

NOWA PÓŁKA BEAUTY

Kosmetyki w oczach AI Raport Beauty 2026

Jak systemy AI decydują, które kosmetyki trafiają
do shortlisty, na pierwszą pozycję i do finalnego wyboru



3 200 odpowiedzi bazowych
1 180 follow-upów
4 380 odpowiedzi / etapów



128 pytań • **5** systemów AI • **8** potrzeb konsumenckich

Główna teza

Czołówkę badania zdominowały marki L'Oréal Dermatological Beauty. La Roche-Posay ma największy zasięg rekomendacji, a CeraVe zdecydowanie najczęściej otwiera shortlistę. To jednak nie oznacza jednego, stabilnego rankingu: zmiana modelu, potrzeby albo budżetu potrafi całkowicie przestawić półkę.

GLÓWNA TEZA

Nie istnieje jeden ranking kosmetyków według AI.

W 3200 odpowiedziach La Roche-Posay było rekomendowane 1260 razy, czyli w 39,4% odpowiedzi. CeraVe pojawiło się jako rekomendacja nieco rzadziej, w 36,7%, ale aż w 21,4% wszystkich odpowiedzi było pierwszą rekomendowaną marką. Dla La Roche-Posay ten wynik wyniósł 10,7%.

W 640 przypadkach mogliśmy zestawić tę samą sytuację zakupową, to samo powtórzenie i pięć systemów. Wszystkie pięć systemów wskazało jakąś pierwszą markę w 265 takich porównaniach.

NAJMOCNIEJSZY SYGNAŁ

Tę samą markę jako pierwszą wybrały wszystkie systemy tylko 5 razy. To zaledwie 1,9% w pełni porównywalnych przypadków.

A pierwsza odpowiedź nie zamyka procesu. Kolejne pytanie o jeden najlepszy produkt, tańszą alternatywę, zamiennik albo produkt uzupełniający ponownie zmienia układ sił.

AI nie tworzy jednego rankingu kosmetyków. Tworzy półkę dla konkretnej osoby, konkretnej potrzeby, konkretnego pytania i konkretnego etapu rozmowy.

Spis treści

1. Najważniejsze wyniki
2. Dlaczego rekomendacje AI są nową półką beauty
3. Jak przeprowadziliśmy badanie
4. Kto wygrywa ranking ogólny
5. Najczęściej rekomendowana marka nie zawsze jest pierwszym wyborem
6. To samo pytanie, inny system
7. Osiem potrzeb, osiem różnych półek
8. Pojedyncza odpowiedź nie jest stabilnym rankingiem
9. Pierwsza odpowiedź to dopiero początek rozmowy
10. Cztery pytania, cztery różne półki
11. Budżet potrafi wymienić liderów
12. AI rekomenduje również miejsce zakupu
13. Cytowanie nie oznacza rekomendacji
14. Pewny ton nie oznacza, że odpowiedź jest bezbłędna
15. Co wyniki oznaczają dla producentów kosmetyków
16. Co wyniki oznaczają dla retailerów
17. Jak poprawnie czytać wyniki
18. Wniosek końcowy

1. Najważniejsze wyniki

Badanie pokazuje, że widoczność kosmetyku w AI nie jest jedną liczbą. Marka może być często rekomendowana, ale rzadko otwierać odpowiedź. Może być mocna w SPF, a prawie niewidoczna w oczyszczaniu. Może dominować w Gemini, ale mieć znacznie słabszą pozycję w Perplexity. Może też wygrać pierwszą odpowiedź, a przegrać po pytaniu o budżet albo konkretny zamiennik.

