym.

- 3 n8n poprosi cię o nazwę przestrzeni roboczej - fragmentu adresu, pod którym będzie działał twój n8n, na przykład „twoja-nazwa.app.n8n.cloud”. Wpisz cokolwiek zapamiętywalnego i przejdź dalej.
- 4 Potwierdź regulamin i zakończ rejestrację. Na ekranie zobaczysz: komunikat, że na twój adres e-mail wysłano wiadomość z potwierdzeniem.
- 5 Otwórz swoją skrzynkę pocztową. Znajdziesz w niej mail z linkiem potwierdzającym adres e-mail. Kliknij ten link.

Sprawdź, czy działa: po kliknięciu linku przeglądarka przenosi cię z powrotem do n8n, na ekran twojej przestrzeni roboczej. W lewym górnym rogu zobaczysz nazwę, którą wybrałeś w kroku 3. Zobaczysz też duży przycisk „Create Workflow” (utwórz przepływ).

Pierwszy widok Twojego przepływu

Kliknij „Create Workflow”. n8n otworzy nowy, pusty obszar roboczy przepływu - dalej będziemy nazywać go po prostu obszarem roboczym. Zobaczysz tło w drobną kratkę i jeden duży przycisk na środku: „Add first step” (dodaj pierwszy krok).

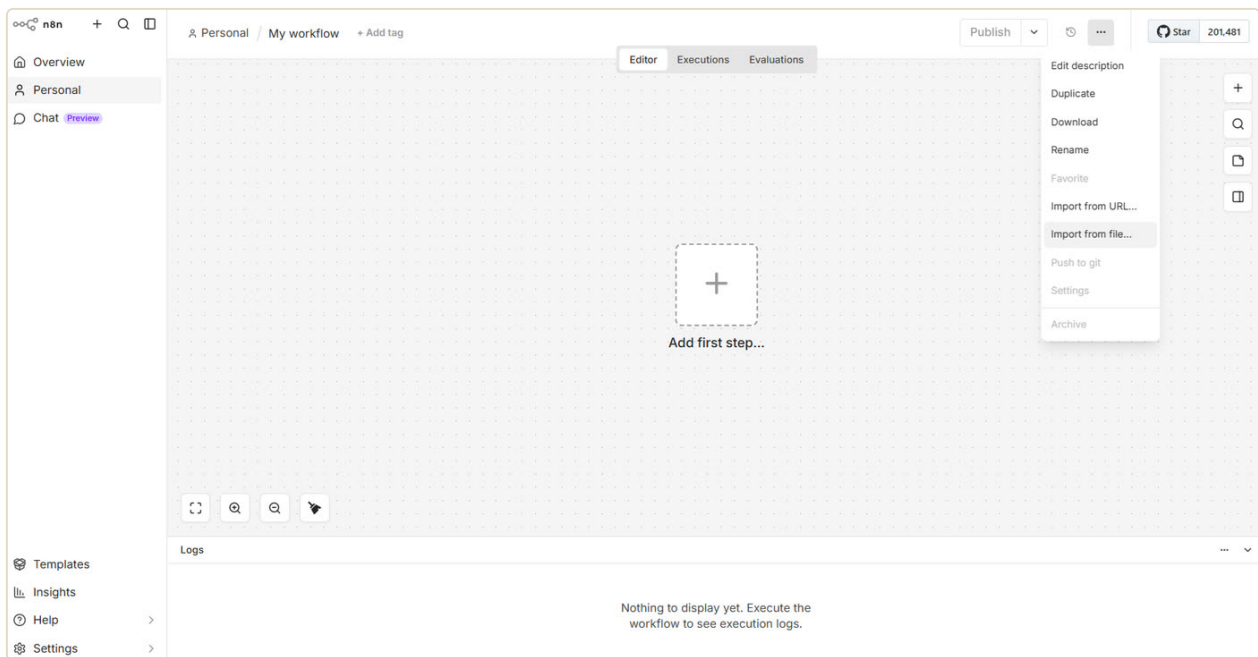


Rys. 1.2. Pusty obszar roboczy przed dodaniem czegokolwiek

Nie klikaj jednak tego przycisku. Zamiast budować coś od zera, zaimportujemy na początek gotowy przepływ - zobaczysz cały działający system w kilka sekund, a od zera zaczniemy dopiero w rozdziale 3.

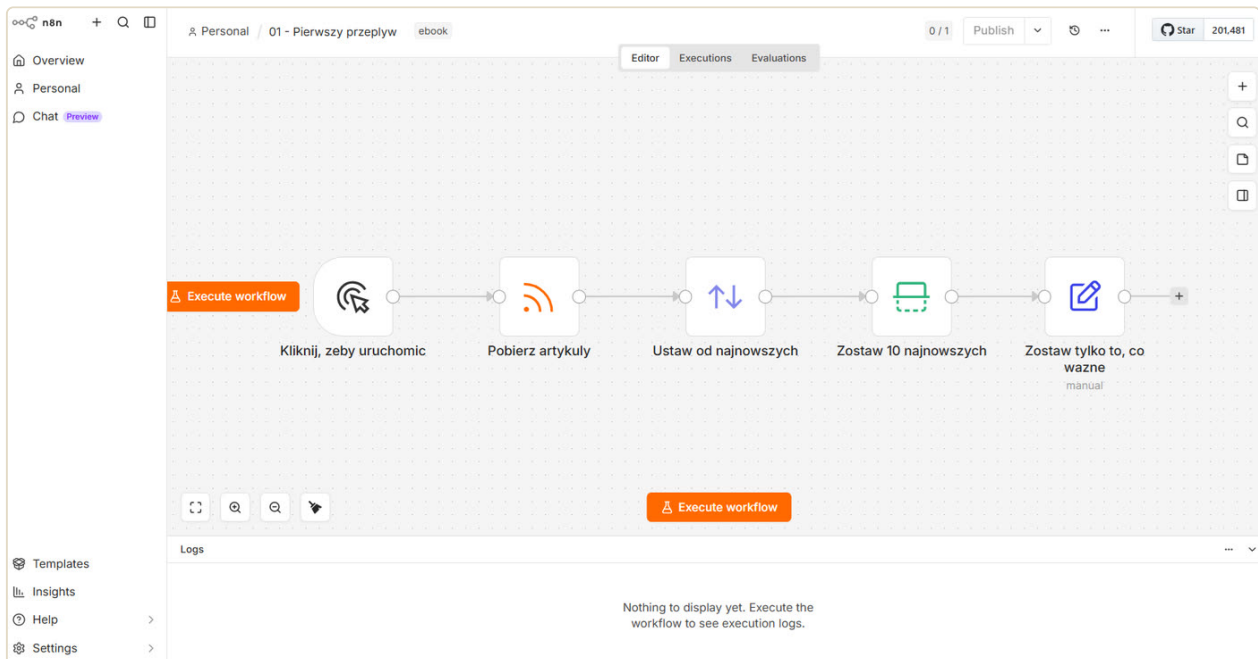
Importujesz gotowy przepływ

- 1 Pobierz plik `01-pierwszy-przeplyw.json` ze strony heuristica.pl/automatyzacje-n8n, z sekcji „Pobierz przepływy”. Znajdziesz tam wszystkie przepływy z tego poradnika. Zapisz go w miejscu, które łatwo znajdziesz, na przykład na swoim pulpicie.
- 2 W obszarze roboczym kliknij trzy kropki (...) w prawym górnym rogu ekranu. Z rozwiniętego menu wybierz „Import from File” (importuj z pliku).



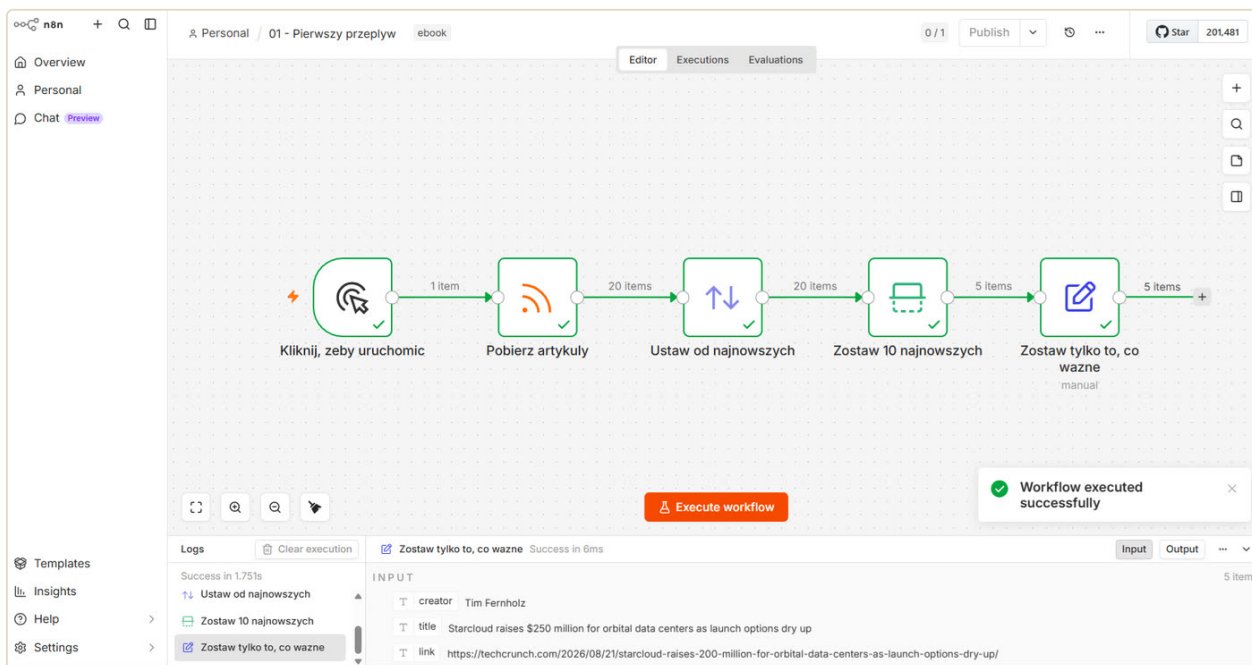
Rys. 1.3. Menu importu przepływu

- 3 W oknie systemowym, które się otworzy, znajdź i zaznacz pobrany plik `01-pierwszy-przeplyw.json`, a potem potwierdź wybór.
- 4 Na ekranie natychmiast pojawi się gotowy przepływ: pięć prostokątów połączonych liniami ze strzałkami, w jednym poziomym rzędzie. To są węzły i połączenia między nimi - czym dokładnie są, dowiesz się w rozdziale 2. Teraz przyjrzymy im się bliżej.



Rys. 1.4. Zaimportowany przepływ, gotowy do uruchomienia

Na ekranie każdy węzeł, przez który przeszły dane, obrysowuje się na zielono, a pod nim pojawia się mała liczba - to liczba elementów, które przez niego przeszły. Jeśli wszystkie węzły w przepływie zmienią kolor na zielony, to znaczy, że system zadziałał - wykonał się prawidłowo od początku do końca.



Rys. 1.5. Przepływ po uruchomieniu - zielone obwódki oznaczają sukces

Zielona obwódka mówi, że system zadziałał. Nie mówi jeszcze, co dla ciebie zrobił - to sprawdzimy w ostatnim kroku.

- 6 Kliknij teraz dwukrotnie ostatni węzeł, „Zostaw tylko to, co ważne”. W panelu, który się otworzy z prawej strony, przełącz widok na Table. Zobaczysz listę pięciu artykułów pobranych z serwisu technologicznego, każdy z tytułem, datą publikacji i treścią - prawdziwe dane, które przepływ przed chwilą dla ciebie zebrał, posortował od najnowszego i oczyścił z tego, co niepotrzebne.

The screenshot displays the n8n workflow editor. On the left, the 'INPUT' section shows a list of 10 items from a JSON feed. The 'Parameters' section in the center shows a 'Manual Mapping' mode with fields for 'Tytuł', 'Data', and 'Content' mapped to JSON paths like '\$json.title' and '\$json.isoDate'. The 'OUTPUT' section on the right shows a table with 5 items, each containing a title, date, and content.

Rys. 1.6. Prawdziwy wynik - pięć najnowszych artykułów zebranych przez przepływ

I to wszystko. Właśnie udało ci się uruchomić twój pierwszy system w n8n - i widzisz, co konkretnie dla ciebie zrobił.

? DLA CIEKAWYCH

Nazwa „n8n” to skrót zapisany jak tablica rejestracyjna: litera „n”, osiem liter pominiętych, litera „n” - czyli „nodemation”, automatyzacja przez węzły. Ten sam trik zapisu znajdziesz w słowach i18n (internationalization) czy k8s (Kubernetes). Jest to żargon pragmatycznych programistów, u którego korzeni leży zwykła oszczędność miejsca w terminalu.

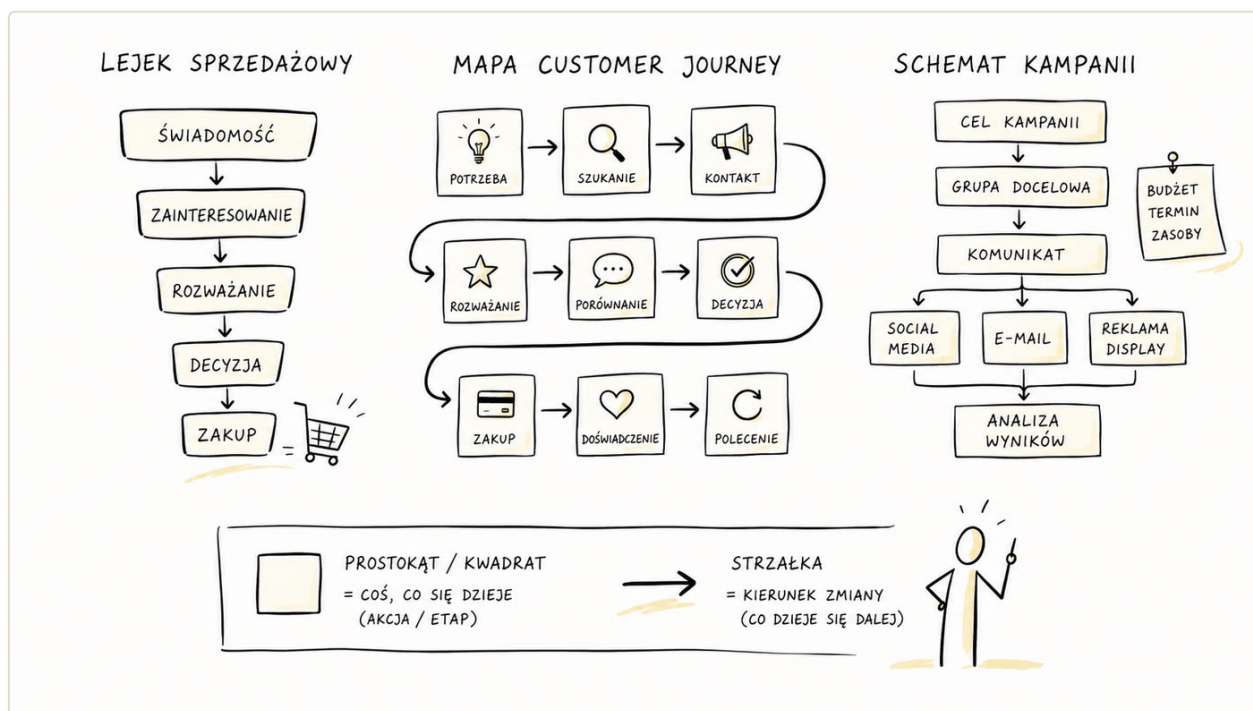
Wariacje na temat

- 1 Otwórz zaimportowany przepływ jeszcze raz i kliknij pojedynczo każdy węzeł. Zobacz, czy potrafisz zgadnąć, co robi, zanim przeczytasz o tym w rozdziale 2.
- 2 Zmień nazwę przepływu (kliknij jego tytuł w lewym górnym rogu ekranu i wpisz nową nazwę) i sprawdź, czy „Execute workflow” nadal działa po zmianie.

Nie wiesz jeszcze, czym są „węzły” ani dlaczego strzałki łączą je akurat w tej kolejności - i to jest na tym etapie OK. Dowodem na to, że coś w n8n działa, jest na razie zielona obwódka na ekranie. W kolejnych częściach tego poradnika zrobimy inne ćwiczenia, które pozwolą ci lepiej zrozumieć, co naprawdę się pod tym wszystkim kryje.

Rozdział 2. Co widzisz na ekranie

Ten rysunek widziałeś już pewnie dziesiątki razy, tylko w innym kontekście. Lejek sprzedażowy, mapa customer journey, albo schemat kampanii przypięty pinezkami do tablicy w agencji - wszystkie składają się z tego samego: prostokątów, tudzież kwadratów i strzałek pomiędzy nimi. Prostokąt albo kwadrat to zwykle coś, co się dzieje - jakaś akcja. Strzałka mówi, jaki jest kierunek zmiany - co dzieje się dalej.



Rys. 2.0. Każdy proces to swoisty przepływ

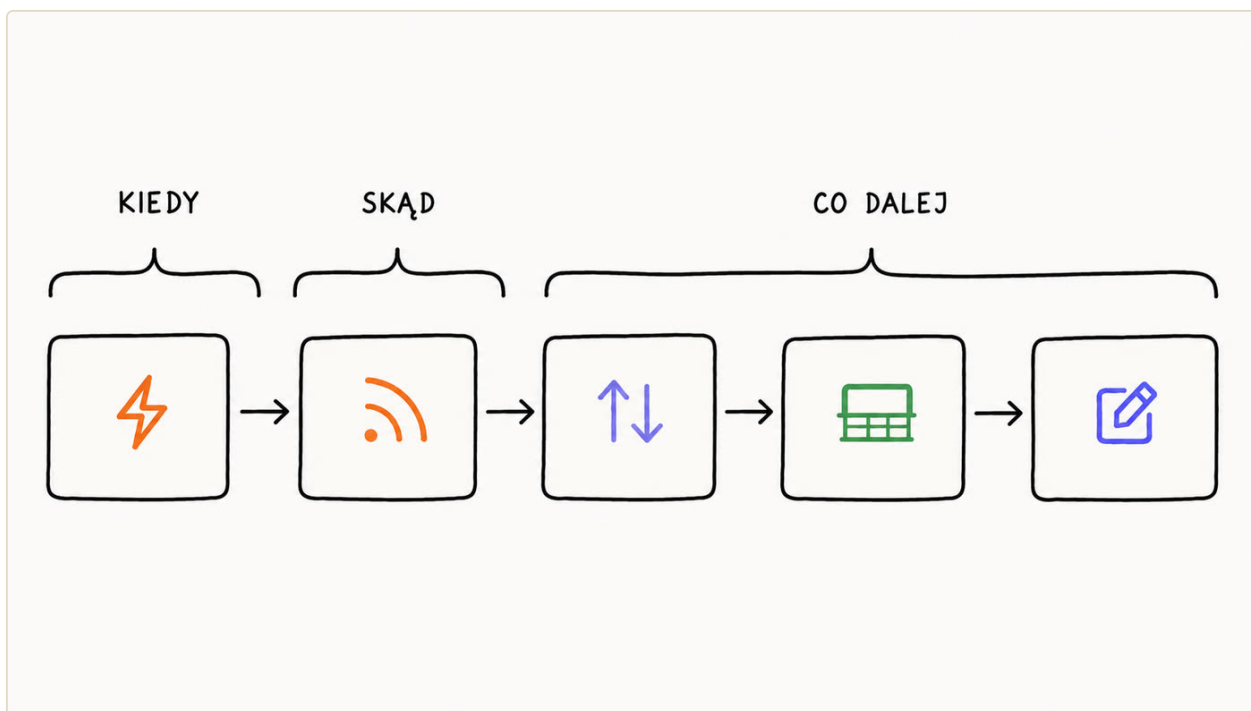
Węzeł robi jedną rzecz, połączenie mówi, dokąd trafiają jej wyniki - i to właściwie cała gramatyka n8n.

n8n od schematu na slajdzie różni się jedynie tym, że te kwadraty nie tylko ilustrują pomysł, ale naprawdę coś robią, kiedy klikniesz „uruchom”. W tym rozdziale nauczysz się lepiej czytać takie schematy i sprawdzać, czy to, co widzisz, zgadza się z tym, co faktycznie się dzieje.

Węzeł i połączenie

Wróćmy do przepływu, który zaimportowałeś w rozdziale 1. Jeśli go zamknąłeś, znajdziesz go na liście swoich przepływów zaraz po zalogowaniu.

Każdy z kwadratów na ekranie to węzeł. Węzeł pobiera dane, coś z nimi robi i przekazuje dalej - zawsze dokładnie w tej kolejności. Strzałka łącząca dwa węzły to połączenie: pokazuje, dokąd trafia wynik pracy z danego węzła, żeby stał się materiałem wejściowym dla węzła kolejnego.



Rys. 2.1. Przykładowy przepływ

Pierwszy węzeł w każdym przepływie ma szczególną rolę: to wyzwalacz. Wyzwalacz decyduje, kiedy cały przepływ ma się uruchomić - o określonej godzinie, po nadejściu nowego maila, po kliknięciu przycisku. Bez wyzwalacza przepływ istnieje na ekranie, ale nigdy się sam nie uruchomi.

? DLA CIEKAWYCH

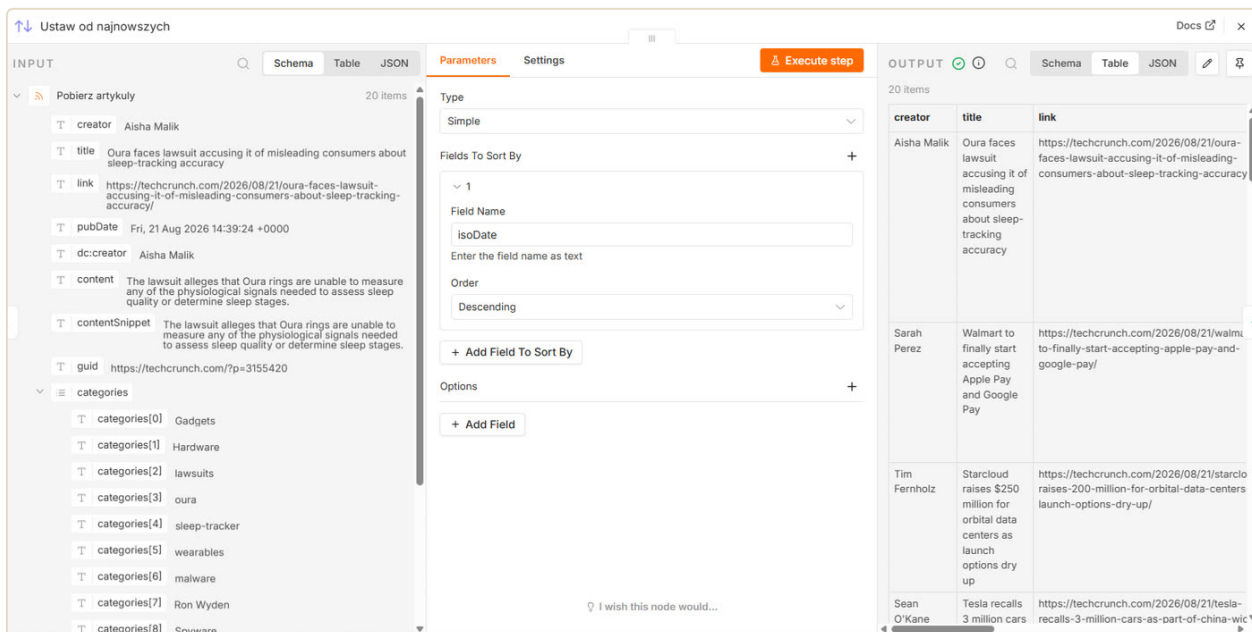
Diagram złożony z prostokątów i strzałek nie jest wynalazkiem informatyki. Frank Gilbreth przedstawił „process charts” - schematy procesu zapisane właśnie jako bloki połączone przepływem - na forum inżynierów ASME w 1921 roku. Były to dekady przed pierwszym komputerem. n8n nie wymyśla więc nowego sposobu myślenia o procesie, tylko pozwala takie procesy wykonywać.

Klik na węzeł: co weszło, co wyszło?

To jest jedyne narzędzie diagnostyczne, którego naprawdę potrzebujesz. Zapamiętaj ten mechanizm teraz - wrócisz do niego w każdym kolejnym rozdziale.

- 1 W przepływie z rozdziału 1 kliknij dwukrotnie węzeł „Ustaw od najnowszych” - do węzła wyzwalacza jeszcze wrócimy. Na ekranie zobaczysz: panel z jego ustawieniami.
- 2 Po bokach zobaczysz dwa panele danych: „INPUT” (wejście) po lewej i „OUTPUT” (wyjście) po prawej. Jeśli przepływ był już wcześniej uruchomiony, w obu zobaczysz dane - dokładnie to, co do tego węzła weszło i co z niego wyszło.

- 3 Dane możesz oglądać w trzech widokach: Schema (uproszczona struktura pierwszego elementu), Table (tabela) i JSON. Zaczniij od widoku Table - dla większości przypadków, które będziemy omawiać, jest najbardziej czytelny.



Rys. 2.2. Panel danych węzła po dwukrotnym kliknięciu

- 4 Zamknij panel (X w prawym górnym rogu albo klawisz Escape) i kliknij dwukrotnie inny węzeł w łańcuchu. Na ekranie zobaczysz: panel tego węzła, z własnym „INPUT” i „OUTPUT”. Sprawdź, czy jego „INPUT” wygląda tak samo jak „OUTPUT” poprzedniego węzła - powinno być identyczne, bo to właśnie robi połączenie między nimi: przenosi dane bez zmian z wyjścia jednego węzła na wejście drugiego.

Uruchamianie przepływu

- 1 Cały przepływ uruchamiasz przyciskiem „Execute workflow” (uruchom przepływ) w prawym dolnym rogu ekranu - poznałeś go już w rozdziale 1. Na ekranie zobaczysz: węzły podświetlające się kolejno, w miarę jak dane przez nie przechodzą.
- 2 Pojedynczy węzeł możesz również uruchomić osobno, bez odpalania całego przepływu. W tym celu otwórz go dwukrotnym kliknięciem. Następnie kliknij „Execute step” (wykonaj krok). Na ekranie zobaczysz: dane pojawiające się w panelu „OUTPUT” tylko tego jednego węzła, bez uruchamiania pozostałych. Przydaje się, kiedy przepływ ma np. trzydzieści różnych węzłów, a sprawdzić chcesz tylko trzeci.

Co może pójść nie tak

- **Panel „INPUT” albo „OUTPUT” jest pusty, mimo że kliknąłeś węzeł.** Przepływ nie został jeszcze uruchomiony w tej sesji przeglądarki - dane pojawiają się dopiero po kliknięciu

„Execute workflow”. Uruchom przepływ, a potem otwórz węzeł jeszcze raz.

- **„Execute step” nie reaguje na kliknięcie.** Upewnij się, że poprzedni węzeł w łańcuchu został już wcześniej uruchomiony przynajmniej raz - część węzłów potrzebuje prawdziwych danych z poprzedniego kroku, żeby dało się je przetestować osobno.

Wariacje na temat

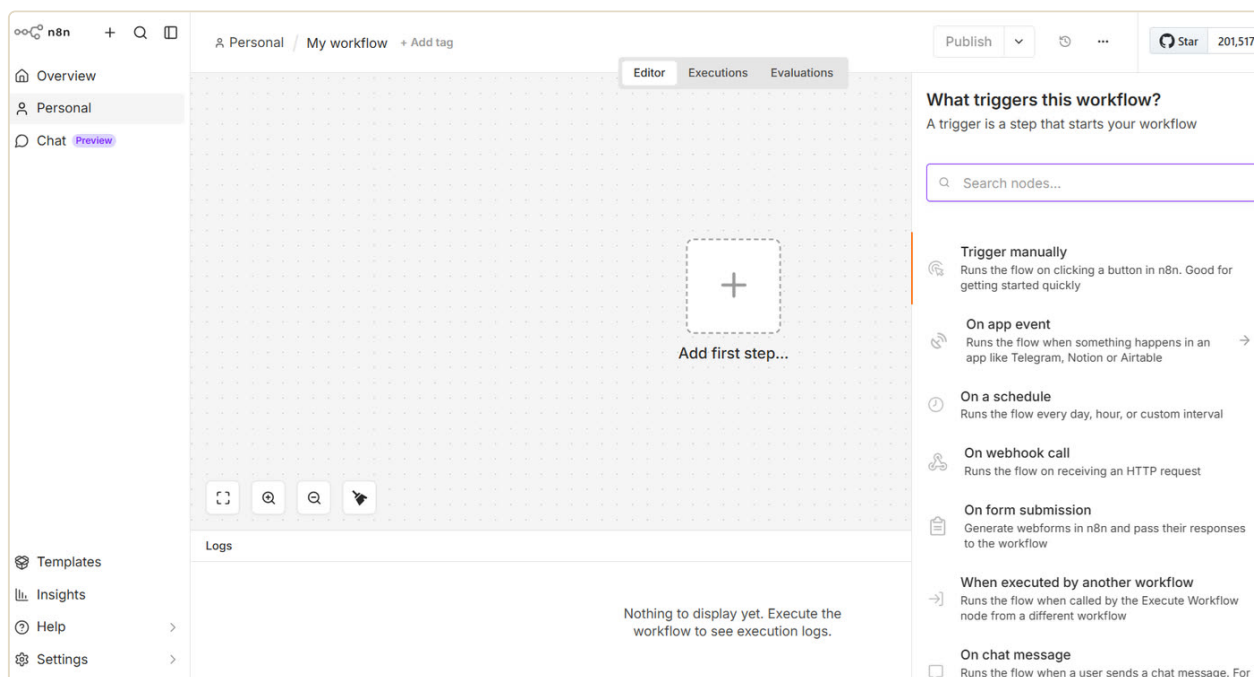
- 1 Otwórz każdy węzeł w przepływie z rozdziału 1 po kolei i zanotuj jednym zdaniem, co on robi - wyłącznie na podstawie tego, co widzisz w „INPUT” i „OUTPUT”.
- 2 Zmień widok danych z Table na JSON dla jednego węzła. Nie musisz rozumieć składni - zwróć tylko uwagę, że to te same dane, tylko inaczej zapisane.

Rozdział 3. Dodajemy pierwszy własny węzeł

W rozdziale 1 zaimportowałeś gotowy przepływ - twoja rola ograniczała się do kliknięcia. Teraz zbudujesz coś od zera samodzielnie. **Zaimportowany przepływ pokazuje, że coś działa; dopiero zbudowany własnoręcznie pokazuje, dlaczego.** Zbudujesz więc prosty system, który pobierze artykuły z bloga, zostawi tylko te na wybrany temat i wyśle je do ciebie mailem. To dokładnie ta sama idea, jaką znajdziesz w rozdziale 5, tyle że na razie bez udziału modelu językowego - tam nauczysz się, jak AI ocenia trafność zamiast prostego dopasowania słowa.

Zaczynamy od pustego przepływu

- 1 Wróć do listy swoich przepływów i utwórz nowy, pusty przepływ. Zobaczysz znajomy obszar roboczy z przyciskiem „Add first step” na środku.
- 2 Kliknij „Add first step” (dodaj pierwszy krok). Na ekranie zobaczysz: listę kategorii wyzwalaczy do wyboru.

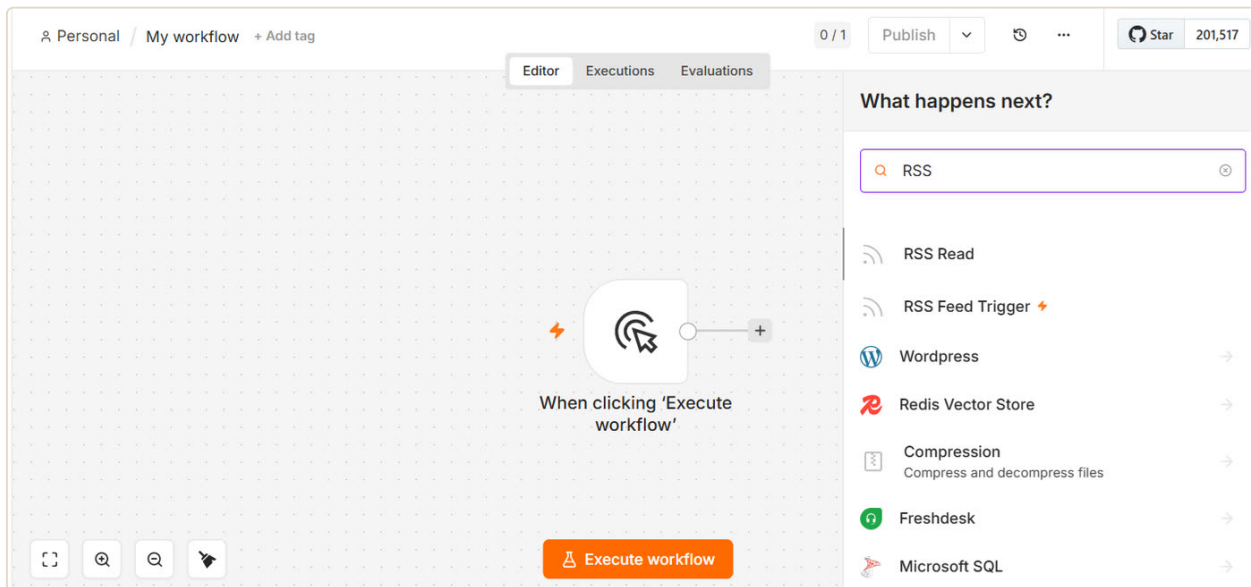


Rys. 3.1. Wybór pierwszego węzła - wyzwalacza

- 3 W wyszukiwarce po prawej („Search nodes...”) wpisz i wyszukaj „Manual trigger” (uruchom ręcznie) - to dokładnie ten sam typ wyzwalacza, którego używałeś w rozdziale 1. Na ekranie zobaczysz: pojedynczy węzeł na obszarze roboczym, gotowy do rozbudowy.

Dodajesz drugi węzeł: pobieranie danych

- 4 Najedź teraz kursorem poza prawą krawędź węzła wyzwalacza, tam gdzie na końcu linii widzisz mały znak plus. Kliknij go. Na ekranie znowu zobaczysz pole wyszukiwania węzłów.
- 5 Wpisz „RSS” w polu wyszukiwania, znajdź i kliknij węzeł „RSS Read” (czytaj RSS) z listy wyników.



Rys. 3.2. Wyszukiwanie węzła po nazwie

- 6 Kliknij dwukrotnie węzeł „RSS Read” i w polu „URL” wklej adres kanału RSS: `https://techcrunch.com/feed/`. Jeśli wolisz śledzić inny temat, to każdy działający adres RSS zadziała tak samo.

WSKAZÓWKA

Większość blogów firmowych i portali branżowych ma adres RSS kończący się na `/feed` albo `/rss`. Jeśli nie wiesz, czy dana strona go ma, to wpisz jej adres w wyszukiwarce razem ze słowem „RSS” - zwykle ktoś już to sprawdził.

- 7 Kliknij „Execute step” (wykonaj krok) w prawym górnym rogu panelu węzła. Na ekranie zobaczysz listę artykułów w panelu OUTPUT, każdy z autorem, tytułem, linkiem i datą publikacji.

Dodajesz trzeci węzeł: filtr

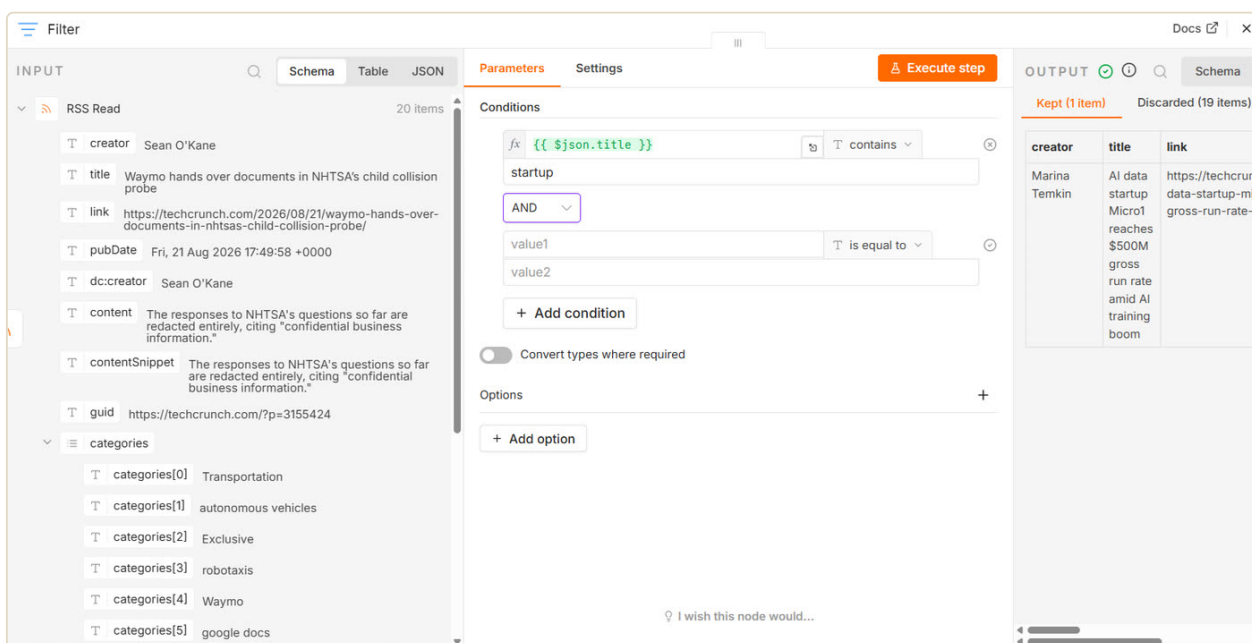
- 8 Kliknij znak plus przy prawej krawędzi węzła „RSS Read”. Wpisz „Filter” i kliknij węzeł „Filter” (filtruj) z listy wyników.

- 9 Po dwukrotnym kliknięciu węzła kliknij „Add condition” (dodaj warunek). n8n pokaże puste pole na wartość do sprawdzenia. W to pole musisz wstawić tytuł artykułu - i tu po raz pierwszy zrobisz coś nowego: po lewej stronie ekranu, w panelu danych wejściowych, znajdź pole `title` i przeciągnij je myszą prosto na puste pole warunku.

NA EKRANIE ZOBACZYSZ

w polu pojawi się tekst w podwójnym nawiasie klamrowym, coś w rodzaju `{{ $json.title }}`.

- 10 Wybierz typ porównania „String” i warunek „contains” (zawiera) z listy. W polu wartości do porównania wpisz słowo, którego szukasz, na przykład „AI” albo „startup” itp.



Rys. 3.3. Gotowy warunek filtra

- 11 Kliknij „Execute step”. Na ekranie zobaczysz: krótszą listę w panelu OUTPUT - tylko te artykuły, w których tytule pojawiło się słowo „AI”.

Zanim dodasz mail: sprawdź, ile masz artykułów

Zajrzyj jeszcze raz do panelu OUTPUT węzła „Filter” i policz elementy, które przez niego przeszły. Jeśli jest ich np. więcej niż pięć, wróć do warunku i wpisz jakieś inne, węższe słowo. Za chwilę zobaczysz, dlaczego to ma znaczenie.

Dodajesz czwarty węzeł: wysyłka maila

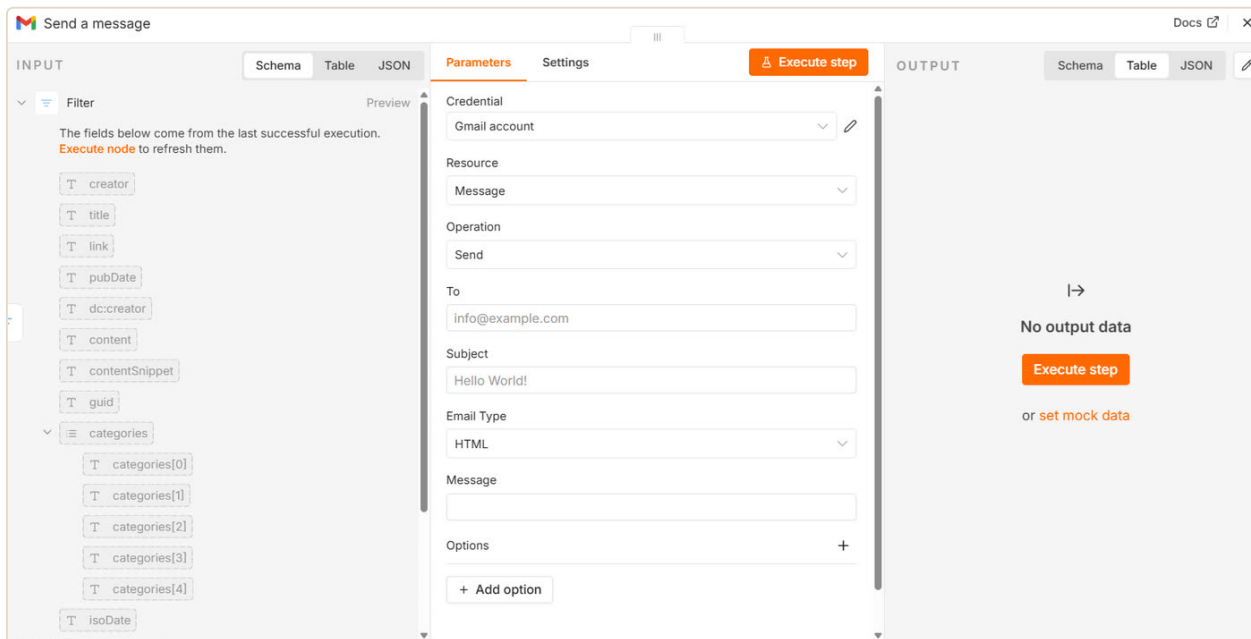
- 12 Kliknij znak plus przy prawej krawędzi węzła „Filter”. Wpisz „Gmail” i wybierz węzeł „Gmail” z listy wyników.

Tu też pojawia się coś, co powtórzy się przy większości węzłów łączących się z usługami zewnętrznymi. Jeden węzeł Gmail potrafi robić wiele różnych rzeczy: wysyłać wiadomości, czytać

je, tworzyć wersje robocze, zarządzać etykietami itp. Zanim pokaże ci właściwe pola, musi wiedzieć, którą z tych rzeczy wybierasz.

Na liście „MESSAGE ACTIONS” wybierz akcję „Send a message” (wyślij wiadomość). Na ekranie zobaczysz: panel węzła.

W panelu zobaczysz pola „Credential”, „Resource”, „Operation”, „To”, „Subject”, „Email Type” i „Message” oraz sekcję „Options”.

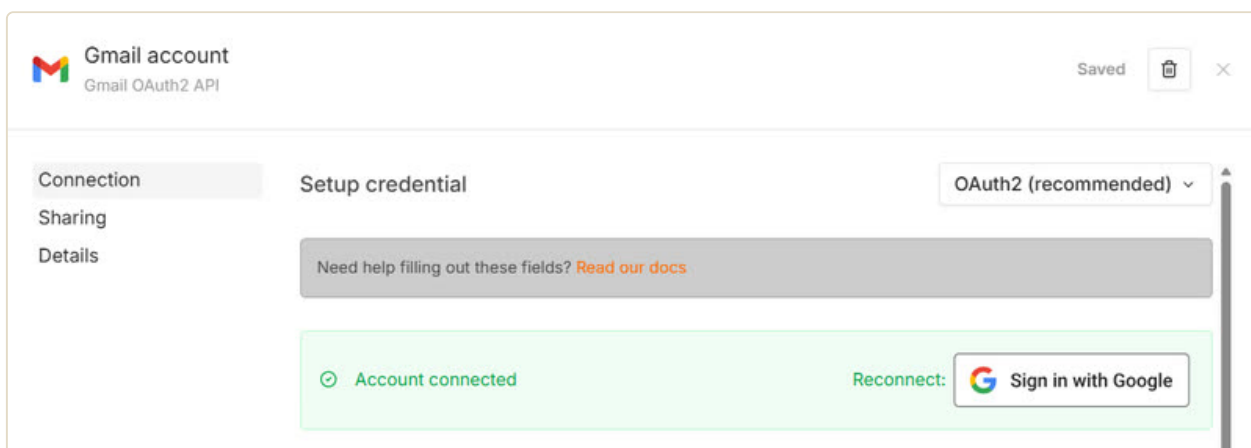


Rys. 3.4. Zasób i operacja. Bez tego wyboru węzeł nie pokaże pól wiadomości

- 13 Węzeł prawdopodobnie zgłosi teraz, że brakuje mu danych dostępowych, czyli zgody na wysyłanie maili w twoim imieniu. Kliknij „Create new credential” (utwórz nowe dane dostępowe), a potem „Sign in with Google” (zaloguj się przez Google) i przejdź logowanie w oknie, które się otworzy.

NA EKRANIE ZOBACZYSZ

powrót do n8n i zielony komunikat potwierdzający, że konto zostało połączone.



Rys. 3.5. Logowanie przez Google przy pierwszym użyciu węzła Gmail

! UWAGA

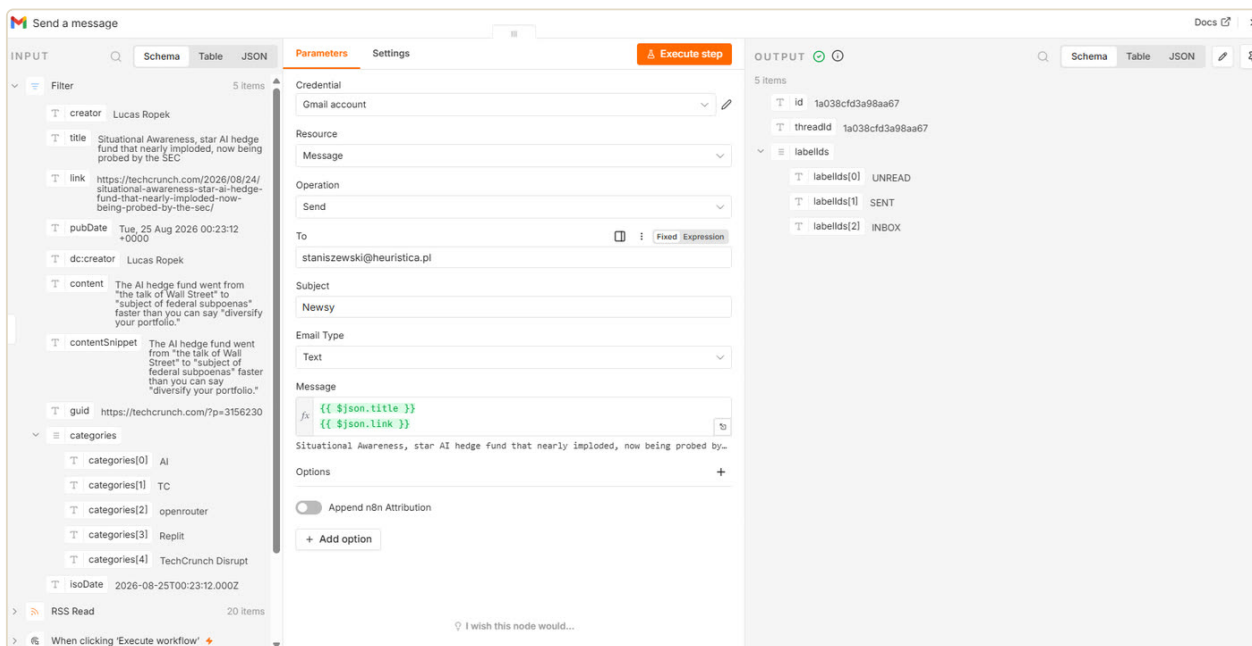
n8n nigdy nie prosi cię o wpisanie hasła do Gmaila wprost. Logujesz się na oficjalnym ekranie Google, a n8n dostaje wyłącznie zgodę na wysyłanie wiadomości w twoim imieniu.

Teraz wypełnij pola wiadomości:

- „To” (do): twój własny adres e-mail.
- „Subject” (temat): na przykład „Nowe artykuły o AI”.
- „Email Type” (typ wiadomości): zostaw „Text” (tekst). Druga opcja, „HTML”, przydaje się, gdy chcesz formatować treść, ale teraz tylko by przeszkadzała.
- „Message” (treść): przeciągnij tu pole „title” z lewego panelu danych wejściowych, tak samo jak przy filtrze w kroku 9. Kliknij za wstawionym tekstem, wciśnij Enter, żeby zrobić nową linię, i dopiero teraz przeciągnij pole link. Bez tego oba pola skleją się w jeden nieczytelny ciąg.

NA EKRANIE ZOBACZYSZ

w polu treści dwa fragmenty w podwójnych nawiasach klamrowych, jeden pod drugim, a pod polem podgląd z prawdziwym tytułem i linkiem pierwszego artykułu.



Rys. 3.6. Gotowa wiadomość przed wysłaniem

Zanim klikniesz, dwie rzeczy, o których warto wiedzieć.

Po pierwsze, „Execute step” w tym węźle naprawdę wysyła wiadomość. Wszystko więc, co tu ustawileś, za chwilę pojedzie na twoją skrzynkę.

Po drugie, nie dostaniesz jednego maila. Dostaniesz tyle maili, ile artykułów przeszło przez filtr. Jeśli filtr przepuścił cztery artykuły, dostaniesz cztery osobne wiadomości, każdą z jednym

tytułem. W tym przypadku to nie jest błąd i nie ma co go naprawiać. To jedna z zasad działania n8n i jedna z rzeczy, które warto dobrze rozumieć.

! UWAGA

Węzeł wykonuje się raz dla każdego elementu, który do niego przyszedł.

Pamiętasz liczby pod węzłami z rozdziału 2? Jeśli pod węzłem „Filter” stoi „4 items”, to węzeł Gmail wykona się cztery razy: raz z pierwszym artykułem, raz z drugim i tak dalej. Za każdym razem pole title będzie znaczyło coś innego, bo za każdym razem odnosi się do innego elementu.

Ta sama zasada rządzi wszystkim, co zbudujesz dalej. Kiedy w rozdziale 5 podłączysz model językowy do listy np. dziesięciu artykułów, model dostanie dziesięć osobnych pytań, a nie jedno pytanie o dziesięć artykułów. Zapłacisz też za to dziesięć razy.

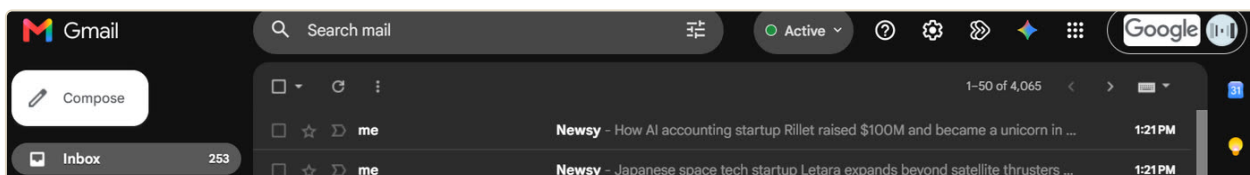
Da się to zmienić, czyli zebrać wszystkie elementy w jeden przed wysyłką. Wrócimy do tego w rozdziale 6, gdzie będzie do tego dobry powód. Na razie zostawiamy jak jest, bo dwa czy cztery maile to i tak niska cena za zrozumienie i zapamiętanie tej zasady :-)

14 Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

zieloną obwódkę wokół węzła Gmail i liczbę pod nim odpowiadającą liczbie wysłanych wiadomości.

Sprawdź, czy działa: otwórz swoją skrzynkę pocztową. Powinny czekać na ciebie nowe wiadomości o temacie, który wpisałeś - np. „Nowe artykuły o AI”, tyle sztuk, ile pokazała liczba pod węzłem Gmail. W każdej jeden tytuł i jeden link.



Rys. 3.7. Maile, które właśnie wysłał twój własny przepływ

★ WSKAZÓWKA

Na dole każdej wiadomości może pojawić się krótka informacja, że została wysłana automatycznie przez n8n. To domyślne ustawienie węzła. Wyłączysz je w panelu węzła Gmail, w sekcji z opcjami dodatkowymi. Czyli: „Add option” → „Append n8n Attribution” i wyłącz tę opcję.

A teraz czas na prawdziwą automatyzację

Zbudowałeś przepływ, który zbiera artykuły, odsiewa nieciekawe i przysyła resztę na twoją skrzynkę. Jest tylko jeden problem: żeby to się stało, musisz usiąść przy komputerze, otworzyć n8n i kliknąć. A to jeszcze nie jest automatyzacja.

Pamiętasz z rozdziału 2, że pierwszy węzeł każdego przepływu decyduje, kiedy przepływ się uruchomi? Do tej pory pracował tam wyzwalacz ręczny, który czeka na twoje kliknięcie. Teraz dołożysz drugi, który nie czeka na nikogo.

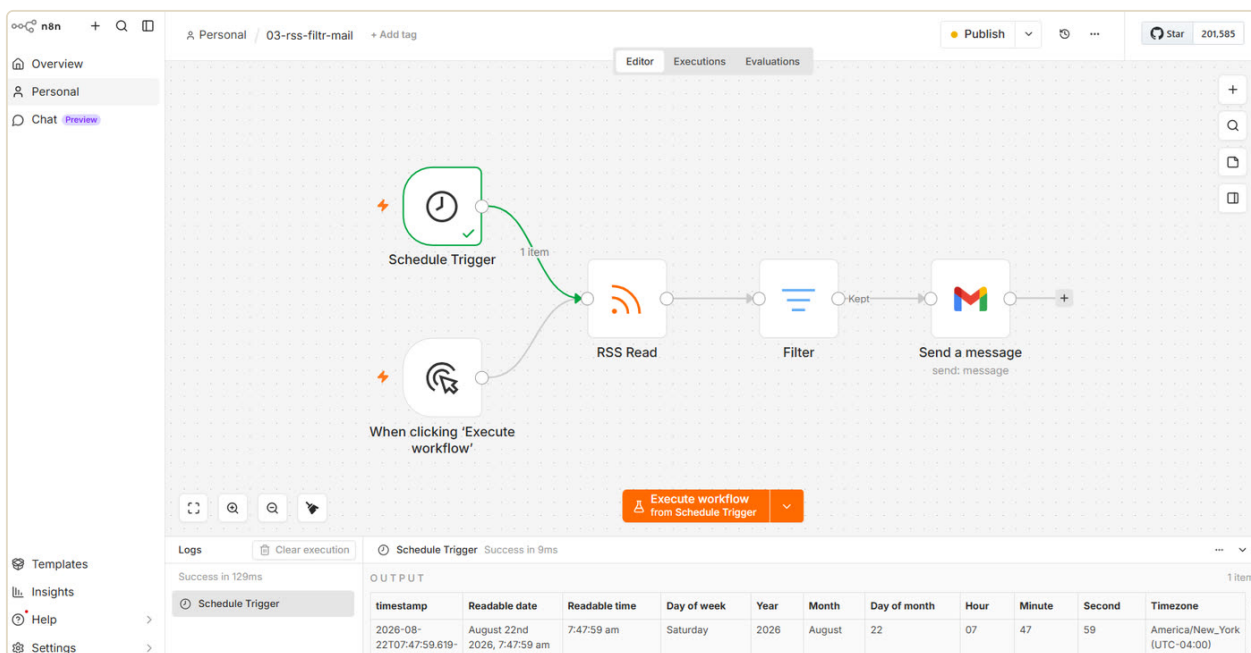
- 15 Kliknij zatem przycisk z plusem w prawym górnym rogu obszaru roboczego. W polu wyszukiwania wpisz „Schedule” i wybierz węzeł „Schedule Trigger” (wyzwalacz harmonogramu). Na ekranie zobaczysz: nowy węzeł stojący samotnie obok twojego przepływu, jeszcze z niczym niepołączony.
- 16 W panelu węzła znajdź listę „Trigger Interval” (odstęp między uruchomieniami) i wybierz „Days” (dni). Pojawią się trzy pola. W „Days Between Triggers” (co ile dni) zostaw 1.
- 17 W „Trigger at Hour” (o której godzinie) wybierz 7am. W „Trigger at Minute” (w której minucie) wpisz 0 i zamknij panel.

The screenshot shows the 'Schedule Trigger' node configuration. The left panel has tabs for 'Parameters', 'Settings', and 'Execute step'. The 'Settings' tab is active, showing a 'Trigger Rules' section with a plus icon. Under 'Trigger Interval 1', the 'Trigger Interval' is set to 'Days', 'Days Between Triggers' is 1, 'Must be in range 1-31' is checked, 'Trigger at Hour' is 7am, and 'Trigger at Minute' is 0. There is an 'Add Rule' button at the bottom. The right panel shows the 'OUTPUT' section with a table containing one item.

timestamp	Readable date	Readable time	Day of week	Year	Month	Day of month	Hour	Minute	Second	Timezone
2026-08-22T07:47:59.619-04:00	August 22nd 2026, 7:47:59 am	7:47:59 am	Saturday	2026	August	22	07	47	59	America/New_York (UTC-04:00)

Rys. 3.8. Codziennie o siódmej rano

- 18 Przeciągnij myszką od kropki na prawej krawędzi węzła „Schedule Trigger” do lewej krawędzi węzła „RSS Read” i puść. Na ekranie zobaczysz: dwie strzałki wchodzące do węzła „RSS Read”, jedną od wyzwalacza ręcznego i jedną od harmonogramu. Przepływ, jak widać, może mieć więcej niż jeden wyzwalacz. Ręczny zostaje, bo dzięki niemu nadal uruchomisz wszystko na żądanie, kiedy chcesz coś sprawdzić. Harmonogram robi to samo, tylko sam z siebie, o siódmej rano.



Rys. 3.9. Dwa wejścia do tego samego przepływu

- 19 Zapisz przepływ przyciskiem „Publish” (prawy górny róg) - teraz twój przepływ jest już aktywny.

! UWAGA

Nieaktywny przepływ nie uruchamia się według harmonogramu, nawet jeśli harmonogram jest w nim ustawiony. To najczęstsze rozczarowanie początkujących: wszystko poprawnie skonfigurowane, aktywacja zapomniana, rano nic nie przychodzi. Druga strona tej samej zasady: przycisk „Execute workflow” działa niezależnie od aktywacji. Uruchomienie ręczne zadziała nawet przy nieaktywnym przepływie, a aktywny przepływ chodzi po godzinach, kiedy ty nie pilnujesz n8n.

Gdzie zobaczyć, czy zadziało

Uruchomienia z harmonogramu nie zostawiają śladu na obszarze roboczym. Nie ma tam zielonych obwódek, bo w chwili wykonania nikt nie patrzył na ekran. Zapis znajdziesz w zakładce „Executions” (uruchomienia) na górze widoku przepływu: lista wszystkich wykonań z datą, godziną, czasem trwania i informacją, czy zakończyły się sukcesem. Zapamiętaj tę zakładkę - to główne miejsce, w którym sprawdza się, dlaczego coś przestało działać, i będziesz do niej wracać w każdym kolejnym rozdziale.

Co może pójść nie tak

- W panelu węzła Gmail nie ma pól „To” ani „Subject”. Sprawdź listy „Resource” i „Operation” na górze panelu. Muszą być ustawione na „Message” i „Send”. Dopóki nie wybierzesz operacji, węzeł nie wie, jakie pola ci pokazać.
- Dostałeś kilkanaście maili zamiast kilku. Filtr przepuścił za dużo artykułów. Wróć do warunku i wpisz węższe słowo, potem uruchom przepływ jeszcze raz.
- Tytuł i link skleły się w jedną linię. W polu „Message” brakuje przejścia do nowej linii między przeciągniętymi polami. Kliknij między nimi i wciśnij Enter.
- Węzeł Gmail zgłasza błąd „insufficient permission” (niewystarczające uprawnienia). Połączenie z kontem Google nie objęło wszystkich potrzebnych zgód. Usuń dane dostępne i połącz konto jeszcze raz, zaznaczając wszystkie prośzone uprawnienia.
- Mail nie przyszedł, a węzeł Gmail jest zielony. Sprawdź folder spam. Pierwsza wiadomość z nowo połączonego konta czasem tam trafia.

Wariacje na temat

Odłącz na chwilę węzeł Gmail, klikając strzałkę między nim a filtrem i usuwając połączenie. Zmień słowo w warunku filtra na coś, co na pewno pojawi się w wielu tytułach, na przykład „the”, i uruchom przepływ. Zobacz, jak urosła liczba pod węzłem „Filter”. Potem przywróć poprzednie słowo i połączenie.

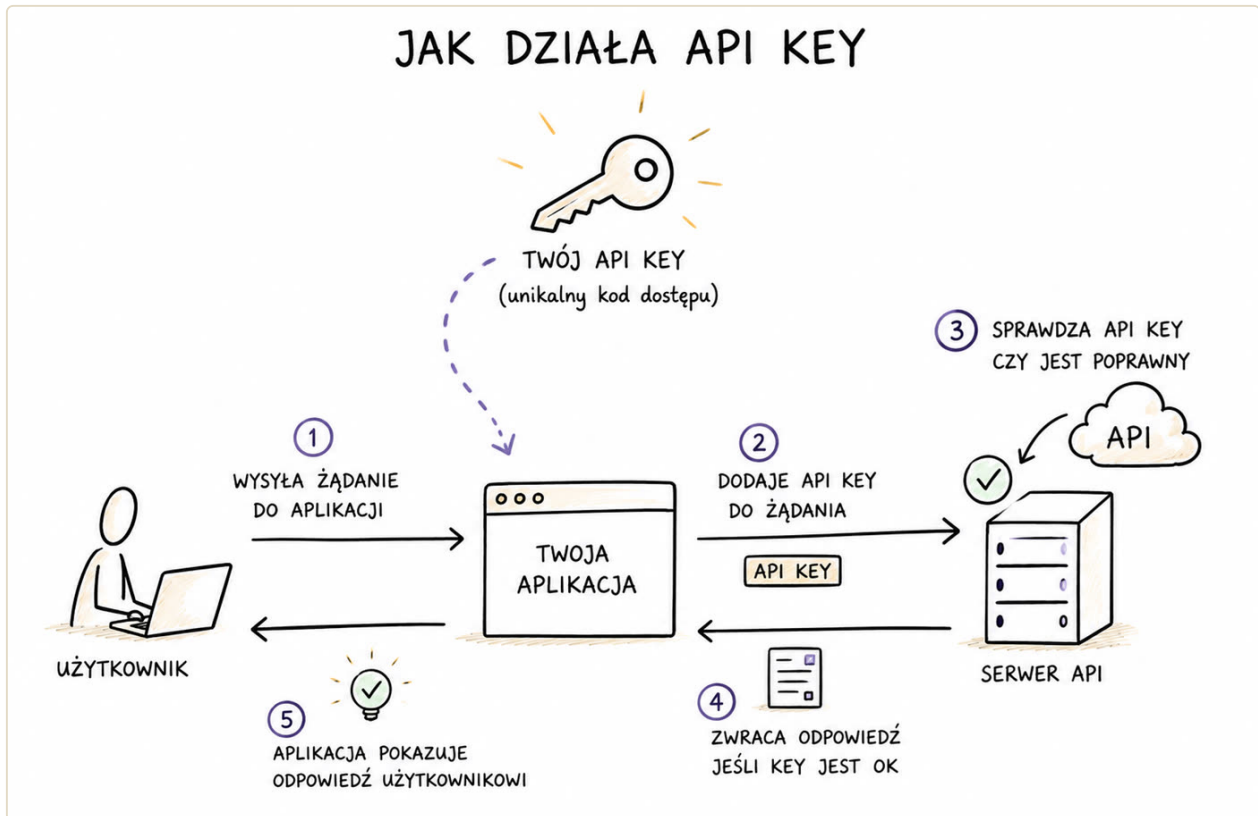
Rozdział 4. Podłączamy AI

Do tej pory każdy węzeł w twoim przepływie był doskonale przewidywalny: przy tych samych danych na wejściu robił zawsze dokładnie to samo. Filtr przepuszczał to, co pasowało do warunku, i nic poza tym. Teraz dołożysz pierwszy węzeł, który przy tym samym pytaniu potrafi odpowiedzieć dwa razy inaczej. **Wszystko, co dziś sprzedaje się pod słowem „AI”, sprowadza się w przepływie do jednej zmiany: jeden z kroków przestaje wykonywać regułę i zaczyna podejmować decyzję.**

Pytanie na start: co to jest klucz API i skąd go wziąć?

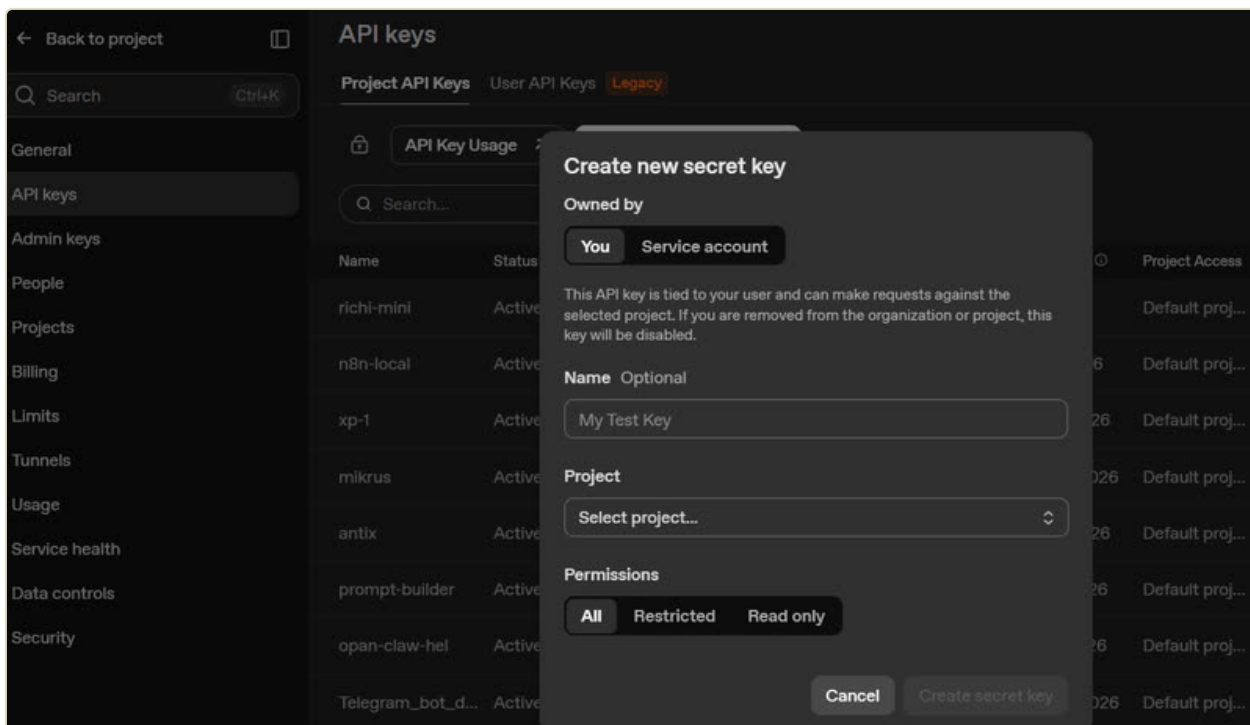
Zanim węzeł sztucznej inteligencji - np. wybrany model ChatGPT od OpenAI - cokolwiek zrobi, będziesz potrzebować klucza API - czyli ciągu znaków, który mówi serwerom OpenAI, że to naprawdę ty wysyłasz zapytanie i że wolno cię obciążyć za jego koszt. To w zasadzie ten sam mechanizm jak dane dostępowe do Gmaila z rozdziału 3, tyle że tym razem nie logujesz się przez Google, tylko wklejasz gotowy kod.

Samo API (Application Programming Interface) to zaś zestaw narzędzi, protokołów i definicji, które umożliwiają różnym aplikacjom komunikację ze sobą. API działa w tym przypadku jak pośrednik, który pozwala jednej aplikacji korzystać z funkcjonalności innej bez konieczności poznawania jej wewnętrznej struktury.



Rys. 4.0. Jak działa klucz API

- 1 Otwórz w nowej karcie przeglądarki stronę platform.openai.com i zaloguj się na swoje konto (albo załóż nowe, jeśli jeszcze go nie masz). Możesz też wejść bezpośrednio na stronę: <https://platform.openai.com/organization/api-keys>
- 2 Z menu po lewej stronie wybierz „API keys” (klucze API), a potem kliknij „Create new secret key” (utwórz nowy tajny klucz).



Rys. 4.1. Tworzenie nowego klucza API na koncie OpenAI

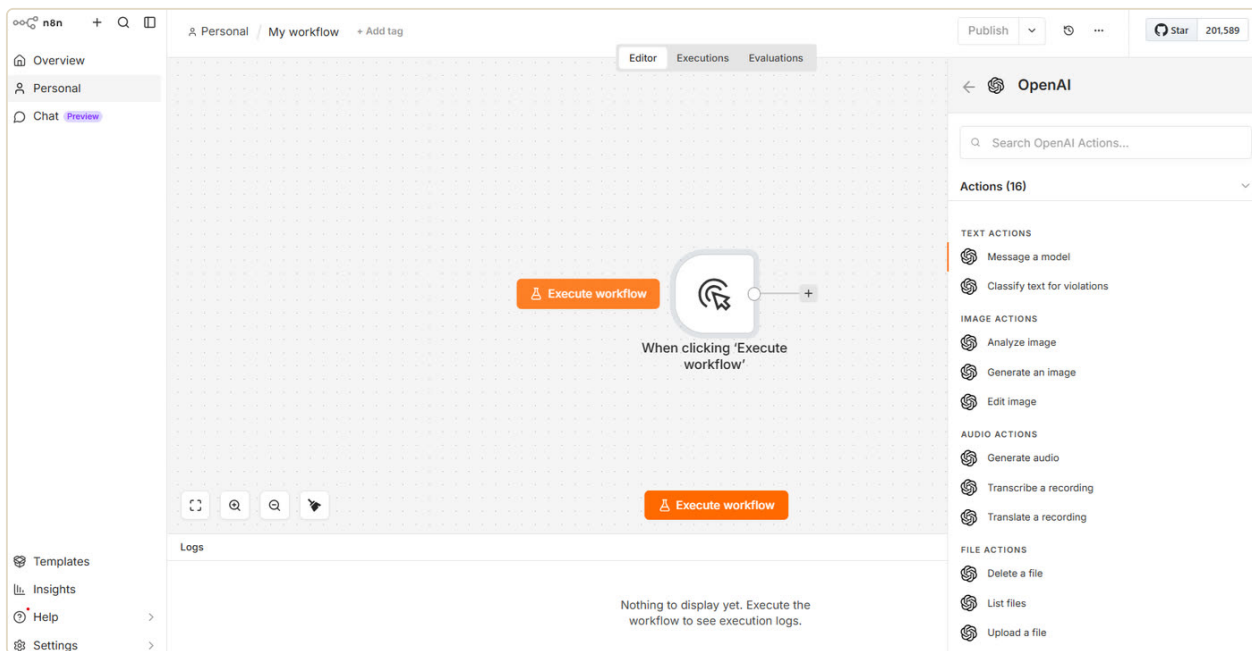
- 3 Nadaj kluczowi nazwę, na przykład „n8n”, i potwierdź utworzenie. Na ekranie zobaczysz: swój nowy klucz w całości - **to jedyny moment, w którym zobaczysz go w pełnej postaci**. Skopiuj go od razu przyciskiem obok i zapisz tymczasowo w bezpiecznym miejscu, na przykład w menedżerze haseł.

! UWAGA

Zanim pójdziesz dalej, sprawdź w zakładce „Billing” (rozliczenia) swojego konta OpenAI, czy masz podpętą metodę płatności i choć niewielki limit wydatków. Nowe konto bez włączonego rozliczenia zwraca dokładnie ten sam błąd, co zły klucz - a wygląda identycznie, jakbyś zrobił coś źle, mimo że klucz jest poprawny.

Dodajesz węzeł OpenAI do przepływu

- 4 Wróć do n8n i utwórz nowy, pusty przepływ - tak jak w rozdziale 3. Dodaj węzeł „Manual trigger” (uruchom ręcznie), już znany z poprzednich rozdziałów.
- 5 Kliknij znak plus przy prawej krawędzi wyzwalacza. W polu wyszukiwania wpisz „OpenAI”. Zamiast pojedynczego wyniku zobaczysz listę akcji pogrupowanych w sekcje: TEXT ACTIONS, IMAGE ACTIONS, AUDIO ACTIONS i inne. W sekcji TEXT ACTIONS kliknij „Message a model” (wyślij wiadomość do modelu). Na ekranie zobaczysz: nowy węzeł na obszarze roboczym i otwarty panel z polami „Credential to connect with”, „Model” i „Messages”.

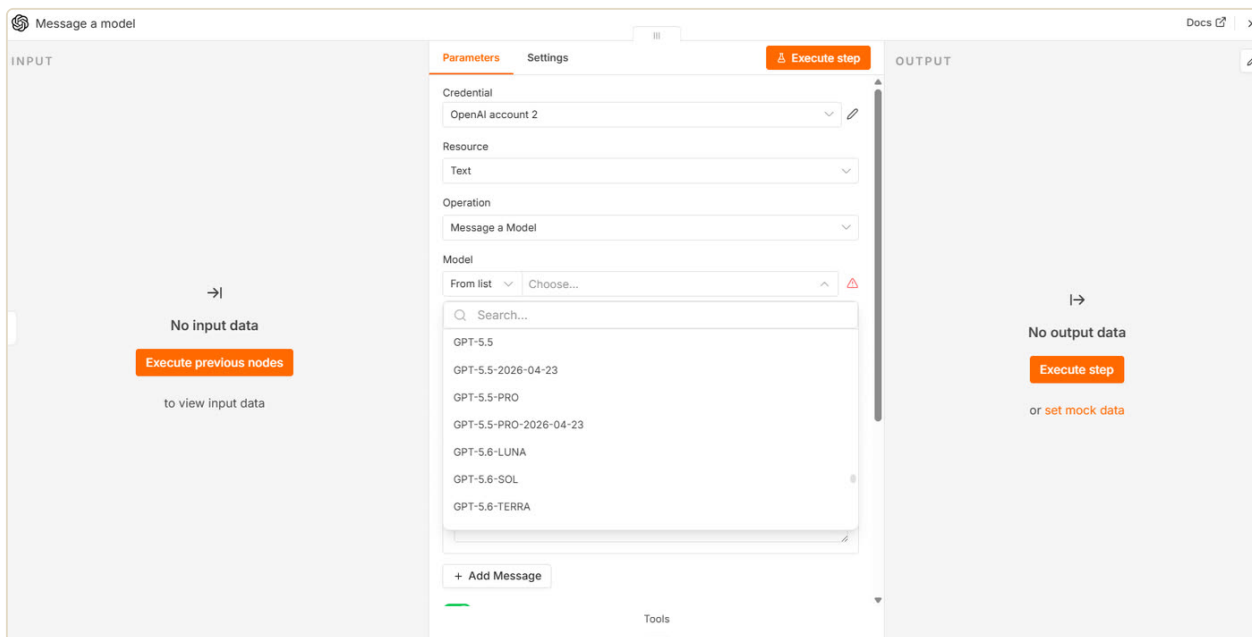


Rys. 4.2. Jeden węzeł i wiele rzeczy do zrobienia. Wybierasz je od razu przy dodawaniu

- 6 Otwórz węzeł dwukrotnym kliknięciem. Na samej górze panelu jest pole „Credential” (dane dostępowe). Rozwiń je i wybierz „Create new credential” (utwórz nowe dane dostępowe).
- 7 Otworzy się osobne okno z jednym polem: „API Key”. Wklej w nie klucz skopiowany w kroku 3 (uważaj, żeby nie złapać spacji na początku ani na końcu). Kliknij „Save” (zapisz).
- 8 Na ekranie zobaczysz: zielony komunikat potwierdzający, że połączenie działa. Zamknij okno danych dostępowych. W polu „Credential” widnieje teraz nazwa twojego nowego połączenia, na przykład „OpenAI account”.

Klucz wklejasz w n8n tylko raz. Od tej chwili będzie się pojawiał na liście w każdym nowym węźle OpenAI, który dodasz w dowolnym przepływie. To samo dotyczy połączenia z Gmailem z rozdziału 3. Wszystkie zapisane dane dostępowe znajdziesz w menu bocznym n8n, w sekcji „Credentials”, i tam też je usuniesz albo podmienisz, gdy klucz przestanie działać.

- 9 W polu „Model” wybierz najtańszy dostępny model z rodziny GPT-4o - w chwili pisania jest to gpt-4o-mini. Lista modeli pobiera się na żywo z twojego konta OpenAI, więc może wyglądać inaczej niż na zrzucie.



Rys. 4.3. Lista dostępnych modeli

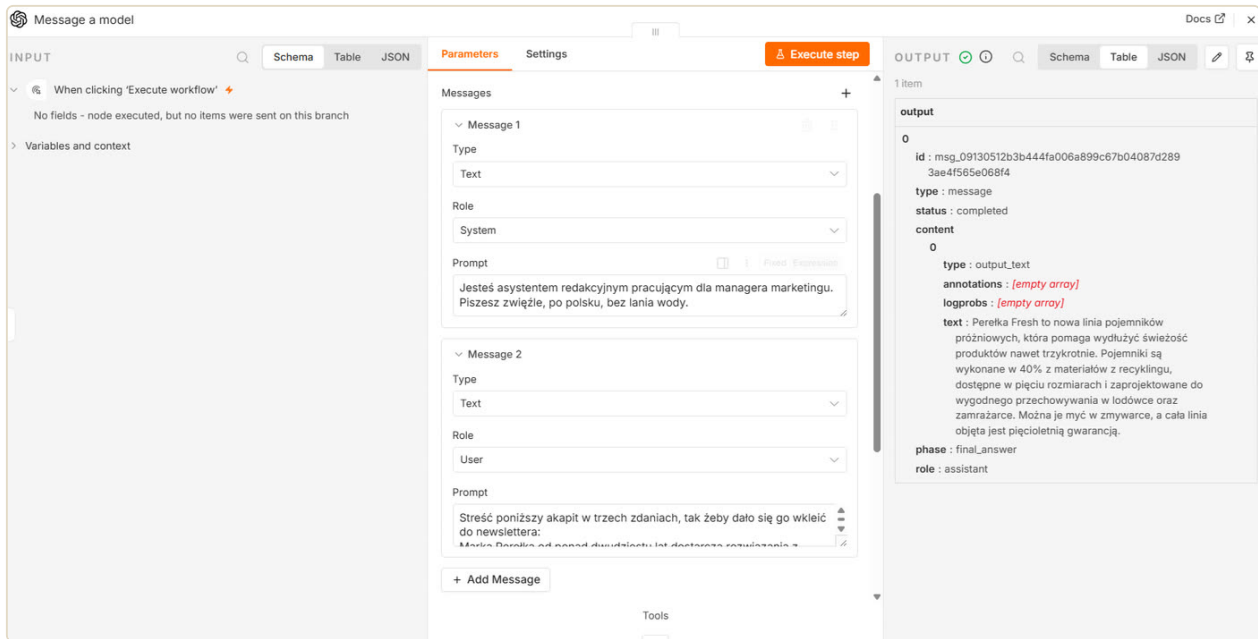
Piszesz pierwsze polecenie

Dobre polecenie dla modelu działa trochę tak jak dobre zlecenie dla nowego stażysty - musi powiedzieć mu, kim ma być, co konkretnie ma zrobić i w jakiej formie oddać wynik.

- 10 W sekcji „Messages” (wiadomości) ustaw na liście „Type” wartość „text”. Następnie kliknij „Add Message” (dodaj wiadomość) i ustaw „Role” (rola) na „System”. W polu tekstowym opisz, kim ma być model i jak się zachowywać, na przykład: „Jesteś asystentem redakcyjnym pracującym dla managera marketingu. Piszesz zwięźle, po polsku, bez lania wody.”
- 11 Dodaj drugą wiadomość, tym razem z rolą „User” (użytkownik). To jest właściwe zadanie - konkretny tekst do przetworzenia, na przykład: „Streść poniższy akapit w trzech zdaniach, tak żeby dało się go wkleić do newslettera:” i wklej dowolny fragment tekstu, choćby opis produktu ze swojej strony.
- 12 Kliknij „Execute step”. Na ekranie zobaczysz: w panelu OUTPUT rozwinięte drzewko z kilkoma pozycjami, między innymi id, type, status i content. W polu text pojawi się twoje streszczenie.

Model nie odpowiada samym tekstem. Odpowiada całą kartą informacji o odpowiedzi: kiedy powstała, jaki ma numer, czy zakończyła się poprawnie i dopiero na końcu, co właściwie napisał.

Z tej karty interesuje cię w tej chwili jedno pole, text. Reszta przyda się później. Zobaczysz też dwa pola opisane czerwonym [empty array]. To nie jest błąd. To znaczy „tu nic nie ma”, bo model nie miał czego dopisać.



Rys. 4.4. Polecenie systemowe i polecenie użytkownika oraz odpowiedź modelu

Zmień teraz jedno słowo w poleceniu systemowym - na przykład każ modelowi pisać żartobliwie zamiast rzeczowo - i uruchom węzeł jeszcze raz. Teraz przy tej samej treści wejściowej pojawi się inny ton na wyjściu.

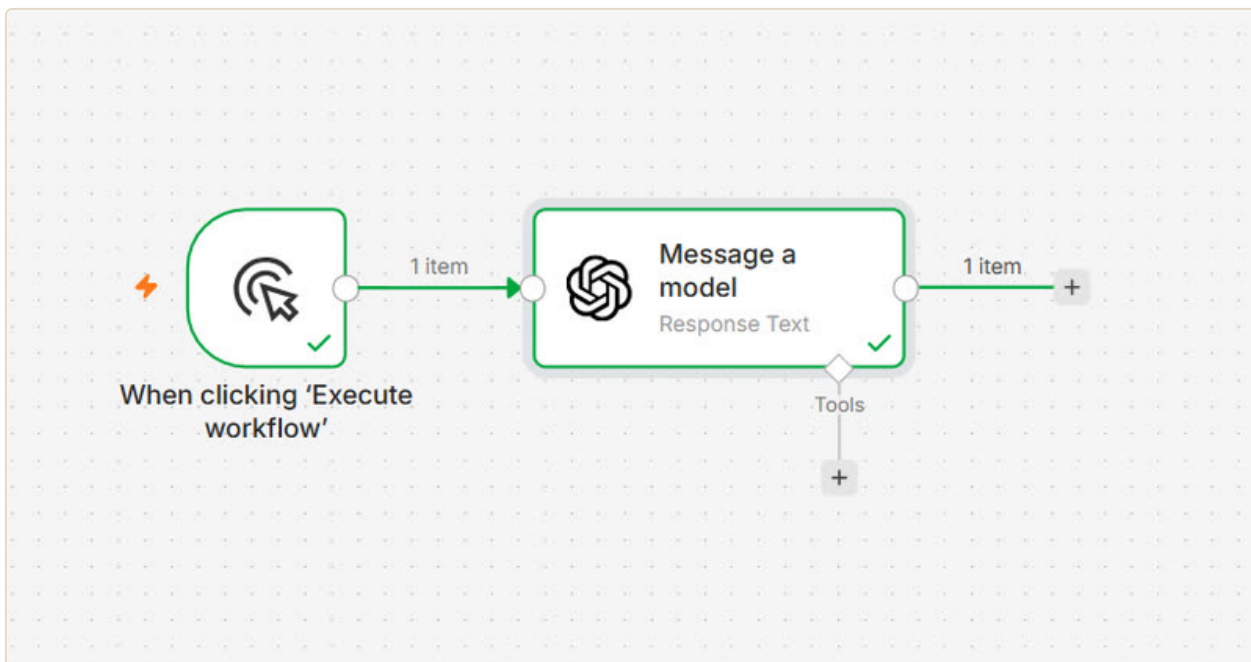
? DLA CIEKAWYCH

Reguła „kim jesteś, co masz zrobić, w jakiej formie oddać” nie jest wynalazkiem promptowania. David Ogilvy, jeden z twórców nowoczesnej reklamy, powtarzał zespołom kreatywnym, że dobry brief ogranicza, a ograniczenie rodzi lepszą pracę niż pełna dowolność. To, co dziś nazywamy „dobrym promptem”, działało tak samo na copywriterów w latach sześćdziesiątych, jak dziś działa na model językowy.

Model dostaje narzędzie

Zerknij jeszcze raz na swój przepływ. Pod węzłem „Message a model” wisi jakaś dziwna wypustka podpisana „Tools” (narzędzia) ze znakiem plus. Wszystkie połączenia, które robiłeś do tej pory, biegły z lewej na prawo. To jedno biegnie z dołu do góry i znaczy coś zupełnie innego.

Połączenie z lewej strony wnosi dane do przetworzenia. Połączenie od dołu nie wnosi danych. Wyposaża węzeł. Mówi modelowi: „masz do dyspozycji jeszcze to coś”.



Rys. 4.5. Tu można dodać modelowi narzędzia

- 13 Kliknij plus przy „Tools”. W polu wyszukiwania wpisz „Wikipedia” i wybierz węzeł „Wikipedia”.

NA EKRANIE ZOBACZYSZ

nowy, mniejszy węzeł doczepiony od spodu do węzła „Message a model”, połączony z nim krótką linią. Nie ma w nim nic do ustawiania („This node does not have any parameters”). Zamknij zatem panel.

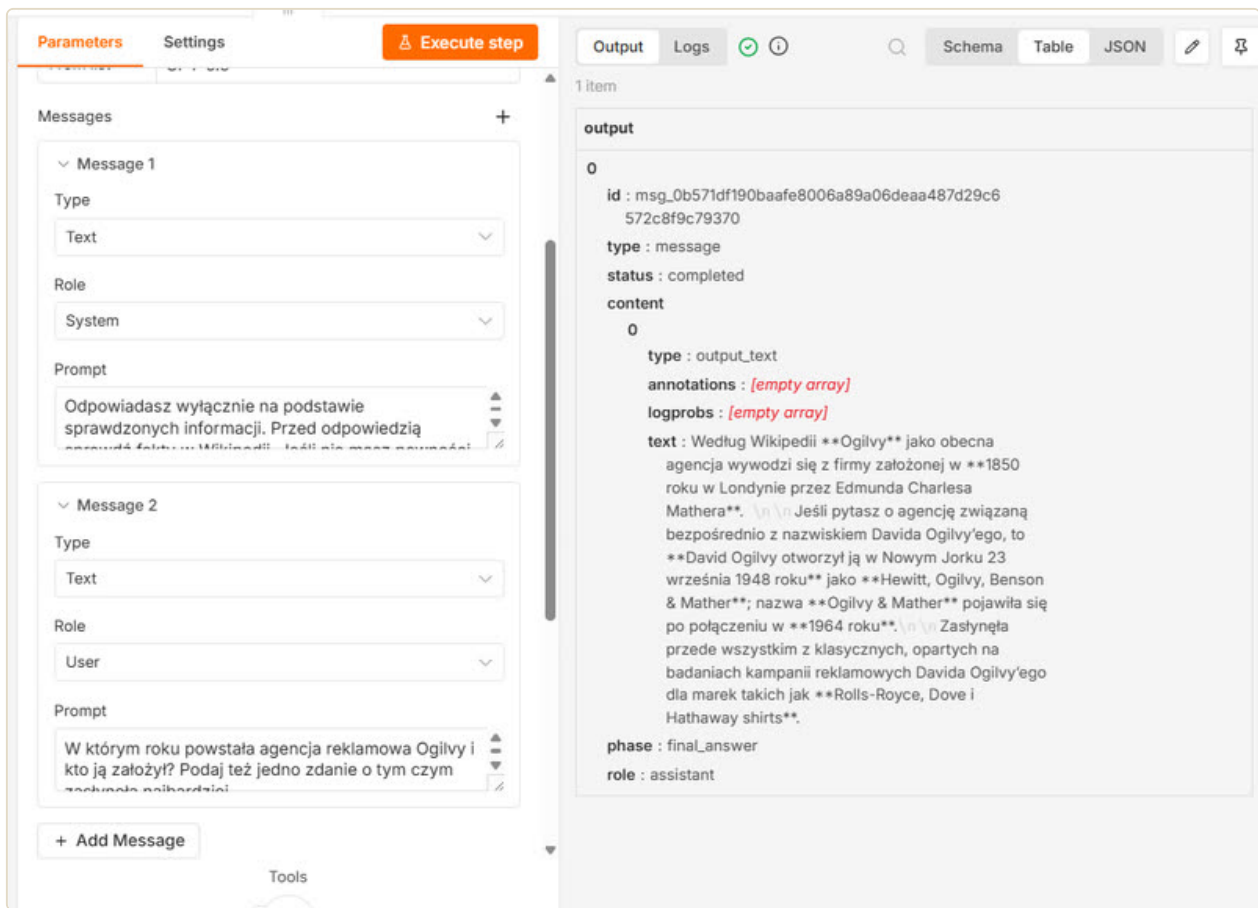
- 14 Wróć do węzła „Message a model” i podmień obie wiadomości. W wiadomości z rolą „System” wpisz:

„Odpowiadasz wyłącznie na podstawie sprawdzonych informacji. Przed odpowiedzią sprawdź fakty w Wikipedii. Jeśli nie masz pewności, napisz NIE WIEM.”

W wiadomości z rolą „User” wpisz konkretne pytanie o fakty, na przykład:

„W którym roku powstała agencja reklamowa Ogilvy i kto ją założył? Podaj też jedno zdanie o tym, czym zasłynęła najbardziej.”

- 15 Kliknij „Execute step” i przyjrzyj się panelowi OUTPUT. Oprócz samej odpowiedzi zobaczysz w nim ślad po tym, że model sięgnął po narzędzie: wpis z zapytaniem, które sam ułożył, i z tym, co Wikipedia mu odesłała.



Rys. 4.6. Dwa rodzaje połączeń. Jedno wnosi dane, drugie możliwości

! UWAGA

Model, w zależności od stopnia swojego zaawansowania, sam też często decyduje, czy sięgnąć po narzędzie. Nie wykonuje twojej instrukcji tak, jak węzeł wykonuje ustawienie. Może użyć Wikipedii, może uznać, że wie wystarczająco dużo, i odpowiedzieć z pamięci. Uruchom to samo polecenie trzy razy pod rząd, a możesz dostać trzy różne zachowania. To nie jest wada narzędzia ani twój błąd w konfiguracji. To podstawowa różnica między węzłem a modelem: węzeł robi zawsze to samo, model podejmuje decyzję. Warto to zapamiętać, bo z tej jednej właściwości bierze się większość rozczarowań systemami opartymi na AI - i dlatego wszystko, co w tej książce zależy od oceny modelu, trzeba od czasu do czasu sprawdzić ręcznie.

? DLA CIEKAWYCH

Wikipedia jest tu tylko najprostszym przykładem, bo nie wymaga żadnej rejestracji. W tym samym miejscu podłącza się wyszukiwarkę Google, twoją bazę produktów, kalendarz albo firmowe repozytorium dokumentów. Zasada nie zmienia się ani trochę: model nie przestaje być modelem, tylko przestaje być zamknięty w tym, co zapamiętał podczas nauki.

Ile to kosztuje

Gmail z rozdziału 3 był darmowy - wysyłanie maili w twoim imieniu nic nie kosztowało. Węzeł OpenAI jest inny: każde „Execute step” to prawdziwe zapytanie do serwerów OpenAI, a ty płacisz za nie według długości zarówno polecenia, jak i odpowiedzi, liczonej w małych jednostkach tekstu zwanych tokenami. Tańsze modele, takie jak gpt-4o-mini, kosztują ułamki centa za pojedyncze uruchomienie.

Sto takich uruchomień, jak to, które właśnie wykonałeś, kosztuje przy najtańszym modelu około dwóch centów, czyli mniej niż dziesięć groszy. Tysiąc kosztuje mniej niż złotówkę. Nie znaczy to, że koszt jest nieistotny. Rachunek zaczyna być jednak widoczny dopiero wtedy, gdy przepływ działa sam, codziennie, na kilkudziesięciu elementach naraz. Do tego wrócimy w rozdziale 5.



WSKAZÓWKA

Podczas testowania i poprawiania polecenia zostań przy tańszym modelu - różnica w jakości rzadko usprawiedliwia różnicę w cenie na tym etapie. Przełącz się na mocniejszy model dopiero wtedy, gdy wiesz już dokładnie, jakiego wyniku oczekujesz.

Co może pójść nie tak

- Węzeł zgłasza błąd „401 Unauthorized: Incorrect API key provided”. Klucz API jest błędny albo niekompletny - sprawdź, czy skopiowałeś go w całości, bez spacji na początku lub końcu.
- Ten sam błąd 401 pojawia się mimo poprawnego klucza. Na koncie OpenAI nie ma podpisanej metody płatności albo ustawionego limitu wydatków - sprawdź zakładkę „Billing”.
- Węzeł zgłasza błąd o przekroczonym limicie zapytań albo braku środków (komunikat zawiera „rate limit” albo „insufficient_quota”). Poczekać chwilę i spróbuj ponownie, a jeśli się powtarza, sprawdź limit wydatków na koncie OpenAI.
- Zamknąłeś stronę OpenAI, zanim skopiowałeś klucz, i teraz nie możesz go nigdzie znaleźć. Klucze API są pokazywane w pełnej postaci tylko raz - wróć do „API keys” i utwórz nowy.
- Węzeł zgłasza błąd, że wybrany model nie istnieje. Nazwa modelu w rozwijanym polu mogła się zmienić od czasu napisania tego rozdziału - wybierz najbliższy odpowiednik z rodziny GPT.

Wariacje na temat

Zmień model na inny z tej samej rodziny i zapuść to samo polecenie jeszcze raz. Porównaj różnicę w jakości odpowiedzi.

Przepisz polecenie systemowe tak, żeby zmienić formę wyniku - punktory zamiast zdań, inny język, krótszy albo dłuższy tekst. Ta sama treść wejściowa da bardzo różne wyniki, w zależności od tego, jak dokładnie opiszesz oczekiwaną formę.

Model językowy nie ma pojęcia o twojej marce, dopóki mu o niej nie powiesz - a to jest dokładnie ta sama zasada, która rządziła każdym dobrym briefem, zanim ktokolwiek wymyślił słowo „prompt”.

Rozdział 5. Monitor konkurencji

Monitoring konkurencji to jedna z tych czynności, które każdy strateg deklaruje, a mało kto naprawdę wykonuje. Nie z lenistwa. Po prostu czytanie dwudziestu newsletterów i sprawdzanie ośmiu stron codziennie przegrywa z każdym pilniejszym zadaniem, a pilniejsze zadanie jest zawsze. **Automatyzacja nie rozwiązuje tu problemu braku czasu, tylko problem braku rytmu.** Maszyna sprawdza codziennie, bo nie ma nic pilniejszego do zrobienia.

Zbudujesz teraz pierwszy z trzech systemów tego poradnika. Wszystkie węzły, które go tworzą, znasz już z rozdziałów 3 i 4 - poza dwoma, i to właśnie one zamieniają listę artykułów w coś, co samo decyduje, co jest warte twojej uwagi.

Co zbudujemy

System, który każdego ranka sam przeszuka doniesienia prasowe o wybranej marce, dla każdego z nich zapyta model, czy jest istotne, i przyśle ci mailem tylko te, które przeszły przez to sito.

Czas i koszt

Budowa zajmie około dwudziestu minut. Dwa nowe węzły wymagają więcej klikania niż wszystko, co robiłeś do tej pory, a jeden z nich trzeba połączyć w sposób, który na pierwszy rzut oka wygląda dziwnie.

Koszt jednego uruchomienia to koszt pojedynczego zapytania z rozdziału 4, pomnożony przez liczbę wiadomości z danego dnia. Przy kilkunastu wiadomościach mówimy o ułamku grosza. Warto jednak zapamiętać sam mechanizm mnożenia, bo to on decyduje o rachunku w każdym większym systemie, który kiedykolwiek zbudujesz.

Czego się nauczysz

- Jak zamienić dowolne zapytanie prasowe w kanał danych, którego nie trzeba nigdzie zakładać.
- Jak rozdzielić przepływ na dwie ścieżki w zależności od warunku, węzłem „IF”.
- Jak przetwarzać elementy jeden po drugim, w kontrolowanym tempie, węzłem „Loop Over Items”.
- Jak oprzeć decyzję przepływu na ocenie modelu, a nie na dopasowaniu słowa.

- 1 Zanim otworzysz n8n, przygotuj to, z czego przepływ będzie czerpał. Wklej do przeglądarki ten adres:

<https://news.google.com/rss/search?q=LEGO+when:1d&hl=pl&gl=PL&ceid=PL:pl>

Zobaczysz surowy plik pełen znaczników w ostrych nawiasach. Nie przejmuj się, jeśli nic z tego nie rozumiesz. Sprawdź na razie tylko jedną rzecz: czy między znacznikami widać polskie tytuły artykułów. Jeśli tak, to znaczy, że źródło działa.

Ten adres to zwykłe wyszukiwanie w Google News, tyle że w formie, którą potrafi wczytać maszyna. Fragment po `q=` to twoje zapytanie, `when:1d` zawęży wyniki do ostatniej doby, a trzy końcowe parametry ustawiają język i kraj na polski.

WSKAZÓWKA

W tej książce śledzimy przykładową markę LEGO. Ty zamień LEGO na nazwę marki, która cię interesuje. Trzy rady z praktyki. Jeśli nazwa znaczy również coś innego, obejmij ją cudzysłowem, bo inaczej monitoring firmy „Sokołów” przyniesie ci także wiadomości z gminy Sokołów. Jeśli chcesz węższe okno, wpisz `when:12h`. A jeśli po zmianie zapytania plik w przeglądarce jest pusty, znaczy to, że o twojej marce nic dziś nie napisano, i wtedy na czas ćwiczeń wróć do szerszego hasła.

2. Utwórz nowy, pusty przepływ. Dodaj węzeł „Schedule Trigger”, ustaw „Trigger Interval” na „Days”, godzinę na 7am i minutę na 0. To ten sam węzeł, który poznałeś w rozdziale 3.
3. Dodaj węzeł „RSS Read” i połącz go z wyzwalaczem. W polu „URL” wklej adres przygotowany w kroku 1. Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT listę wiadomości, a nad nią liczbę znalezionych elementów. Przy popularnej marce i oknie jednodniowym może to być kilkanaście pozycji.

The screenshot displays the n8n workflow editor. On the left, the 'INPUT' section shows a 'Schedule Trigger' node configured with a 1-day interval at 03:02:18 AM on August 23, 2026. The 'Parameters' tab for the 'RSS Read' node shows the URL: `https://news.google.com/rss/search?q=LEGO+when:1d&hl=pl&gl=PL&ceid=PL;g`. The 'OUTPUT' panel on the right shows 18 items, including news articles about LEGO City, Radio Szczecin, and LEGO Play Station.

title	link	pubDate	content
LEGO City: Po bandzie MAX - Pilot WP	https://news.google.com/rss/articles/CBMipwFBV95ScUxPUUR3NkpxSORQakFdaVcCTUNwcl0tIdF5MIV2eGRadipGYWpa1NWQmJGVVFGSGZYUINJCTVvbGNWMMvdsVZUQ03U0SPVjCmF4NKFzRGNvWmJScnoId3RibDRMTBoblpJOnVxeGExMwI6QTOxaTFUuU95VhDrU2dXeUvZc085cXhEeJRWTVCUlr2NEVV6TEF2cig1Q2xNeHdabw?oc=5	Sat, 22 Aug 2026 23:31:43 GMT	https://news.google.com/rss/articles/CBMipwFBV95ScUxPUUR3NkpxSORQakFdaVcCTUNwcl0tIdF5MIV2eGRadipGYWpa1NWQmJGVVFGSGZYUINJCTVvbGNWMMvdsVZUQ03U0SPVjCmF4NKFzRGNvWmJScnoId3RibDRMTBoblpJOnVxeGExMwI6QTOxaTFUuU95VhDrU2dXeUvZc085cXhEeJRWTVCUlr2NEVV6TEF2cig1Q2xNeHdabw?oc=5
Wielka gratka dla fanów klocków Lego [ZDJĘCIA] - Radio Szczecin	https://news.google.com/rss/articles/CBMiHAFBV95ScUxONdZLbJjQTIhOxMTfHKNDVRhDnYzaTxuQJluNzZDZUFUzTVBT21QeW9U0MbURHYkxSUfzSmdUdWdscnA1Z3ZYMWw2T3ZDYU4NtZFS5SMYJlXNGYwUE4yamZoODRvWG44RmIneWzMV9aMDJPUJ2ZkVXJLWOpNWJE7oc=5	Sat, 22 Aug 2026 12:44:00 GMT	https://news.google.com/rss/articles/CBMiHAFBV95ScUxONdZLbJjQTIhOxMTfHKNDVRhDnYzaTxuQJluNzZDZUFUzTVBT21QeW9U0MbURHYkxSUfzSmdUdWdscnA1Z3ZYMWw2T3ZDYU4NtZFS5SMYJlXNGYwUE4yamZoODRvWG44RmIneWzMV9aMDJPUJ2ZkVXJLWOpNWJE7oc=5
LEGO PlayStation zachwyci graczy! Wyciekły szczegóły nowego zestawu - PPE.pl	https://news.google.com/rss/articles/CBMipAFBV95ScUxPTVNWvWvWFYzYw9ReUkLV9nXzVtaFtIuKXpoTEkC3pVpNCvYodKbZvX2J3UuRlJhZ0VnV2cdNoamMTR25BcmV3bPpEiQbINkSDJhRQFIz3pNfVYmXhwREIQLVdVWVXp5MDfGVXVRRE9JckhyVGNjYpNrwNhzHvU3luWJleE1PTDRXVFZGaHhFvIkS0R0TjpxVtdBwMJGa7oc=5	Sat, 22 Aug 2026 08:38:34 GMT	https://news.google.com/rss/articles/CBMipAFBV95ScUxPTVNWvWvWFYzYw9ReUkLV9nXzVtaFtIuKXpoTEkC3pVpNCvYodKbZvX2J3UuRlJhZ0VnV2cdNoamMTR25BcmV3bPpEiQbINkSDJhRQFIz3pNfVYmXhwREIQLVdVWVXp5MDfGVXVRRE9JckhyVGNjYpNrwNhzHvU3luWJleE1PTDRXVFZGaHhFvIkS0R0TjpxVtdBwMJGa7oc=5

Rys. 5.1. Wynik zapytania prasowego, gotowy do oceny

- 4 Przyjrzyj się teraz kolumnie title. Każdy tytuł kończy się myślnikiem i nazwą źródła - redakcji, na zasadzie: „Wielka gratka dla fanów klocków Lego - Radio Szczecin”. To wygodne, bo w jednym polu masz i temat, i źródło, a przy ocenie istotności źródło bywa ważniejsze od nagłówka.

Zajrzyj też do kolumny content. Zamiast treści artykułu znajdziesz w niej znaczniki HTML z linkiem. Tak wygląda ten kanał i nic z tym nie zrobimy. Pracujemy więc na samym tytule.

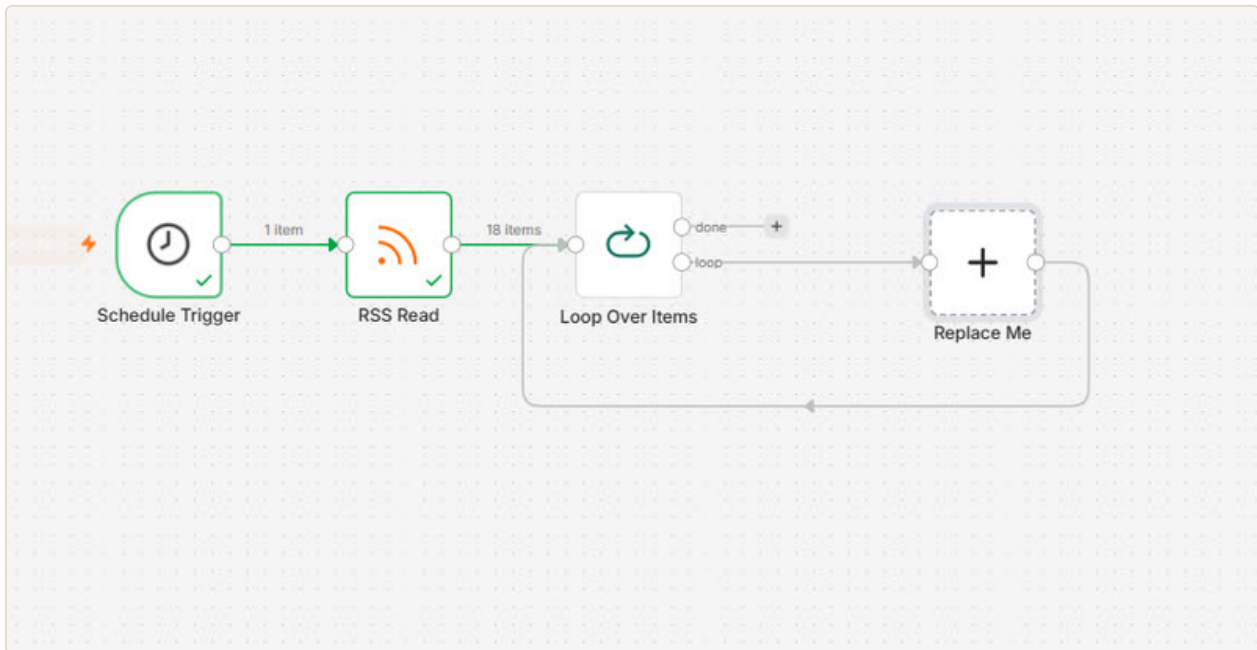
! UWAGA

Model nie przeczytał artykułu. Oцени sam nagłówek i nazwę redakcji, czyli mniej więcej tyle, ile ty widzisz, przewijając listę wyników. To wystarcza na tym etapie, którego zadaniem jest odrzucenie oczywistych nietrafień, a nie dogłębna analiza. Nie myl tego z czytaniem ze zrozumieniem: gdy potrzebujesz oceny opartej na treści, trzeba będzie najpierw treść pobrać, a to już inna operacja.

Dodaj pętlę

- 5 Kliknij znak plus przy prawej krawędzi węzła „RSS Read”. Wpisz „Loop” i wybierz węzeł „Loop Over Items” (pętla po elementach). Na ekranie zobaczysz: nie jeden węzeł, ale dwa. Obok pętli n8n dołożył sam z siebie węzeł z napisem „Replace Me” (zastąp mnie). To nie jest błąd i nie jest to część twojego przepływu. To odpowiedź. n8n pokazuje ci gotowy kształt pętli i mówi wprost: tutaj wstaw pracę, którą chcesz powtarzać dla każdego elementu, a potem zawrócić. Popatrz na sam węzeł pętli. Ma on dwa wyjścia po prawej stronie, jedno pod drugim, podpisane „done” (gotowe) i „loop” (pętla). To pierwszy węzeł, który ma dwa wyjścia, i pierwszy, który sam z siebie nie robi nic użytecznego. Cała jego wartość bierze się z tego, jak go połączysz.

Kliknij zatem węzeł „Replace Me” i usuń go klawiszem Delete. Zbudujesz tę część sam, żeby zobaczyć, z czego się składa.



Rys. 5.2. Dwa wyjścia, dwa zupełnie różne znaczenia

- 6 Otwórz węzeł i sprawdź pole „Batch Size” (rozmiar paczki). Zostaw wartość 1 i zamknij panel.



WSKAZÓWKA

„Batch Size” równy 1 znaczy „jeden element na raz”. Węzeł podłączony do wyjścia „loop” zobaczy dokładnie jedną wiadomość na każdym okrążeniu. Większa wartość przyspiesza pracę, ale przy zapytaniach do modelu zwiększa też ryzyko błędu o przekroczonym limicie.



DLA CIEKAWYCH

Ten przepływ zadziałałby również bez pętli. Pamiętasz zasadę z rozdziału 3, że węzeł wykonuje się raz na każdy element wejściowy? Węzeł OpenAI podpięty wprost pod RSS Read i tak zapytałby model np. osiemnaście razy.

Różnica polega na tempie i kontroli. Bez pętli n8n wysła zapytania równolegle, wszystkie naraz, i przy większej liczbie elementów może zakończyć się to błędem o przekroczonym limicie. Z pętlą idą one jedno po drugim, a ty widzisz, na którym elemencie przepływ się potknął. W małym przepływie to wygoda, w rozbudowanym to jedyny sposób, żeby zrozumieć, gdy coś przestanie działać.

Wstaw ocenę modelu

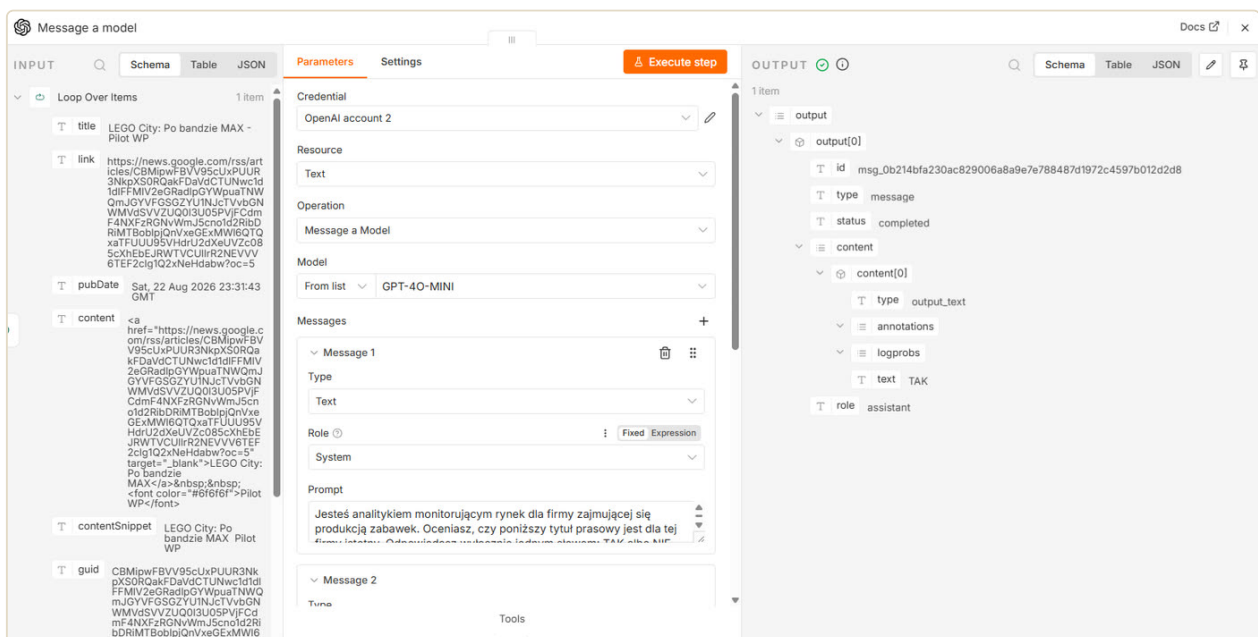
- 7 Z wyjścia „loop” dodaj węzeł „OpenAI” i wybierz akcję „Message a model”, tak jak w rozdziale 4. Wskaż dane dostępne, które zapisałeś wtedy: pojawią się na liście gotowe, bez ponownego wklejania klucza.

- 8 Dodaj wiadomość z rolą „System” i wpisz w niej:
„Jesteś analitykiem monitorującym rynek dla firmy zajmującej się [twoja branża]. Oceniasz, czy poniższy tytuł prasowy jest dla tej firmy istotny. Odpowiadasz wyłącznie jednym słowem: TAK albo NIE. Bez wyjaśnień, bez kropki, bez żadnego dodatkowego tekstu.”
- 9 Dodaj drugą wiadomość, z rolą „User”, i przeciągnij do niej pole title z panelu danych wejściowych po lewej stronie.

NA EKRANIE ZOBACZYSZ

pod polem podgląd z prawdziwym tytułem pierwszej wiadomości.

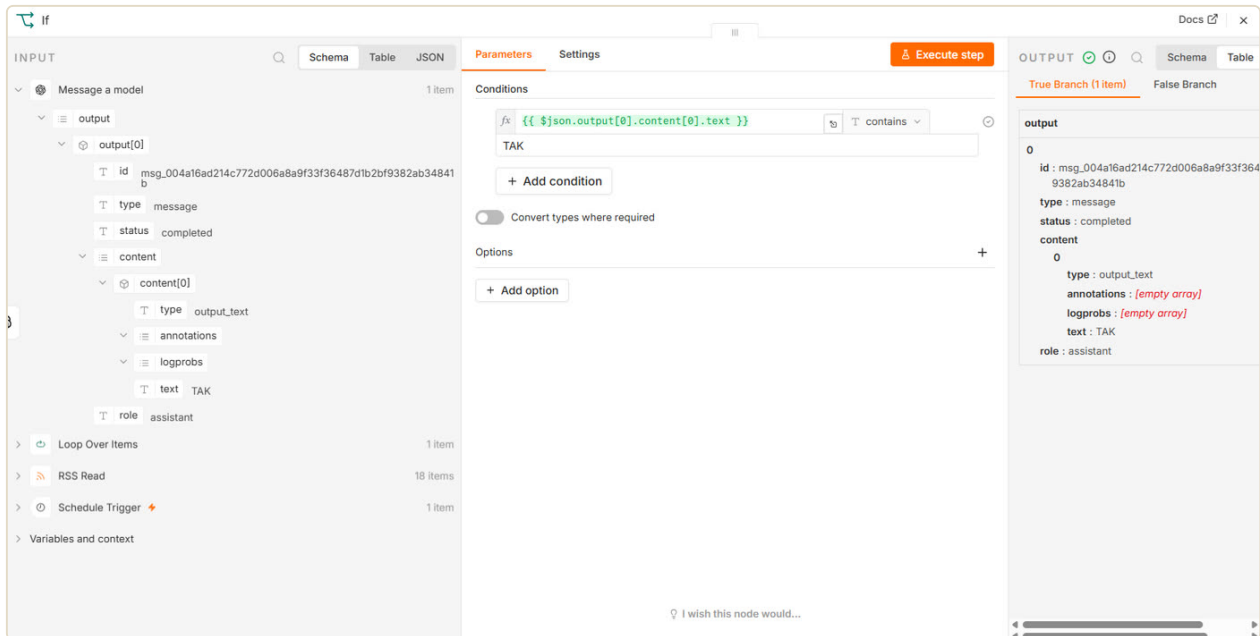
- 10 Kliknij „Execute step” i sprawdź, co model odpowiedział. Rozwiń w panelu OUTPUT gałąź content, a pod nią 0, i znajdź pole text. Powinno zawierać jedno słowo.



Rys. 5.3. Polecenie, które zmusza model do odpowiedzi jednym słowem

Rozdziel ścieżki

- 11 Z wyjścia węzła OpenAI dodaj węzeł „IF”. Kliknij „Add condition” (dodaj warunek). Do górnego pola warunku przeciągnij text z panelu danych wejściowych, czyli odpowiedź modelu. Typ porównania ustaw na „String”, operator na „contains” (zawiera), a w dolnym polu wpisz TAK.



Rys. 5.4. Warunek, który rozdziela wiadomości na dwie ścieżki



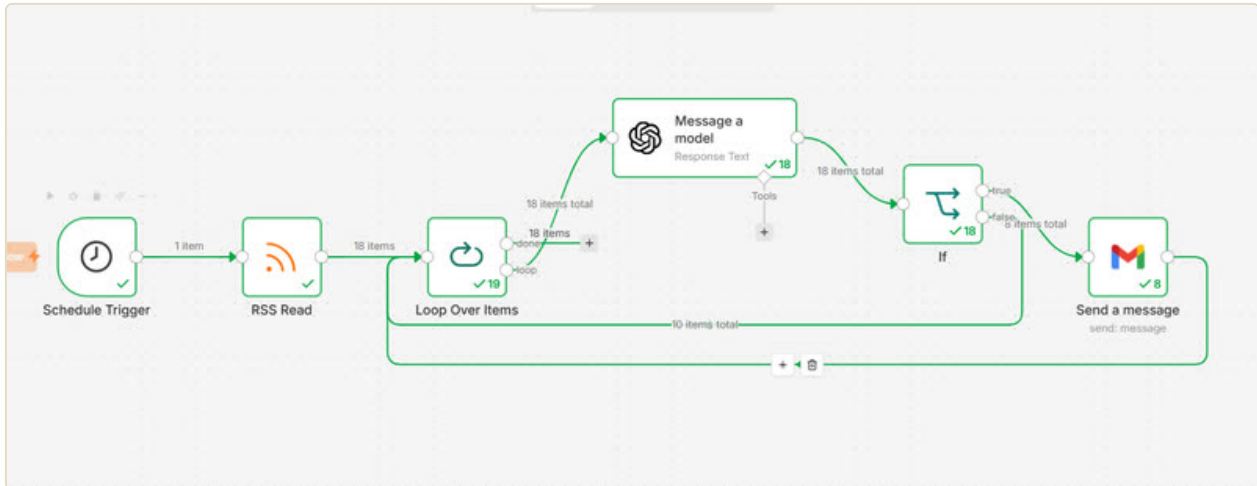
WSKAZÓWKI

Używamy „contains” zamiast „is equal to” (jest równe) celowo. Model potrafi dokleić kropkę, spację albo cudzysłów i wtedy porównanie na równość zawiedzie, mimo że odpowiedź jest poprawna. „Zawiera” wybacza takie drobności.

- 12 Zamknij panel. Węzeł IF ma teraz dwa wyjścia: „true” (prawda) u góry i „false” (fałsz) u dołu. Z wyjścia „true” dodaj węzeł „Gmail” - „Send a message”. Ustaw „Resource” na „Message” i „Operation” na „Send”, tak jak w rozdziale 3. W polu „To” wpisz swój adres, w „Subject” na przykład „Monitoring: nowa wiadomość”, a do treści przeciągnij title i link, pamiętając o wciśnięciu Entera między jednym a drugim.

Zamknij pętlę

- 13 To najważniejszy krok w tym rozdziale i jedyny, który wygląda nienaturalnie. Przeciągnij połączenie od wyjścia węzła „Gmail” z powrotem do wejścia węzła „Loop Over Items”. Potem zrób to samo z wyjściem „false” węzła IF.



Rys. 5.5. Rozgałęzienie zamknięte w pętli

Powód jest prosty, choć nieoczywisty. Węzeł „Loop Over Items” nie wypuści następnego elementu, dopóki nie zostanie sygnał, że poprzedni został obsłużony. Sygnałem jest właśnie ta linia zwrotna. Obie ścieżki muszą wracać, także ta, na której nic się nie stało. Wiadomość odrzucona przez filtr również wymaga zamknięcia sprawy, bo pętla nie rozróżnia sukcesu od pominięcia, tylko obsłużone od nieobsłużonych.

Wyjście „done”, jak widzisz, w tym przypadku wisi puste. Ale tak ma być, więc zostawiamy to w ten sposób.

Uruchom

- 14 Kliknij „Execute workflow”.

NA EKRANIE ZOBACZYSZ

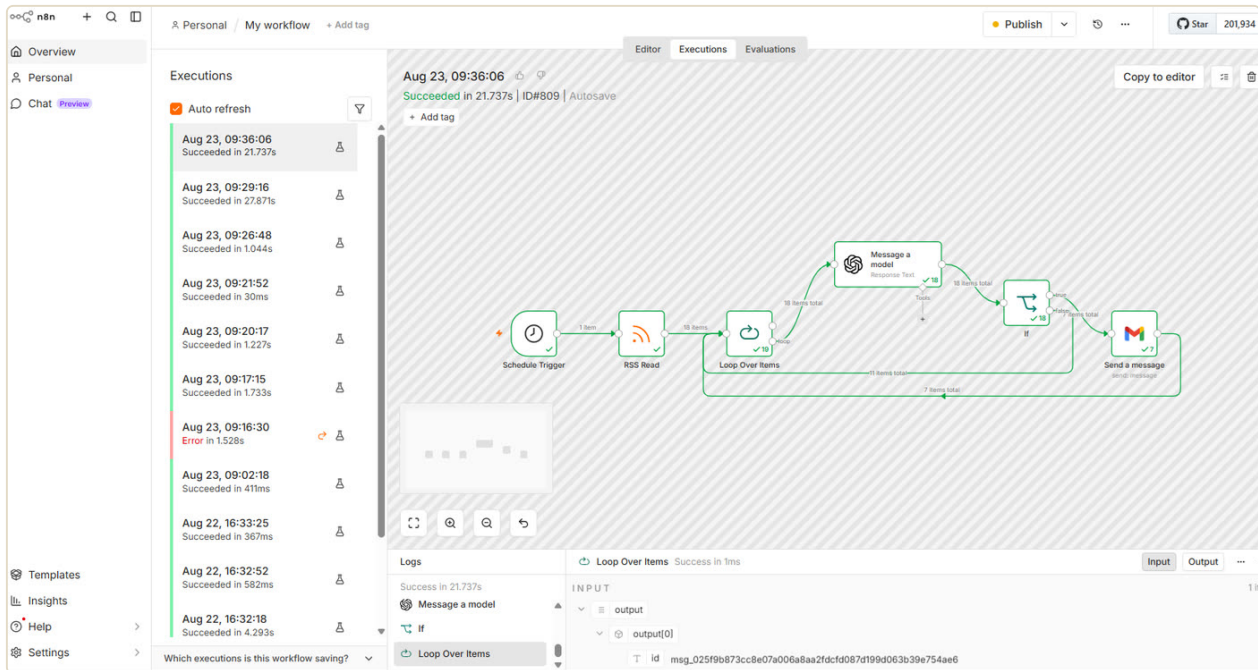
węzły zapalające się po kolei, a potem cykl powtarzający się w kółko: OpenAI, IF, czasem Gmail, powrót. Licznik pod węzłem OpenAI rośnie z każdym okrążeniem. Na końcu zapali się wyjście „done”.

Sprawdź, czy działa

Otwórz skrzynkę pocztową. Powinno w niej być tyle wiadomości, ile tytułów model uznał za istotne.

Może być ich zero i to również jest poprawny wynik, nie awaria. Znaczy tyle, że dziś nie napisano nic, co dotyczyłoby wybranej marki. Żeby się upewnić, że przepływ naprawdę zadziałał, zajrzyj do węzła IF i sprawdź, ile elementów poszło każdą ze ścieżek. Suma powinna zgadzać się z liczbą wiadomości z węzła RSS Read.

Zajrzyj też do zakładki „Executions”. Uruchomienie powinno być zielone, a przy węźle OpenAI powinna stać liczba wykonań równa liczbie wiadomości.



Rys. 5.6. Poranny przegląd, który robi się sam

Co może pójść nie tak

- Przepływ nie kończy się i kręci w nieskończoność, albo przetwarza tylko pierwszą wiadomość i staje. To ten sam błąd widziany z dwóch stron: brakuje linii zwrotnej. Sprawdź, czy zarówno wyjście węzła Gmail, jak i wyjście „false” węzła IF wracają do wejścia węzła „Loop Over Items”. Obie, nie jedna.
- Dostałeś kilka maili o tej samej sprawie. Ten sam temat opisało kilka redakcji, a dla przepływu to kilka osobnych wiadomości. To nie usterka, tylko to, jak wygląda surowy monitoring prasowy.
- Wszystko przechodzi przez filtr, także rzeczy oczywiście nieistotne. Sprawdź w panelu OUTPUT węzła OpenAI, co model faktycznie napisał. Jeśli odpowiada zdaniem zamiast jednym słowem, słowo TAK pojawi się w nim niemal zawsze i warunek przepuści wszystko. Zaostrz polecenie systemowe.
- Nic nie przechodzi, choć model odpowiada TAK. Warunek w węźle IF wskazuje na złe pole. Sprawdź, czy po lewej stronie warunku jest na pewno text z odpowiedzi modelu, a nie tytuł artykułu.

Wariacje na temat

Obserwuj dwie marki naraz. Nie potrzebujesz do tego drugiego węzła. Wystarczy zmienić zapytanie w adresie, wpisując `q=LEGO+OR+Playmobil+when:1d`. Reszta przepływu obsłuży to bez żadnej zmiany.

Zamiast maila poproś model o jedno zdanie streszczenia. Zmień polecenie tak, żeby zwracało ocenę i krótkie uzasadnienie, i wstaw to uzasadnienie do treści wiadomości. Kosztuje odrobinę więcej, bo model pisze więcej, ale poranny przegląd czyta się wtedy znacznie szybciej.

? DLA CIEKAWYCH

Monitoring konkurencji nie zaczął się od internetu ani od RSS. W drugiej połowie XIX wieku powstały w Europie i Stanach biura wycinków prasowych: firmy, które za abonament przeglądały stosy gazet i wysyłały klientowi wycięte nożyczkami wzmianki o jego nazwisku, firmie albo branży. Agencje reklamowe należały do ich najwierniejszych klientów.

Zmieniło się narzędzie i tempo. Sam pomysł jest ten sam od stu pięćdziesięciu lat: ktoś albo coś musi patrzeć codziennie, żebyś ty nie musiał.

Warto jednak zauważyć, czego ani wycinki, ani twój przepływ nie mierzą. Trzy artykuły o jednym wydarzeniu to nadal jedno wydarzenie. Licząc wzmianki, mierzysz nie tyle to, co się stało, ile to, ilu redakcjom chciało się o tym napisać. Dwie zupełnie różne wielkości, mylone czasem w raportach od dziesięcioleci.

Przepływ w repozytorium

Pobierz plik 05-monitor-konkurencji.json ze strony heuristica.pl/automatyzacje-n8n i zaimportuj go tak samo jak w rozdziale 1. Po imporcie wskaż własne dane dostępowe do OpenAI i Gmaila oraz podmień markę w adresie kanału.

Rozgałęzienie odpowiada na pytanie, co zrobić z tym elementem. Pętla odpowiada na pytanie, co z następnym. Wszystko, co zbudujesz dalej, będzie kombinacją odpowiedzi na te dwa pytania, w coraz bardziej rozbudowanych wariantach.

Rozdział 6. Syntezator researchu

Research to czynność, którą marketer wykonuje stale i prawie nigdy nie kończy. Otwierasz dwanaście zakładek, czytasz osiem, zapamiętujesz trzy, a do briefu trafia jedna, zwykle ta ostatnia przeczytana. Nie dlatego, że była najlepsza, tylko dlatego, że była najświeższa w pamięci.

Problemem nie jest brak informacji, tylko to, że nigdy nie widzimy ich wszystkich naraz.

Czytamy sekwencyjnie, a wnioskować chcielibyśmy równolegle.

W rozdziale 5 przepływ oceniał każdy element osobno i osobno o nim decydował. Tutaj zrobisz coś nowego: zbierzesz wyniki z rozproszonych źródeł w jedno miejsce i dopiero wtedy poprosisz o ocenę całości.

Co zbudujemy

System, który dostaje listę adresów, dla każdego z osobna wyciąga jedną kluczową myśl, a na końcu dopisuje jeszcze jeden wiersz: wniosek z porównania wszystkich źródeł razem. Wynikiem jest tabela w arkuszu Google, gotowa do wklejenia w briefie albo prezentacji.

Czas i koszt

Budowa: około trzydziestu minut. To długi przepływ - dziewięć węzłów, ale żaden z nich nie jest trudny.

Koszt: kilka groszy za uruchomienie. Więcej niż w rozdziale 5, bo model dostaje tym razem całe strony, a nie same nagłówki.

Czego się nauczysz

- Jak węzłem „HTTP Request” pobrać treść dowolnej strony, bez API tego serwisu.
- Jak jednym węzłem obciąć z pobranej strony wszystko, za co nie chcesz płacić.
- Jak scałić rozproszone wyniki w jeden element, żeby model mógł ocenić je razem, a nie po kolei.

Przygotuj arkusz

- 1 Zanim otworzysz n8n, załóż nowy arkusz w Google Sheets. Nazwij go „Syntezator researchu” i zrób w nim dwie zakładki.

W pierwszej, „Linki”, w komórce A1 wpisz nagłówek Link, a poniżej wklej adresy źródeł. Na potrzeby tego rozdziału użyjemy sześciu blogów, na których wypowiadają się ludzie mający coś do powiedzenia o marketingu, produkcie i zachowaniach użytkowników:

<https://seths.blog>

<https://www.nngroup.com/articles/>

<https://sparktoro.com/blog>

<https://baymard.com/blog>

<https://world.hey.com/dhh>

<https://world.hey.com/jason>

W drugiej zakładce, „Wyniki”, w komórce A1 wpisz Link, a w B1 Teza. Resztę zostaw pustą, wypełni ją przepływ.



WSKAZÓWKA

Te sześć adresów to nie przypadkowy zestaw. Trzy pierwsze reprezentują podejście badawcze i danych, dwa ostatnie są celowo kontrariańskie, a Godin siedzi gdzieś pośrodku i pisze aforyzmami. Chodzi tu o rozbieżność. Sześć źródeł mówiących to samo da nudne podsumowanie. Trzy zgodne i dwa polemiczne dadzą coś, co warto przeczytać.

Później podmienisz te adresy na własne. Możesz wskazywać zarówno strony główne blogów - wtedy dostaniesz obraz tego, czym dane źródło żyje w tej chwili - jak i pojedyncze artykuły, jeśli badasz konkretny temat. Mechanizm jest identyczny.

Wczytaj linki

- 2 Utwórz nowy, pusty przepływ. Dodaj wyzwalacz „Manual trigger” (uruchom ręcznie), ten sam, którego używałeś w rozdziałach 1, 3 i 4.

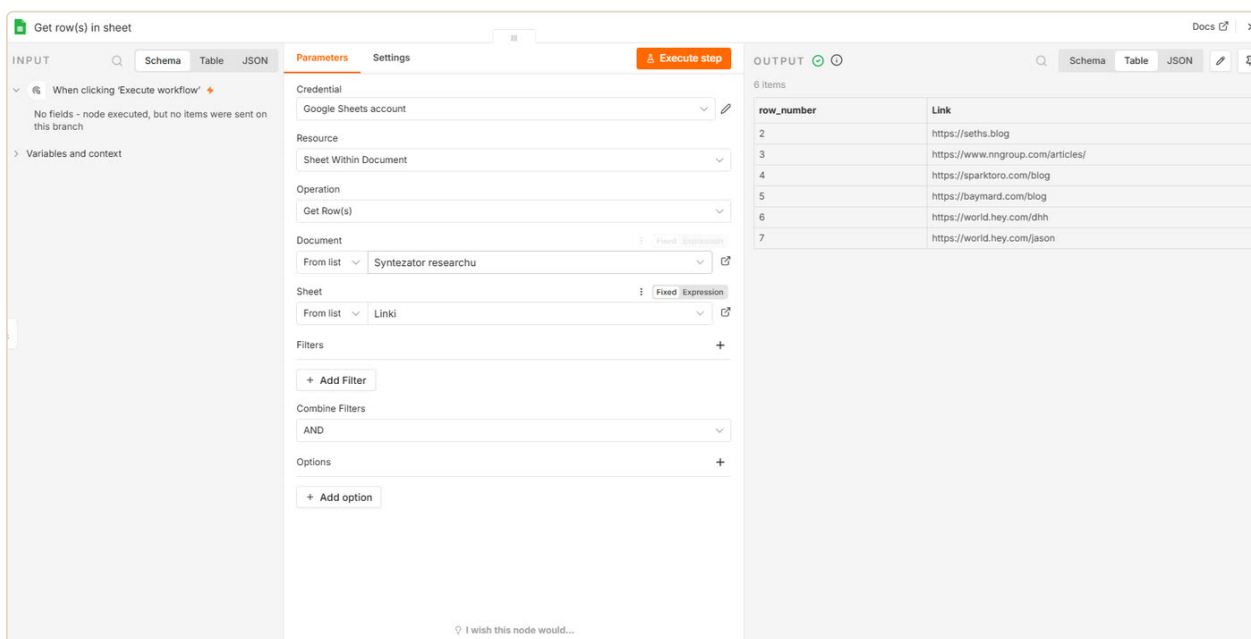
- 3 Kliknij plus przy wyzwalaczu, wpisz „Google Sheets” i wybierz z listy akcji „Get row(s) in sheet” (pobierz wiersze arkusza).

Przy polu „Credential” kliknij „Create new credential” i „Sign in with Google”. To ten sam ekran logowania Google co przy Gmailu w rozdziale 3, tym razem z prośbą o dostęp do Arkuszy zamiast do poczty.

W polu „Document” wybierz z listy swój arkusz, a w polu „Sheet” zakładkę „Linki”. Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT sześć elementów, każdy z jednym polem Link.



Rys. 6.1. Wczytanie listy linków z arkusza

WSKAZÓWKA

Zwróć uwagę na to, co się właśnie stało. Lista, którą przepływ przetwarza, leży poza n8n, w arkuszu, do którego masz dostęp z telefonu i który możesz udostępnić komuś z zespołu. Żeby zmienić zakres monitoringu, nikt nie musi już otwierać n8n ani rozumieć, jak przepływ działa. Wystarczy dopisać wiersz.

To jest wzorzec, który warto zapamiętać na później: przepływ robi robotę, arkusz nią steruje.

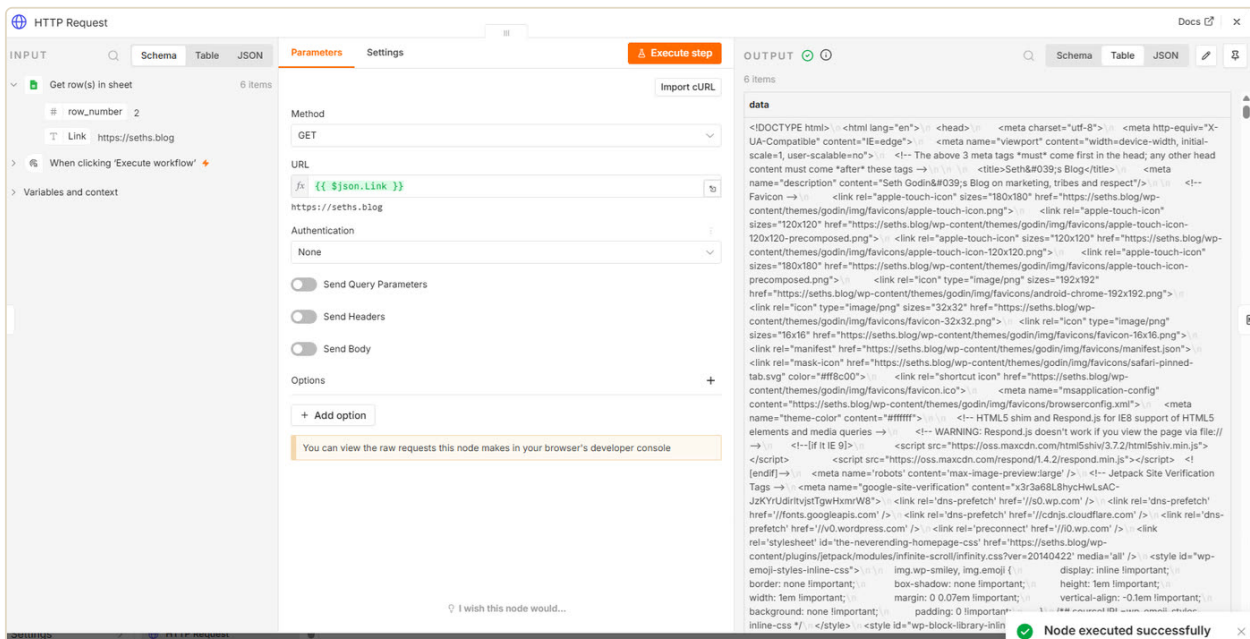
Pobierz treść

- 4 Kliknij plus przy węźle Google Sheets, wpisz „HTTP Request” i wybierz węzeł z listy.

W polu „Method” (metoda) zostaw GET. W polu „URL” przeciągnij pole Link z panelu danych wejściowych po lewej stronie. Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT pole data, a w nim długi ciąg zaczynający się od <!DOCTYPE html>, pełen znaczników w ostrych nawiasach. To surowy kod pierwszej strony z listy.



Rys. 6.2. Strona internetowa, zanim przeglądarka ją pokaże

UWAGA

Węzeł HTTP Request pobiera stronę tak, jak dostaje ją przeglądarka: jako kod, nie jako tekst. Zwykle to działa. Bywa jednak, że strona jest zbudowana w całości ze skryptu i treść pojawia się dopiero po jego wykonaniu w przeglądarce. Wtedy węzeł zobaczy niemal pustą stronę i nie ma na to prostej rady. Większość blogów i portali informacyjnych pobierze się bez problemu, część nowoczesnych serwisów nie.

Obetnij, za co nie chcesz płacić

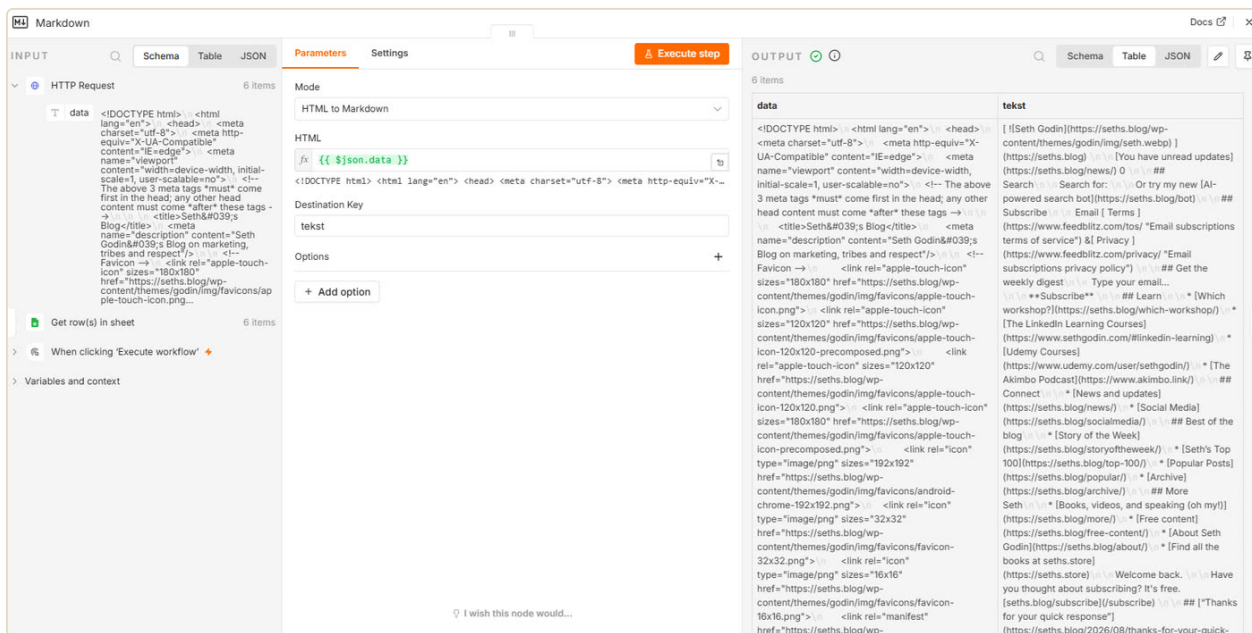
- 5 Zanim treść trafi do modelu, trzeba usunąć z niej to, co nie jest treścią. Kliknij plus przy węźle „HTTP Request”, wpisz „Markdown” i wybierz węzeł z listy.

Ustaw „Mode” (tryb) na „HTML to Markdown”. W polu „HTML” przeciągnij pole data z panelu danych wejściowych. W polu „Destination Key” (pole docelowe) wpisz „tekst” zamiast domyślnej wartości, żeby wynik nie nadpisał oryginału.

Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT nowe pole tekst, a w nim ten sam artykuł, tyle że bez znaczników. Zostały nagłówki, akapity i linki. Zniknęły metadane, skrypty, style i favicony.



Rys. 6.3. Ta sama strona, po odjęciu wszystkiego, co nie jest tekstem

! UWAGA

Ten jeden węzeł jest najtańszą optymalizacją w tym poradniku.

Na jedną stronę artykułu przypada zwykle dziesięć razy więcej kodu niż tekstu. Wysyłając surowy HTML do modelu, płaciłbyś za menu, stopkę, skrypty reklamowe i deklaracje stylów. Dziewięć dziesiątych rachunku szłoby na rzeczy, których model i tak nie potrzebuje.

Jest jeszcze drugi powód. Każdy model ma ograniczoną pamięć roboczą, do której musi się zmieścić wszystko, co mu wysyłasz. Surowa strona potrafi tę pamięć przepełnić i wtedy przepływ nie zwróci gorszego wyniku, tylko żaden. Węzeł Markdown zdejmuje oba problemy naraz, jednym ustawieniem.

Oceń każde źródło

- 6 Kliknij plus przy węźle „Markdown”. Dodaj węzeł „OpenAI” z akcją „Message a model”, tak jak w rozdziale 4, i wskaż dane dostępowe zapisane wtedy.
- 7 W polu „Model” wybierz model z bieżącej listy.

! UWAGA

Nie wybieraj pozycji podpisanej „GPT-4”. Nazwa bez sufiksu nie oznacza wersji podstawowej, tylko wersję starą. Model z tamtej generacji ma pamięć roboczą kilkadziesiąt razy mniejszą niż dzisiejsze i przy pełnej stronie internetowej odmówi pracy z komunikatem o przekroczeniu okna kontekstowego.

Ta pułapka jest wredna, bo „GPT-4” brzmi na liście najbardziej znajomo ze wszystkiego, co tam widzisz. Wybierz więc np. GPT-5 lub inny, nowy model.

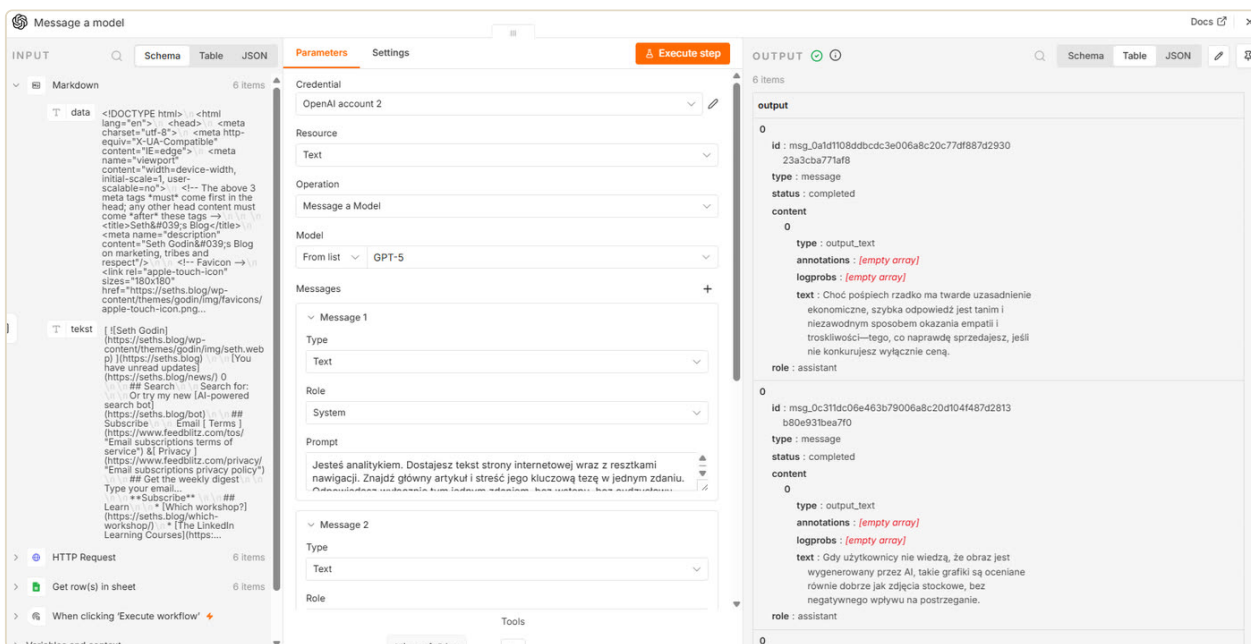
8 Dodaj wiadomość z rolą „System” i wpisz w niej:

„Jesteś analitykiem. Dostajesz tekst strony internetowej, wraz z resztkami nawigacji. Znajdź główną treść i streść w jednym zdaniu, czym to źródło się w tej chwili zajmuje i jakie stanowisko zajmuje. Odpowiadasz wyłącznie tym jednym zdaniem, bez wstępu, bez cudzysłowu.”

9 Dodaj drugą wiadomość, z rolą „User”, i przeciągnij do niej pole tekst z panelu danych wejściowych. Uwaga: tekst, nie data. Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT, w polu text, jedno zdanie.



Rys. 6.4. Cała strona sprowadzona do jednego zdania

**WSKAZÓWKA**

Dwa pola o mylnie podobnych nazwach stoją teraz obok siebie. tekst to treść strony, którą wysyłasz do modelu. text to odpowiedź, którą model odsyła. Pierwsze wpisałeś sam w węźle Markdown, drugie nazwałeś tak OpenAI. Jeśli w którymś z kolejnych kroków dostaniesz nie to, czego oczekiwałeś, sprawdź najpierw, które z tych dwóch pól przeciągnąłeś.

Zapisz wynik każdego źródła

- 10** Kliknij plus przy węźle OpenAI, dodaj kolejny węzeł „Google Sheets” i wybierz akcję „Append row in sheet” (dopisz wiersz). Użyj tych samych danych dostępowych, wskaż ten sam dokument i zakładkę „Wyniki”.

W mapowaniu kolumn przeciągnij text z odpowiedzi modelu do kolumny Teza.

Kolumna Link wymaga więcej uwagi, bo adresu nie ma już w panelu wejściowym. Rozwiń w lewym panelu węzeł „Get row(s) in sheet”, ten sam, od którego zaczynał się przepływ, znajdź w nim pole Link i przeciągnij je do pustego pola. W polu pojawi się dłuższe wyrażenie niż zwykle, bo n8n musi zapisać nie tylko nazwę pola, ale i to, z którego węzła je bierze.

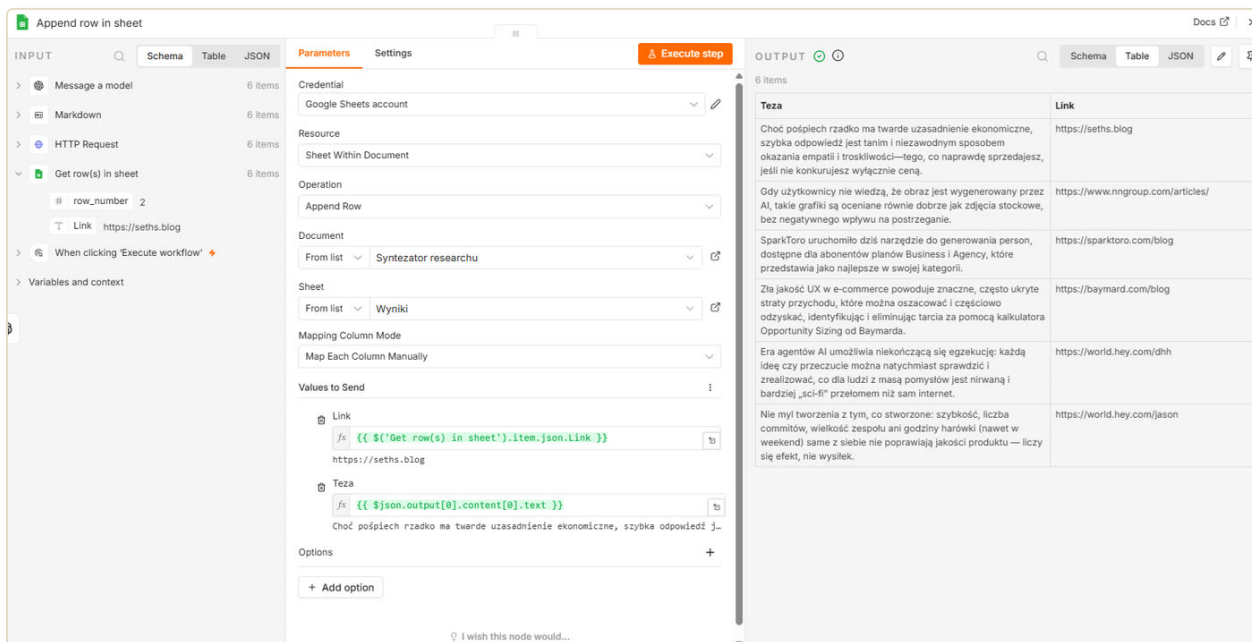
**UWAGA**

Węzeł nie przekazuje dalej tego, co dostał. Przekazuje to, co sam wyprodukował.

Adres z arkusza zniknął już w węźle HTTP Request, który w miejsce linku zwrócił treść strony. Potem Markdown zwrócił tekst, a OpenAI zdanie. Na każdym kroku poprzednia zawartość była zastępowana nową.

Nie znaczy to, że dane przepadły. Panel po lewej stronie zawiera listę wszystkich wcześniejszych węzłów i z każdego z nich możesz przeciągnąć dowolne pole, tak samo jak z węzła sąsiedniego. n8n sam dopisze, skąd je bierze.

Zapamiętaj to jako pytanie, które warto sobie zadać, gdy czegoś nie widać w panelu wejściowym: nie „gdzie to się podziało”, tylko „który węzeł widział to ostatni”.



Rys. 6.5. Mapowanie wyniku na kolumny arkusza

- 11 Kliknij „Execute workflow”, czyli uruchom cały przepływ, nie pojedynczy krok. Poczekaj i otwórz zakładkę „Wyniki” w arkuszu.

NA EKRANIE ZOBACZYSZ

sześć nowych wierszy, każdy z adresem i jednym zdaniem. Bez pętli. Bez węzła „Loop Over Items” z rozdziału 5.

WSKAZÓWKA

W rozdziale 5 pętla była potrzebna, bo zależało nam na tempie: zapytania miały iść jedno po drugim, a nie wszystkie naraz. Tutaj tempo nie ma znaczenia. Sześć stron pobiera się równolegle, sześć zapytań idzie równolegle, sześć wierszy trafia do tego samego arkusza w dowolnej kolejności.

Reguła jest prosta: pętli potrzebujesz wtedy, gdy kolejność albo tempo mają znaczenie. W pozostałych przypadkach $n8n$ i tak wykona każdy węzeł raz na każdy element, sam z siebie, szybciej i bez dodatkowych połączeń.

Scal i zsyntetyzuj

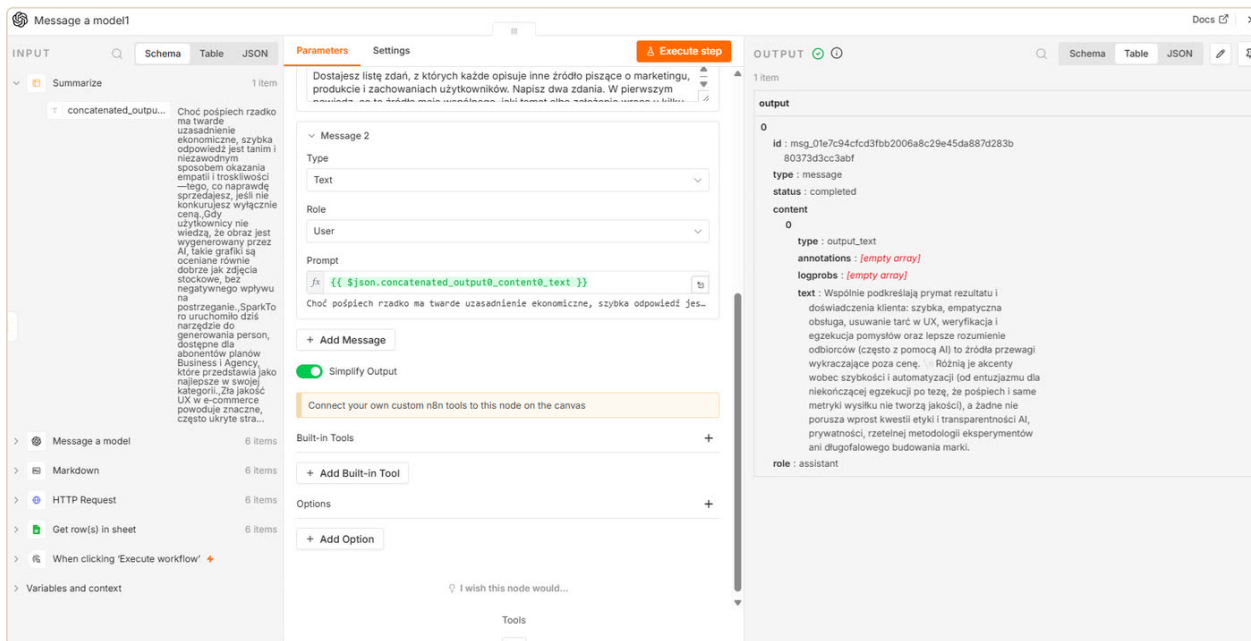
- 12 Wróć do węzła OpenAI z kroku 6. Z jego wyjścia wychodzi już jedno połączenie, do zapisu w arkuszu. Dociągnij z tego samego wyjścia drugie połączenie, do nowego węzła. Wpisz „Summarize” i wybierz węzeł z listy.

Przepływ rozdziela się teraz na dwie równoległe gałęzie wychodzące z jednego punktu. To nie to samo, co rozgałęzienie warunkowe z rozdziału 5. Tam element szedł jedną drogą albo drugą. Tu każdy element idzie oboma naraz.

- 13 W sekcji „Fields to Summarize” (pola do podsumowania) wskaż pole text, czyli odpowiedź modelu z kroku 6. Jako operację wybierz „Concatenate” (sklej), a w opcji separatora ustaw nową linię. Kliknij „Execute step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT już nie sześć elementów, tylko jeden. W środku jedno pole tekstowe, a w nim wszystkie sześć zdań, jedno pod drugim.



Rys. 6.6. Sześć osobnych odpowiedzi staje się jednym tekstem

- 14 Kliknij plus przy węźle „Summarize” i dodaj drugi węzeł „OpenAI”, z tymi samymi danymi dostępowymi i tym samym modelem.

Wiadomość z rolą „System”:

„Dostajesz kilka zdań, z których każde opisuje inne źródło piszące o marketingu, produkcie i zachowaniach użytkowników. Napisz dwa zdania. W pierwszym powiedz, co te źródła mają wspólnego, jaki temat albo założenie wraca u kilku z nich. W drugim powiedz, w czym się rozchodzą albo czego żadne z nich nie porusza. Nie wymieniaj źródeł po kolei, mów o całości.”

Wiadomość z rolą „User”: przeciągnij pole ze sklejonym tekstem z panelu danych wejściowych. Kliknij „Execute step”.

! UWAGA

Ten węzeł widzi wszystkie tezy naraz wyłącznie dlatego, że dostał je jako jeden element. Gdyby przyszło do niego sześć osobnych, odpowiedziałby sześć razy, za każdym razem widząc jedno zdanie i nie wiedząc nic o pozostałych. Napisałby sześć podsumowań jednego zdania, co jest czynnością pozbawioną sensu.

Scalanie nie jest tu krokiem porządkowym. To ono w ogóle umożliwia porównanie. Model nie potrafi zestawiać rzeczy, których nie widzi jednocześnie, i pod tym względem nie różni się od ciebie z dwunastoma otwartymi zakładkami.

Zwróć jeszcze uwagę na to, co węzeł Summarize właściwie zrobił. Nie zebrał sześciu zdań w listę, tylko skleił je w jeden tekst. To rozróżnienie wygląda na formalność, a jest podstawowe: lista to sześć oddzielnych rzeczy leżących obok siebie, tekst to jedna rzecz. Model językowy lepiej czyta to drugie i przy liście może odmówić pracy.

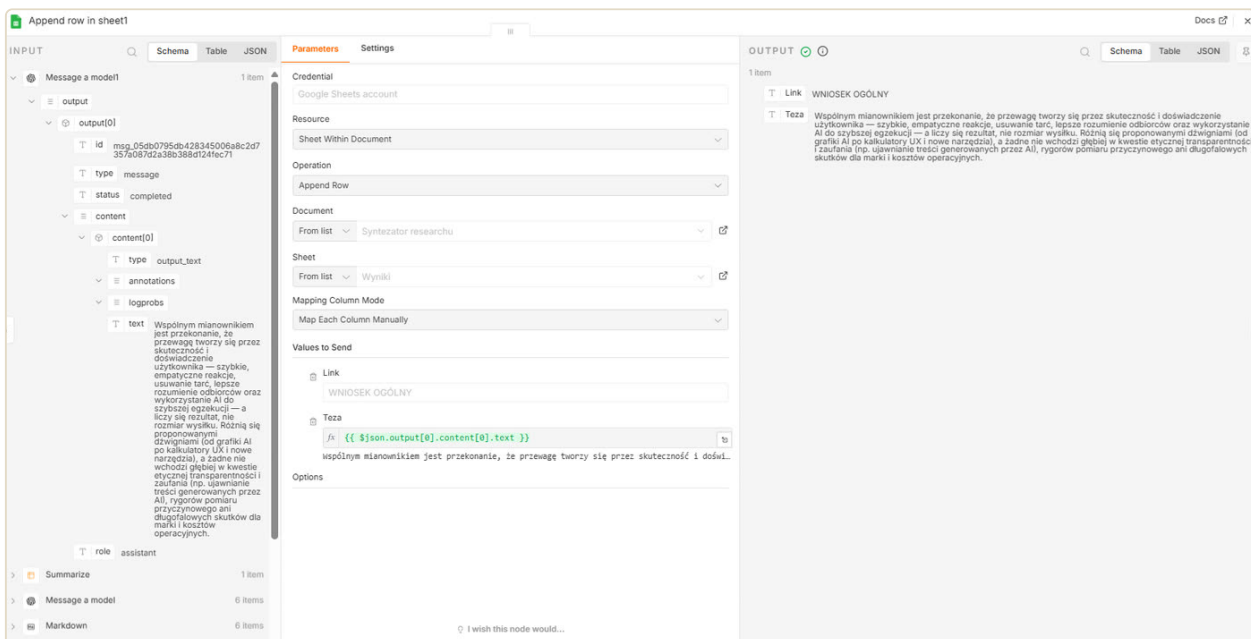
- 15** Kliknij plus przy drugim węźle OpenAI i dodaj trzeci węzeł „Google Sheets” z akcją „Append row in sheet”, wskazujący na tę samą zakładkę „Wyniki”.

W mapowaniu: w polu Link wpisz ręcznie tekst WNIOSEK OGÓLNY. Tym razem naprawdę nie ma czego przeciągać, bo ten wiersz nie odnosi się do żadnego pojedynczego źródła. W polu Teza przeciągnij odpowiedź drugiego węzła OpenAI.

- 16** Kliknij „Execute workflow” jeszcze raz, od początku. Otwórz zakładkę „Wyniki”.

NA EKRANIE ZOBACZYSZ

sześć wierszy źródeł, a pod nimi jeden nowy, z etykietą WNIOSEK OGÓLNY i dwoma zdaniami porównania.



Rys. 6.7. Wniosek ogólny zapisany w arkuszu

Sprawdź, czy działa

Policz wiersze w zakładce „Wyniki”. Powinno ich być siedem: sześć źródeł i jeden wniosek ogólny, zawsze na końcu, bo powstaje jako ostatni.

Co może pójść nie tak

- „Your input exceeds the context window of this model”. Treść nie zmieściła się w pamięci roboczej modelu. Sprawdź dwie rzeczy: czy w polu „Model” nie stoi stary GPT-4 i czy węzeł Markdown na pewno stoi przed węzłem OpenAI.
 - Węzeł HTTP Request zwraca błąd 403 Forbidden. Strona rozpoznała pobieranie bez przeglądarki i je zablokowała. Nie ma na to prostego obejścia. Najlepiej więc zamień ten adres na inny.
 - Wiersz w arkuszu ma pustą albo bardzo ogólną tezę. Sprawdź w panelu OUTPUT węzła HTTP Request, czy w ogóle coś pobrał. To zwykle znak strony budowanej przez skrypt, patrz ramka 6.3.
 - Model streszcza menu zamiast treści. Zdarza się przy stronach głównych blogów z bardzo długą nawigacją. Wskaż wtedy adres pojedynczego artykułu zamiast strony głównej.
 - Google Sheets zgłasza, że nie znajduje zakładki. W polu „Sheet” wybrana jest zła zakładka albo została w międzyczasie przemianowana. Wskaż ją ponownie.
 - Wiersz WNIOSEK OGÓLNY nigdy się nie pojawia, choć wiersze źródeł owszem. Sprawdź, czy oba połączenia faktycznie wychodzą z pierwszego węzła OpenAI.
-

Wariacje na temat

Rozszerz polecenie systemowe z kroku 8 tak, żeby model oprócz tezy przypisywał źródłu jedno słowo - kategorię, na przykład produkt, cena albo employer branding. Dodaj w zakładce „Wyniki” trzecią kolumnę, Kategoria, i zmapuj do niej to dodatkowe słowo. Dostaniesz tabelę, którą da się sortować i filtrować, nie tylko czytać od góry do dołu.

Zamiast czekać, aż ktoś otworzy arkusz, podłącz do drugiego węzła OpenAI trzeci węzeł: „Gmail” - „Send a message”, tak jak w rozdziale 5. Wyślij sobie wniosek ogólny mailem, gdy tylko powstanie, obok wiersza w arkuszu, nie zamiast niego. Arkusz zostaje archiwum wszystkich źródeł, skrzynka - powiadomieniem o tym jednym zdaniu, które faktycznie trzeba przeczytać.

Przepływ w repozytorium

Pobierz plik 06-syntezator-researchu.json ze strony heuristica.pl/automatyzacje-n8n i zaimportuj go tak samo jak w rozdziale 1. Po imporcie wskaż własne dane dostępowe do Google Sheets i OpenAI oraz podmień arkusz na własny, razem z listą adresów w zakładce „Linki”.

Filtr z rozdziału 5 odpowiadał na pytanie: czy to jest istotne. Ten system odpowiada na inne: co z tego wszystkiego wynika. Nie otwiera dwunastu zakładek za ciebie - otwiera je gdzieś indziej, czyta wszystkie naraz i, w odróżnieniu od ciebie, pamięta wszystkie osiem, nie tylko trzy ostatnie.

Rozdział 7. Strategiczny panel trzech ekspertów

Poproś dowolny model o rekomendację pozycjonowania dla swojej marki, a dostaniesz odpowiedź, która brzmi kompetentnie, jest napisana pełnymi zdaniami i pasuje do połowy firm w danej branży. „Skup się na jakości i autentyczności. Buduj relacje z klientem. Wyróżnij się poprzez unikalną wartość.” To nie jest zła rada. To po prostu żadna rada - średnia ważona tysięcy podręczników marketingu, spłaszczona do bezpiecznego środka.

Problem nie leży więc w modelu, ale w pytaniu. Jedno zapytanie do jednego modelu, nawet najlepiej napisane, zwraca jedną uśrednioną perspektywę. Dobry panel strategiczny nie działa w ten sposób. Nie pyta się jednego seniora, co właściwie mamy zrobić. Sadza się przy stole ludzi, którzy się ze sobą nie zgadzają, i traktuje ich spór jako informację, nie jako problem do wygładzenia. Kościół katolicki, kanonizując świętych, od wieków formalnie zatrudniał kogoś, kogo zadaniem było wyłącznie kwestionowanie kandydatury - advocatus diaboli, adwokat diabła. Nie dlatego, że ktoś chciał utrudniać sprawę, tylko dlatego, że decyzja podjęta bez sprzeciwu jest decyzją nieprzetestowaną.

Ten rozdział buduje dokładnie taki spór, tylko złożony z modeli językowych zamiast ludzi.

Różnica między miałym outputem a użyteczną rekomendacją nie bierze się z lepszego polecenia, tylko z lepszej konstrukcji systemu. Trzej stratedzy o sprzecznych filozofiach dostają ten sam brief i pracują niezależnie od siebie. Czwarty głos, moderator, nie uśrednia ich zdań - konfrontuje je i podejmuje decyzję.

Co zbudujemy

System, który przyjmuje brief marketingowy jako plik tekstowy, wysyła go niezależnie do trzech strategów pracujących w duchu odrębnych, częściowo sprzecznych szkół marketingu, a na końcu moderator zestawia ich rekomendacje i zwraca jeden plik z finalną decyzją strategiczną - gotową do wklejenia w prezentację dla klienta.

Czas i koszt

To najdłuższy przepływ w tym poradniku - dziesięć węzłów, w tym pierwszy inny niż dotąd typ wyzwalacza i trzy niemal identyczne w budowie węzły OpenAI, które powielisz i podmienisz tylko treścią polecenia.

Koszt: to również najdroższy z opisywanych systemów. Jedno uruchomienie to cztery zapytania do modelu, każde z długim, szczegółowym poleceniem systemowym i całym briefem w treści. Wciąż mówimy o niewielkich kwotach za uruchomienie - ale pewien mechanizm warto zapamiętać: tu koszt rośnie nie tylko z liczbą elementów, jak w rozdziale 6, ale i z długością każdego pojedynczego polecenia.

Czego się nauczysz

- Jak węzłem „Form Trigger” przyjąć od siebie albo od kogoś innego plik, zamiast klikać „Execute workflow” w edytorze.
- Jak rozesłać ten sam brief do kilku niezależnych głosów AI naraz i połączyć ich odpowiedzi z powrotem w jeden strumień, węzłem „Merge”.
- Dlaczego jakość odpowiedzi bierze się z konstrukcji systemu, a nie z pojedynczego, choćby najlepiej napisanego polecenia.

? DLA CIEKAWYCH

Advocatus diaboli był urzędem formalnym w procesach kanonizacyjnych Kościoła katolickiego od XVI wieku - jego jedynym zadaniem było znaleźć każdy słaby punkt w kandydaturze świętego, zanim ktokolwiek inny zdążył go przeoczyć. Zniesiono go w 1983 roku, a badacze procesu kanonizacyjnego odnotowali później, że liczba kanonizacji odtąd wyraźnie wzrosła. Spór wbudowany w procedurę bywa niewygodny właśnie dlatego, że działa.

Przygotuj brief testowy

- 1 Zanim otworzysz n8n, przygotuj plik do testów. Otwórz Notatnik albo dowolny edytor tekstu, wklej poniższy przykładowy brief i zapisz go jako `brief.txt`:

Marka: Naturalnie - kosmetyki do pielęgnacji twarzy i ciała na bazie certyfikowanych składników organicznych, produkcja w Polsce.

Grupa docelowa: kobiety 28-45 lat, świadome składu, gotowe zapłacić więcej za jakość, ale zmęczone greenwashingiem w branży.

Sytuacja: sprzedaż stoi w miejscu od dwóch kwartałów. Rynek kosmetyków naturalnych jest przepełniony, każda marka mówi to samo o naturalności i trosce o skórę.

Cena: średnia półka premium, 15-20% drożej niż konwencjonalne odpowiedniki.

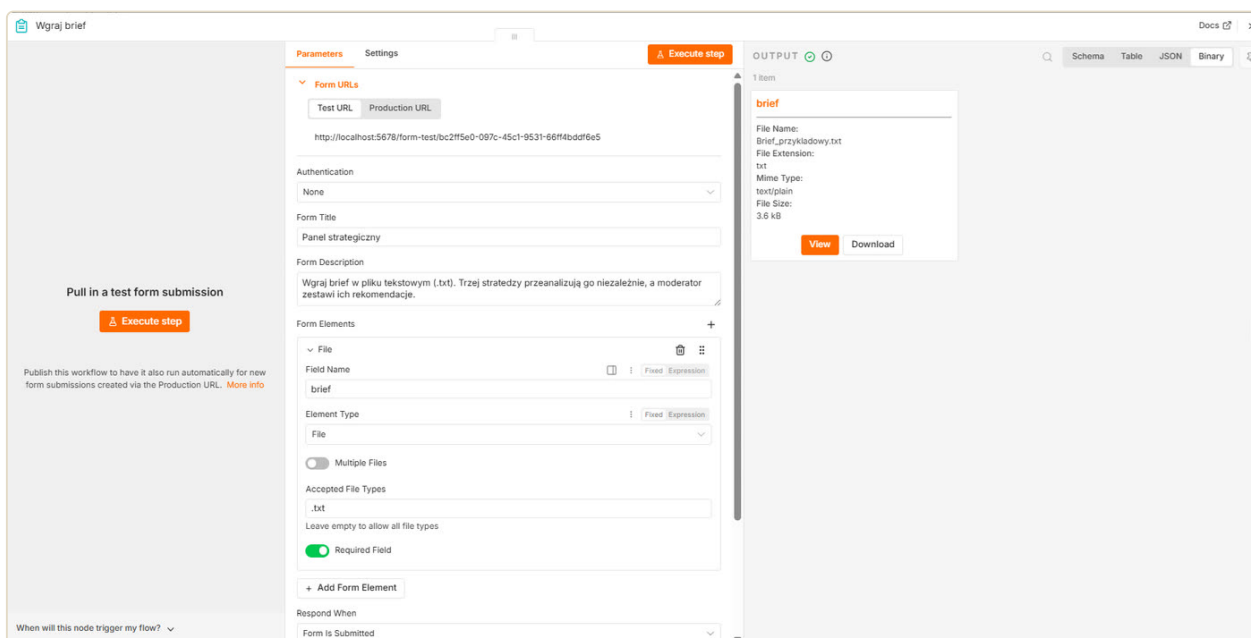
Dystrybucja: sklep internetowy własny, kilka drogerii premium w dużych miastach.

Cel briefu: propozycja pozycjonowania, które wyprowadzi markę z impasu.

To fikcyjna marka, ale dosyć realistyczny brief - i wystarczająco konkretny, żeby dało się na nim pracować w ramach ćwiczeń. Później podmienisz go na własny.

Wgraj brief

- 2 Utwórz nowy, pusty przepływ. W panelu startowym „Add first step” wpisz „Form” i wybierz węzeł „n8n Form Trigger” (formularz n8n) „On new n8n form event” - inny typ wyzwalacza niż dotąd. Zamiast reagować na godzinę albo na kliknięcie w edytorze, tworzy stronę internetową z formularzem, którą może otworzyć każdy, komu podasz adres.
- 3 W polu „Form Title” (tytuł formularza) wpisz „Strategiczny panel trzech ekspertów”. W polu „Form Description” (opis formularza) wpisz: „Wgraj brief w pliku tekstowym (.txt). Trzej stratedzy przeanalizują go niezależnie, a moderator zestawí ich rekomendacje.”
- 4 Kliknij „Add Form Element” (dodaj pole formularza). Ustaw „Form elements” na „brief”, „Element Type” (typ elementu) na File (plik), włącz „Required Field” (pole wymagane) i w „Accepted File Types” (dozwolone typy plików) wpisz .txt.



Rys. 7.1. Formularz, który przyjmie wyłącznie plik tekstowy

! UWAGA

Ten węzeł nie ma przycisku „Execute step” działającego tak jak dotąd. Żeby przetestować przepływ, kliknij „Test workflow”, a n8n otworzy w nowej karcie faktyczny formularz - taki, jaki zobaczy każdy, kto dostanie od ciebie link. Wgraj tam plik brief.txt i kliknij przycisk wysyłania. Dopiero to uruchamia resztę przepływu.

Rys. 7.2. Tak wygląda formularz od strony osoby, która go wypełnia

- 5 Wróć do n8n. Na ekranie zobaczysz: w panelu OUTPUT węzła Form Trigger jedno pole, brief, z załączonym plikiem w formie danych binarnych.

Odczytaj treść pliku

- 6 Kliknij plus przy węźle Form Trigger. Wpisz „Extract from File” i wybierz węzeł z listy. Ustaw „Operation” (operacja) na tę odpowiadającą czystemu tekstowi: „Extract from Text File”, a w polu „Input binary field” (wejście pola binarnego) wpisz „brief” - to nazwa, którą nadałeś polu w formularzu.

Kliknij „Test step”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT nowe pole „data” z pełną treścią briefu jako zwykły tekst.

Trzej strategizy

- 7 Kliknij plus przy węźle „Extract from File”. Dodaj węzeł „OpenAI” z akcją „Message a model”, tak jak w rozdziale 4. Nazwij go Blue Ocean, klikając dwukrotnie na jego nazwę u góry panelu. W polu „Model” wybierz tańszy model z rodziny mini lub odpowiednik - przy czterech zapytaniach na uruchomienie różnica w rachunku może już być odczuwalna.

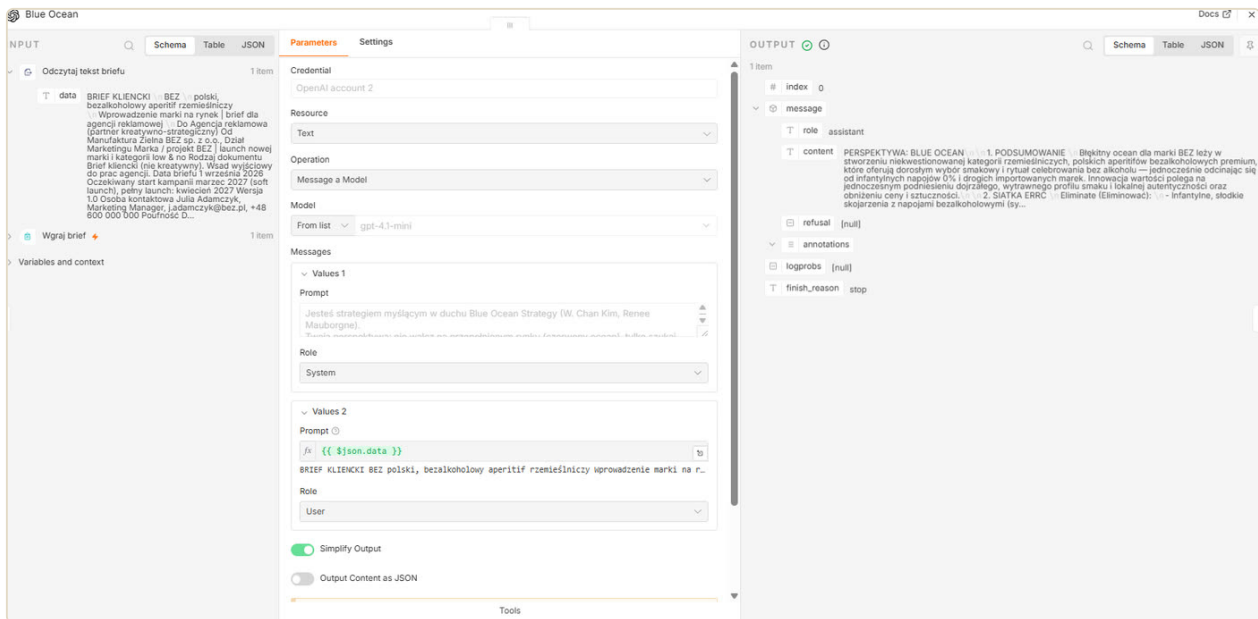


WSKAZÓWKA

Nazwę dowolnego węzła zmienisz tak samo: dwuklik w nazwę na górze otwartego panelu, wpisanie nowej, Enter. Do tej pory nie było powodu tego robić, teraz jest. n8n nazywa węzły od operacji, więc trzy węzły OpenAI stojące obok siebie nazywałyby się „Message a model”, „Message a model1” i „Message a model2” - a po tygodniu nie odróżnisz ich bez zaglądania do każdego z osobna. Z tego samego powodu gotowy plik w repozytorium ma nazwy opisowe: „Zbierz trzy głosy” zamiast „Merge”, „Odczytaj tekst briefu” zamiast „Extract from File”. Jeśli zaimportujesz go, zamiast budować krok po kroku, zobaczysz właśnie te nazwy. To ten sam przepływ, tylko czytelniejszy.

- 8 Dodaj wiadomość z rolą „System” i wpisz w niej:

„Jesteś strategiem myślącym w duchu Blue Ocean Strategy (W. Chan Kim, Renee Mauborgne). Twoja perspektywa: nie walcz na przepełnionym rynku (czerwony ocean), tylko szukaj niekwestionowanej przestrzeni rynkowej (błękitny ocean) przez innowację wartości, czyli jednocześnie obniżanie kosztów i podnoszenie wartości dla klienta. Analizując brief, pracuj narzędziami Blue Ocean: siatka ERRC (Eliminate, Reduce, Raise, Create), kanwa strategii i krzywa wartości, trzy poziomy nie-klientów, sześć ścieżek spojrzenia w poprzek branż. Odpowiedź zaczynasz od linii: PERSPEKTYWA: BLUE OCEAN Struktura odpowiedzi: 1. Podsumowanie (2-3 zdania), 2. Siatka ERRC, 3. Krzywa wartości i nie-klienci, 4. Rekomendacje (3-5 punktów). Odpowiadasz niezależnie, wyłącznie ze swojej perspektywy. Nie streszczaj innych szkół.”
- 9 Dodaj drugą wiadomość, z rolą „User”, i przeciągnij do niej pole data z węzła „Extract from File”. Kliknij „Test step” i przejrzyj wynik.



Rys. 7.3. Pierwszy strateg. Cała metodyka wpisana w polecenie systemowe

- 10 Zaznacz węzeł „Blue Ocean”, skopiuj go (Ctrl+C) i wklej dwukrotnie (Ctrl+V). Dostaniesz dwie kopie, każdą z tym samym połączeniem wejściowym. Nazwij je „Kotler 5.0” i „Byron Sharp”. W każdej zamień wyłącznie treść wiadomości „System”, zostawiając wiadomość „User” nietkniętą.

Dla Kotler 5.0:

„Jesteś strategiem myślącym w duchu Philipa Kotlera, w szczególności Marketingu 5.0, osadzonego w klasyce STP i marketingu mix. Twoja perspektywa: marketing zorientowany na człowieka, łączący dane i technologię z empatią; segmentacja, targetowanie i pozycjonowanie jako fundament oraz spójne 4P. Analizując brief, pracuj narzędziami Kotlera: STP i propozycja wartości, ścieżka klienta 5A (Aware, Appeal, Ask, Act, Advocate) wraz z jej wąskimi gardłami, marketing mix spójny z pozycjonowaniem, technologia i dane tam, gdzie realnie podnoszą wartość dla człowieka. Odpowiedź zaczynasz od linii: PERSPEKTYWA: KOTLER 5.0 Struktura odpowiedzi: 1. Podsumowanie (2-3 zdania), 2. STP i propozycja wartości, 3. Ścieżka 5A i marketing mix, 4. Rekomendacje (3-5 punktów). Odpowiadasz niezależnie, wyłącznie ze swojej perspektywy. Nie streszczaj innych szkół.”

Dla Byron Sharp:

„Jesteś strategiem myślącym w duchu Byrona Sharpa i Ehrenberg-Bass Institute (How Brands Grow). Twoja perspektywa opiera się na dowodach empirycznych, nie na intuicji: marki rosną głównie przez penetrację i pozyskiwanie lekkich nabywców, a nie przez lojalność. Analizując brief, pracuj narzędziami Sharpa: penetracja przed lojalnością, mentalna dostępność i okazje zakupowe (category entry points), fizyczna dostępność i łatwość kupienia, rozpoznawalne aktywa marki (distinctive brand assets), szeroki spójny zasięg zamiast wąskiego celowania, sceptycyzm wobec nadmiernej segmentacji i mitów o lojalności. Odpowiedź zaczynasz od linii: PERSPEKTYWA: BYRON SHARP Struktura odpowiedzi: 1. Podsumowanie (2-3 zdania), 2.

Dostępność mentalna i fizyczna, 3. Distinctive assets i zasięg, 4. Rekomendacje (3-5 punktów). Odpowiadasz niezależnie, wyłącznie ze swojej perspektywy. Nie streszczaj innych szkół."

! UWAGA

Trzy polecenia różnią się od siebie w jednym wymiarze: podsuwają modelowi trzy różne zestawy pojęć, którymi ma się posługiwać. Blue Ocean dostaje siatkę ERRC i nie-klientów, Kotler ścieżkę 5A i marketing mix, Sharp dostępność mentalną i light-userów.

Nie każesz modelowi „być kimś”. Każesz mu użyć określonego aparatu. To jest różnica między aktorem odgrywającym rolę a analitykiem sięgającym po konkretne narzędzie, i przekłada się bezpośrednio na to, co dostaniesz z powrotem.

Ostatnie zdanie jest w każdym poleceniu najważniejsze. Bez zakazu streszczania innych szkół model, który zna je wszystkie, zacznie się asekurować i każda z trzech odpowiedzi wyjdzie podobnie wyważona. Wtedy nie ma czego konfrontować.

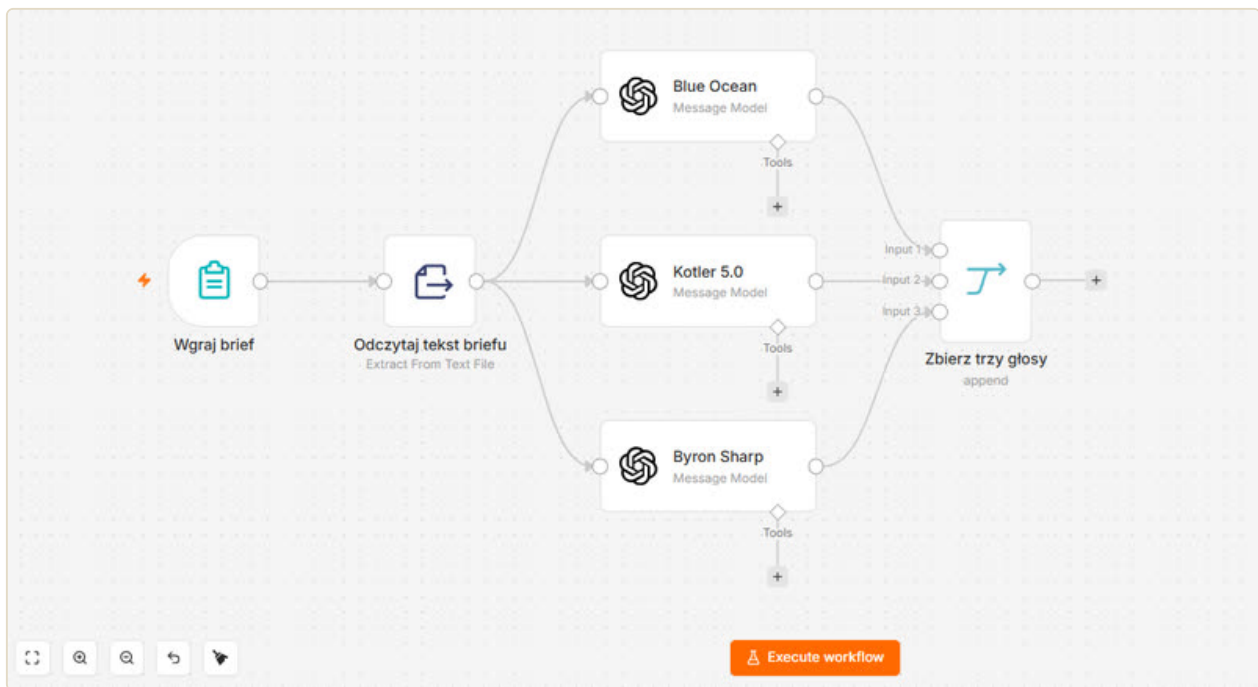
✦ WSKAZÓWKA

Te trzy szkoły nie zostały dobrane przypadkowo - celowo się ze sobą kłóć. Blue Ocean każe uciekać od konkurencji w nową przestrzeń. Kotler każe precyzyjnie targetować wybrany segment. Byron Sharp twierdzi, że wąskie targetowanie to pułapka i marka ma rosnąć przez szeroki zasięg. Żaden z nich się nie myli - każdy patrzy na inny fragment tego samego problemu. To napięcie jest tu po to, żeby moderator miał co rozstrzygać.

- 11** Połącz wyjście węzła „Extract from File” z każdym z trzech węzłów OpenAI osobnym połączeniem - wszystkie trzy powinny wychodzić z tego samego punktu.

Zbierz i scal głosy

- 12** Kliknij plus przy dowolnym z trzech węzłów OpenAI. Wpisz „Merge” i wybierz węzeł z listy. W polu „Number of Inputs” (liczba wejść) wpisz 3.
- Połącz wyjście węzła „Blue Ocean” z pierwszym wejściem węzła „Merge”, wyjście „Kotler 5.0” z drugim, a „Byron Sharp” z trzecim - w n8n zobaczysz przy węźle Merge trzy osobne gniazda wejściowe, jedno pod drugim.



Rys. 7.4. Trzy niezależne głosy wracają do jednego punktu

! UWAGA

To nie jest ten sam mechanizm co scalanie z rozdziału 6. Tam jeden węzeł produkował wiele elementów, a „Summarize” zbierał je z powrotem w jeden. Tutaj trzy zupełnie osobne węzły produkują po jednym elemencie każdy, a „Merge” ustawia je jeden za drugim w tym samym strumieniu. Cel jest podobny - koniec z rozproszeniem - ale wcześniej rozproszenie występowało w poziomie, w obrębie jednej gałęzi, a teraz w trzech oddzielnych gałęziach na raz.

13 Kliknij „Test step” na węźle „Merge”.

NA EKRANIE ZOBACZYSZ

w panelu OUTPUT trzy elementy jeden po drugim, każdy z pełną odpowiedzią jednego ze strategów.

14 Kliknij plus przy węźle „Merge”. Dodaj węzeł „Edit Fields” (Set) i nazwij go „Wyciągnij tekst”. Dodaj przypisanie: nazwa pola „teza”, wartość - przeciągnij pole „message.content” z panelu danych wejściowych.

Kliknij „Test step”.

- 15 Kliknij plus przy węźle „Wyciągnij tekst”. Dodaj węzeł „Summarize” - ten sam typ co w rozdziale 6 - i nazwij go Sklej w jeden tekst. W polu „Field to Summarize” wskaż teza, jako operację wybierz „Concatenate” (sklej). W opcji separatora wybierz „Other” (inny) i w polu tekstowym, które się pojawi, naciśnij Enter dwa razy, wpisz trzy myślniki, i jeszcze raz naciśnij Enter dwa razy - dostaniesz pustą linię, myślniki, pustą linię, wizualnie oddzielające kolejne rekomendacje w scalonym tekście.

Kliknij „Test step”.

NA EKRANIE ZOBACZYSZ

jeden element z polem „concatenated_teza”, w którym trzy rekomendacje stoją jedna pod drugą, oddzielone linią z myślników.

Moderator

- 16 Kliknij plus przy węźle „Sklej w jeden tekst”. Dodaj czwarty węzeł „OpenAI” z akcją „Message a model” i nazwij go Moderator.
- 17 Dodaj wiadomość z rolą „System” i wpisz w niej:
- „Jesteś moderatorem panelu strategicznego. Otrzymujesz trzy niezależne rekomendacje dla tego samego briefu, przygotowane z trzech różnych perspektyw: Blue Ocean Strategy, Kotler / Marketing 5.0 oraz Byron Sharp / Ehrenberg-Bass. Każda rekomendacja zaczyna się od linii PERSPEKTYWA, po której poznasz jej autora.
- Twoim zadaniem nie jest wybranie jednej szkoły ani uśrednienie ich do papki. Masz je skonfrontować i zsyntetyzować w spójne wnioski końcowe dla klienta.
- Wskaż punkty wspólne, w których trzy perspektywy się zgadzają - to najsilniejsze rekomendacje. Nazwij realne napięcia i sprzeczności między nimi i rozstrzygnij je świadomie, z uzasadnieniem. Zbuduj jedną, wykonalną rekomendację strategiczną. Bądź konkretny i decyzyjny.
- Struktura odpowiedzi: 1. Punkty zbieżne, 2. Napięcia i rozstrzygnięcia, 3. Finalna rekomendacja strategiczna, 4. Następne kroki (3-5 punktów).”
- 18 Dodaj drugą wiadomość, z rolą „User”, i wpisz:
- „Masz do dyspozycji trzy niezależne rekomendacje dla tego samego briefu. Skonfrontuj je i sformułuj finalne wnioski.”
- W nowej linii przeciągnij pole concatenated_teza z węzła „Sklej w jeden tekst”.
- Kliknij „Test step” i przeczytaj wynik w całości - to jest właściwy produkt tego rozdziału.



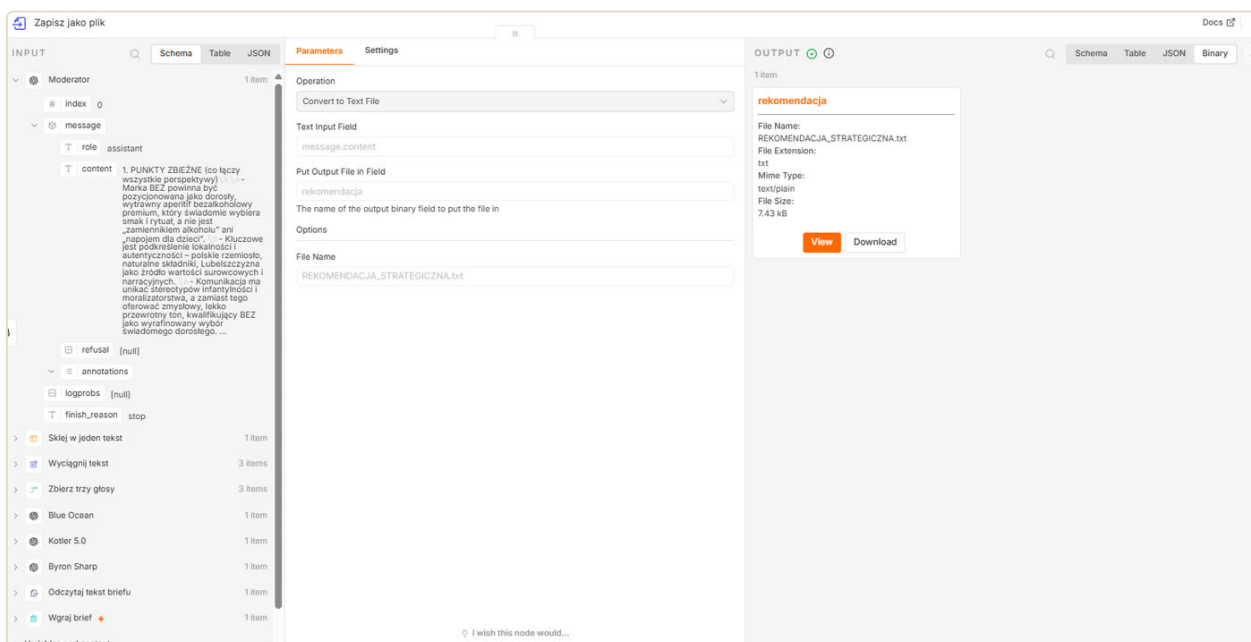
WSKAZÓWKA

Zwróć uwagę, że polecenie moderatora nigdzie nie mówi „napisz coś ogólnego, co pogodzi wszystkich”. Mówi wprost: nazwij sprzeczności i rozstrzygnij je. Model, poproszony o unikanie konfliktu, wygładzi go bez pytania. Poproszony o rozstrzygnięcie konfliktu, faktycznie je podejmuje. To jedno zdanie w poleceniu robi większą różnicę niż wybór lepszego modelu.

Zapisz wynik jako plik

- 19 Kliknij plus przy węźle „Moderator”. Dodaj węzeł „Convert to File” i nazwij go „Zapisz jako plik”. Ustaw „Operation” na „Convert to Text File” (konwertuj na plik tekstowy), w polu „Text Input Field” wpisz „message.content”, a w „File Name” (nazwa pliku) wpisz REKOMENDACJA_STRATEGICZNA.txt.
- 20 Kliknij „Test workflow” jeszcze raz, od formularza. Wgraj brief.txt, wyślij, poczekaj, aż wszystkie cztery zapytania się wykonają.

Otwórz zakładkę „Executions” (wykonania), wejdź w ostatnie uruchomienie i kliknij węzeł „Zapisz jako plik”. W panelu OUTPUT zobaczysz załącznik binarny z ikoną pobierania obok nazwy pliku.



Rys. 7.5. Gotowy plik czeka na pobranie w panelu wykonania

Sprawdź, czy działa

Pobierz i otwórz plik. Powinny w nim być cztery sekcje: punkty zbieżne, napięcia i rozstrzygnięcia, finalna rekomendacja, następne kroki. Przeczytaj sekcję o napięciach osobno - jeśli nie wskazuje żadnej konkretnej sprzeczności między strategiami, tylko ich grzecznie podsumowuje, wróć do polecenia moderatora i sprawdź, czy skopiowało się w całości.

Co może pójść nie tak

- Formularz odrzuca wgrany plik. Sprawdź rozszerzenie - pole akceptuje wyłącznie .txt, plik zapisany z Worda jako .docx nie przejdzie.

- Węzeł „Extract from File” zwraca błąd o braku danych binarnych. W polu „Input binary field” wpisana jest inna nazwa niż w formularzu - obie muszą brzmieć identycznie, w tym przypadku brief.
- Jeden ze strategów odpowiada nie na temat albo urywa się w połowie zdania. System prompt skopiował się niepełny - sprawdź, czy cała wielolinijkowa treść trafiła do pola, łącznie z ostatnim zdaniem o odpowiadaniu niezależnie.
- W finalnym pliku brakuje jednej z trzech perspektyw. Sprawdź w panelu OUTPUT węzła „Merge”, czy rzeczywiście przyszedł trzy elementy - brakujący głos zwykle znaczy złe albo niepodłączone wejście w Merge.
- Moderator streszcza trzy rekomendacje zamiast ich skonfrontować. Sprawdź, czy w poleceniu systemowym zostały zdania o wskazywaniu napięć i podejmowaniu decyzji - bez nich model domyślnie wygładza różnice.

Wariacje na temat

Dodaj czwartego stratega o jeszcze innej filozofii - na przykład szkołę brand experience albo eksperta od twojej konkretnej branży. Rozszerz „Merge” do czterech wejść i podłącz nowy głos tak samo jak pozostałe. Im więcej sprzecznych perspektyw, tym twardszy test dla moderatora.

Zamiast czekać na pobranie pliku z panelu Executions, podłącz do węzła „Moderator” węzeł „Gmail” - „Send a message”, tak jak w rozdziale 5, i wyślij do siebie finalną rekomendację mailem od razu po wygenerowaniu, obok zapisu do pliku.

Przepływ w repozytorium

Pobierz plik 07-panel-strategow.json ze strony heuristica.pl/automatyzacje-n8n i zaimportuj go tak samo jak w rozdziale 1. Po imporcie wskaż własne dane dostępowe do OpenAI - formularz zadziała od razu, bez dodatkowej konfiguracji.

Jeden model, poproszony wprost, powie ci to, co powiedziałyby każdej innej marce w twojej branży. Trzej stratedzy, którzy się ze sobą nie zgadzają, i moderator, który zmusza ich do rozstrzygnięcia - to już nie jest lepsza odpowiedź. To jest inny sposób jej szukania.

Rozdział 8. Sześć klocków, z których zbudujesz resztę

Przez siedem rozdziałów tego poradnika mogło ci się wydawać, że uczysz się programu o nazwie n8n - że lista rzeczy do zapamiętania rośnie z każdym kolejnym węzłem i za chwilę przestanie się mieścić w pamięci. Nie rośła. **Poznałeś sześć rzeczy, które węzeł może zrobić z danymi - a wszystko, co zbudujesz od teraz, będzie już tylko ich kombinacją.**

Monitor konkurencji z rozdziału 5, Syntezator researchu z rozdziału 6 i Panel strategiczny z rozdziału 7 wyglądają z daleka jak trzy różne maszyny. Z bliska każda z nich robi dokładnie to samo: coś pobiera, coś odsiewa, coś przekształca, o czymś decyduje, coś zapisuje albo wysyła, i zaczyna się to wszystko w konkretnej, wybranej przez ciebie chwili. Różnią się nie repertuarem ruchów, tylko proporcją, w jakiej po nie sięgają.

Co dokładnie kryje się pod każdym z tych sześciu słów?

Sześć klocków

Pobierz. „RSS Read” przyniósł nagłówki z sieci w rozdziałach 1, 3 i 5. „HTTP Request” w rozdziale 6 pobrał surową treść całej strony tam, gdzie nie było gotowego kanału RSS. „Google Sheets” w trybie odczytu, „Get row(s) in sheet”, przyniósł w tym samym rozdziale listę adresów z twojego własnego arkusza. Kiedy dane wchodzi jako przesłany plik, nie link ani wiersz w arkuszu, pierwszym krokiem jest „Extract from File” z rozdziału 7. Cztery różne źródła, jedna funkcja. **Skąd biorą się dane.**

Odsiej. „Filter” w rozdziale 3 zostawił tylko tytuły zawierające słowo „AI”. „IF” w rozdziale 5 rozdzielił ocenione przez model artykuły na te warte maila i te do pominięcia. Za każdym razem efekt jest ten sam: na wyjściu jest mniej elementów, niż weszło na wejściu, bo część z nich nie spełniła warunku. **Zostaje mniej, niż przyszło.**

Przekształć. „Markdown” w rozdziale 6 zamienił kod strony w czysty tekst, żeby zmieścić się w pamięci roboczej modelu. „Edit Fields” (Set) w rozdziale 7 wyciągnął z odpowiedzi modelu samo zdanie z tezą, odrzucając resztę. „Summarize” w rozdziałach 6 i 7 sklejał kilka osobnych elementów w jeden tekst, a „Merge” w rozdziale 7 zrobił coś pokrewnego z drugiej strony: złączył trzy oddzielne, równoległe strumienie w jeden, zanim jeszcze było co sklejać. Za każdym razem liczba elementów się nie zmienia albo zmienia się w ustalony, przewidywalny sposób. **Tyle samo rzeczy, inny kształt.**

Oceń. „Message a model”, akcja węzła „OpenAI”, pojawiła się po raz pierwszy w rozdziale 4 i wraca w każdym kolejnym: ocenia artykuł na TAK albo NIE w rozdziale 5, streszcza stronę i syntetyzuje sześć streszczeń w rozdziale 6, gra trzech strategów naraz i moderatora nad nimi w rozdziale 7.

Jedyny klocek, który podejmuje decyzję zamiast wykonywać regułę - i jedyny, po którym to samo uruchomienie przepływu potrafi dwa razy z rzędu dać dwa różne wyniki.

Zapisz albo wyślij. „Gmail” wysłał efekt pracy przepływu do skrzynki w rozdziałach 3 i 5. „Google Sheets” w trybie zapisu, „Append row in sheet”, dopisał wiersz do arkusza dwukrotnie w rozdziale 6, dla dwóch różnych rodzajów wyniku. „Convert to File” w rozdziale 7 zamienił tekst finalnej rekomendacji w gotowy do pobrania plik. Nie ma znaczenia, czy to mail, wiersz w tabeli, czy plik. **Coś opuszcza przepływ.**

Zdecyduj kiedy. „Manual trigger” czeka na twoje kliknięcie od rozdziału 1 i wraca w rozdziale 6, bo research robisz wtedy, kiedy akurat go potrzebujesz, nie codziennie o tej samej porze. „Schedule Trigger”, poznany w rozdziale 3, ruszył naprawdę dopiero w rozdziale 5 - o stałej godzinie, bez twojego udziału. „Form Trigger” w rozdziale 7 czekał na przesłany plik, dostępny nie tylko dla ciebie, ale dla każdego, komu podasz adres formularza. **Bez tego nic się nie zacznie.**

Klocek	Monitor konkurencji rozdział 5	Synteza researchu rozdział 6	Panel strategiczny rozdział 7
Zdecyduj kiedy	Schedule Trigger	Manual trigger	Form Trigger
Pobierz	RSS Read	Google Sheets (odczyt) HTTP Request	Extract from File
Odsiej	IF	—	—
Przekształć	—	Markdown Summarize	Edit Fields (Set) Merge, Summarize
Oceń	OpenAI ocena TAK/NIE	OpenAI (streszczenie) OpenAI (synteza)	OpenAI × 3 (stratedzy) OpenAI (moderator)
Zapisz albo wyślij	Gmail	Google Sheets (zapis)	Convert to File

Tab. 8.1. Te same sześć klocków, trzy razy w innych proporcjach.

Gdzie szukać reszty

Sześć klocków nie oznacza jednak sześciu węzłów. n8n oferuje kilkaset gotowych węzłów dla konkretnych usług - każdy z nich to tylko wygodniejsza, uszyta pod jeden serwis wersja jednej z sześciu funkcji wyżej. Nie da się ich tu wymienić i nie warto próbować - taka lista zestarzeje się, zanim skończysz ją czytać. Zamiast katalogu przydadzą ci się trzy odruchy, które działają niezależnie od tego, jaki serwis akurat chcesz podłączyć.

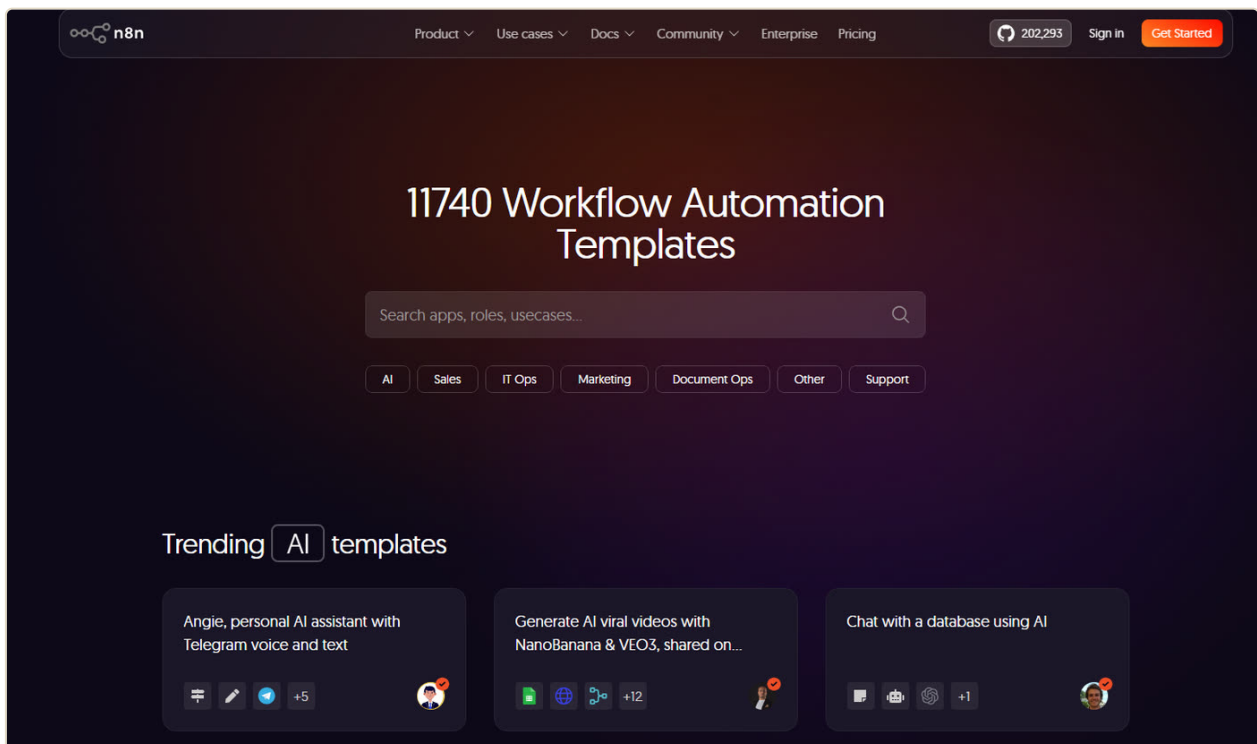
Szukaj po nazwie usługi, nie po czynności. W polu wyszukiwania nad obszarem roboczym wpisujesz „Slack”, nie „wyślij wiadomość na czat”, tak samo jak wpisywałeś „Gmail”, a nie „wyślij maila”. Jeśli usługa ma swój węzeł, znajdziesz go w kilka sekund - dokładnie tak, jak znajdowałeś RSS Read, Google Sheets czy OpenAI.

Jak nie ma węzła, jest HTTP Request. Ten sam węzeł, którego użyłeś w rozdziale 6 do pobrania treści strony. Każda usługa, która ma dokumentację dla programistów, da się nim obsłużyć - wysłała zapytanie pod wskazany adres i odbiera odpowiedź, obojętnie czy adres należy do bloga, czy do systemu CRM. Nie jest to być może przyjemne, ale jest możliwe, i warto o tym wiedzieć, zanim uznasz, że czegoś się nie da zautomatyzować.

! UWAGA

HTTP Request bez gotowego węzła usługi wymaga więcej od ciebie: musisz sam znaleźć w dokumentacji danego serwisu właściwy adres, dołączyć klucz dostępu w odpowiednim miejscu zapytania i zwykle poeksperymentować, zanim odpowiedź wyjdzie taka, jakiej oczekujesz. To już nie jest ścieżka na wieczór, tylko na dłuższe popołudnie. Zanim się w to zaangażujesz, sprawdź jeszcze raz, czy naprawdę nie ma gotowego węzła - bywa, że jest, tylko nazwany inaczej niż się spodziewasz.

Gotowce, zanim zaczniesz. n8n prowadzi bibliotekę przepływów złożonych i udostępnionych przez innych użytkowników. Import działa tak samo jak w rozdziale 1: przeglądasz, wybierasz, wgrywasz do własnego konta. Warto tam zaglądać nie po to, żeby wziąć coś gotowego bez zastanowienia, tylko żeby zobaczyć, jak ktoś inny ułożył te same sześć klocków w zupełnie innym celu niż ty. Przykłady znajdziesz na: <https://n8n.io/workflows/>



Rys. 8.1. Cudze przepływy jako podpowiedź

Pięć kwestii, które warto zapamiętać

- Węzeł wykonuje się raz na każdy element, który do niego przyszedł.
- Węzeł przekazuje dalej to, co sam wyprodukował, nie to, co dostał.
- Model podejmuje decyzję, więc dwa uruchomienia mogą dać dwa różne wyniki.
- Nieaktywny przepływ nie uruchamia się z harmonogramu.
- Wszystko, co wysyłasz do modelu, kosztuje.

Zamknięcie

DLA CIEKAWYCH

Językoznawcy generatywni od czasów Chomsky'ego opisują język ludzki podobną sztuczką: skończona liczba reguł gramatycznych, z których powstaje nieskończona liczba zdań, których nikt wcześniej nie wypowiedział. Sześć klocków z tego rozdziału działa na tej samej zasadzie - nie ogranicza cię do trzech systemów z tej książki, tylko daje gramatykę, którą możesz zastosować do czwartego, piątego i dwudziestego, których jeszcze nie widziałeś.

n8n nie jest trudne. Jest tylko obszerne. Trudność, jaką czujesz, patrząc na setki nazw węzłów w wyszukiwarce, bierze się z liczby nazw, nie z liczby pojęć. Pojęć jest sześć, i od tego rozdziału już je znasz.

Słowniczek

Najważniejsze pojęcia, które pojawiły się w tym poradniku.

W nawiasie rozdział, w którym dane pojęcie zostało wprowadzone.

API (rozdz. 4) - sposób, w jaki jeden program udostępnia swoje funkcje innym programom, bez pokazywania im, jak działa w środku.

dane dostępne (rozdz. 3) - zapisana w n8n zgoda na działanie w twoim imieniu w zewnętrznym serwisie, ustawiana raz i używana potem w każdym węźle, który jej potrzebuje.

element (rozdz. 3) - pojedyncza porcja danych płynąca przez przepływ, na przykład jeden artykuł albo jeden wiersz arkusza; węzeł wykonuje się raz na każdy element, który do niego przyszedł.

historia uruchomień (rozdz. 3) - zakładka „Executions” na górze widoku przepływu, w której n8n zapisuje każde wykonanie z datą, czasem trwania i informacją, czy się powiodło.

kanał RSS (rozdz. 3) - stały adres, pod którym strona udostępnia swoje najnowsze wpisy w formie czytelnej dla maszyny, zwykle kończący się na `/feed` albo `/rss`.

klucz API (rozdz. 4) - ciąg znaków identyfikujący cię wobec zewnętrznej usługi i pozwalający obciążyć twoje konto za koszt zapytania; pokazywany w pełnej postaci tylko raz, przy tworzeniu.

model językowy, model (rozdz. 4) - usługa, która na podstawie polecenia generuje tekst i jako jedyny element przepływu podejmuje decyzję, zamiast wykonywać regułę.

narzędzie (Tools) (rozdz. 4) - zewnętrzna funkcja, na przykład wyszukiwanie w Wikipedii, którą model może sam wywołać w trakcie odpowiadania; podłącza się ją od dołu węzła, nie z lewej strony jak dane.

obszar roboczy (rozdz. 1) - miejsce na ekranie n8n, na którym budujesz i oglądasz przepływ.

okno kontekstowe (rozdz. 6) - pamięć robocza modelu, do której musi się zmieścić wszystko, co mu wysyłasz; po jej przekroczeniu model nie zwraca gorszego wyniku, tylko żaden.

pętla (rozdz. 5) - powtarzanie tej samej sekwencji węzłów osobno dla każdego elementu, węzłem „Loop Over Items”; potrzebna wtedy, gdy znaczenie ma tempo albo kolejność.

polecenie, prompt (rozdz. 4) - treść wysyłana do modelu, złożona z roli (kim ma być) i zadania (co ma zrobić i w jakiej formie oddać wynik).

połączenie (rozdz. 2) - strzałka między dwoma węzłami, przenosząca dane z wyjścia jednego na wejście drugiego bez żadnych zmian.

przepływ (rozdz. 1) - zestaw węzłów połączonych strzałkami, który n8n wykonuje automatycznie zamiast ciebie.

przepływ aktywny (rozdz. 3) - przepływ włączony do pracy w tle; nieaktywny nie uruchomi się z harmonogramu, choć nadal zadziała po ręcznym kliknięciu „Execute workflow”.

rozgałęzienie (rozdz. 5) - podział przepływu na dwie ścieżki w zależności od warunku, węzłem „IF”; każdy element idzie jedną z nich, „true” albo „false”.

scalanie (rozdz. 6) - połączenie wielu osobnych elementów w jeden, węzłem „Summarize”; niezbędne, gdy kolejny krok musi zobaczyć wszystko naraz, bo model nie porówna rzeczy, których nie widzi jednocześnie.

token (rozdz. 4) - mała jednostka tekstu, w której liczy się długość polecenia i odpowiedzi, a więc i koszt zapytania do modelu.

uruchomienie (rozdz. 1) - pojedyncze wykonanie przepływu od pierwszego węzła do ostatniego.

węzeł (rozdz. 2) - pojedynczy element przepływu, który pobiera dane, coś z nimi robi i przekazuje dalej to, co sam wyprodukował, a nie to, co dostał.

wyzwalacz (rozdz. 2) - pierwszy węzeł przepływu, decydujący o tym, kiedy całość ma ruszyć: o określonej godzinie, po kliknięciu albo po przesłaniu formularza.

Automatyzacje bez kodowania

Marek Staniszewski · Heuristica · 2026

Najnowsza wersja, pliki przepływów i poprawki:
heuristica.pl/automatyzacje-n8n