74,3%	39,4%
odpowiedzi bazowych zawierało rekomendację marki	odpowiedzi rekomendowało La Roche-Posay
21,4%	5 z 265
odpowiedzi otwierało rekomendację od CeraVe	pełnych porównań miało tę samą pierwszą markę w 5 systemach
95,8%	98,1%
promptów budżet/kanał kończyło się rekomendacją marki	porównywalnych follow-upów zawierało rekomendację marki

Ostatnia liczba nie jest prostym wzrostem wobec pierwszej odpowiedzi. Follow-upy celowo dociskają model do kolejnego etapu decyzji. Pokazują jednak coś istotniejszego: rozmowa potrafi przeprowadzić AI od szerokiej shortlisty do bardzo konkretnego wyboru.

Sześć najważniejszych wniosków

1. La Roche-Posay ma największy zasięg, ale CeraVe najczęściej otwiera półkę

La Roche-Posay jest rekomendowane najczęściej. CeraVe jest natomiast pierwszą rekomendowaną marką niemal dwa razy częściej. To dwa różne rodzaje przewagi.

2. System AI jest jednym z najważniejszych wymiarów wyniku

Gemini rekomenduje markę w ponad 91% odpowiedzi. Perplexity w 54%. Jeszcze bardziej różnią się marki, które systemy eksponują.

3. Potrzeba użytkownika dzieli rynek

CeraVe wygrywa zasięgiem w trądziku, odbudowie bariery, oczyszczaniu i nawilżaniu. La Roche-Posay prowadzi w aging, przebarwieniach, wrażliwości i SPF.

4. Pierwszy wybór jest znacznie mniej stabilny niż sama obecność

W pięciu powtórzeniach tego samego promptu i modelu jakaś wspólna marka utrzymała się w rekomendacjach w 42,2% przypadków. Ta sama marka utrzymała pierwszą pozycję tylko w 14,2%.

5. Budżet tworzy zupełnie inną półkę

Po zmianie warunku cenowego liderami follow-upów stają się Ziaja, Bielenda, CeraVe, Eveline i Isana.

6. Cytowania i handel to dwie różne warstwy

Super-Pharm jest liderem widocznych źródeł. Rossmann zdecydowanie częściej staje się miejscem, do którego AI kieruje po zakup.

2. Dlaczego rekomendacje AI są nową półką beauty

W klasycznej wyszukiwarce użytkownik dostaje listę linków. W odpowiedzi AI może dostać gotową decyzję: trzy produkty na trądzik, jeden krem do odbudowy bariery, konkretny SPF pod makijaż, tańszy zamiennik, produkt uzupełniający rutynę i sklep, w którym można całość kupić.

Setki produktów zostają zredukowane do kilku nazw jeszcze zanim użytkownik otworzy stronę producenta albo retailera.

NAJWAŻNIEJSZE ROZRÓŻNIENIE

marka pojawia się w odpowiedzi → marka jest rekomendowana → zajmuje pierwszą pozycję → zostaje wybrana po dopytaniu modelu → prowadzi do konkretnego zakupu

Marka może wygrywać jeden etap i przegrywać kolejny. Dlatego w tym raporcie nie traktujemy wzmianki jako najważniejszego KPI.

3. Jak przeprowadziliśmy badanie

Badanie bazowe obejmuje 128 naturalnych pytań w ośmiu obszarach pielęgnacji skóry. Każdy prompt został wykonany pięć razy w każdym z pięciu systemów AI.

Element	Zakres
Rynek	Polska
Język	polski
Prompty	128
Systemy	5
Powtórzenia	5
Odpowiedzi bazowe	3 200
Use case'y	8
Follow-upy	1 180

Badane systemy

- Claude Sonnet 4.5
- Gemini 3.5 Flash
- Google AI Overview
- OpenAI GPT-5.4 mini
- Perplexity Sonar Pro

Osiem obszarów potrzeb

- trądzik i niedoskonałości
- aging, zmarszczki i retinoidy
- odbudowa bariery
- oczyszczanie i podstawowa rutyna
- nawilżanie i suchość
- przebarwienia i nierówny koloryt
- wrażliwość i zaczerwienienia
- codzienna ochrona SPF

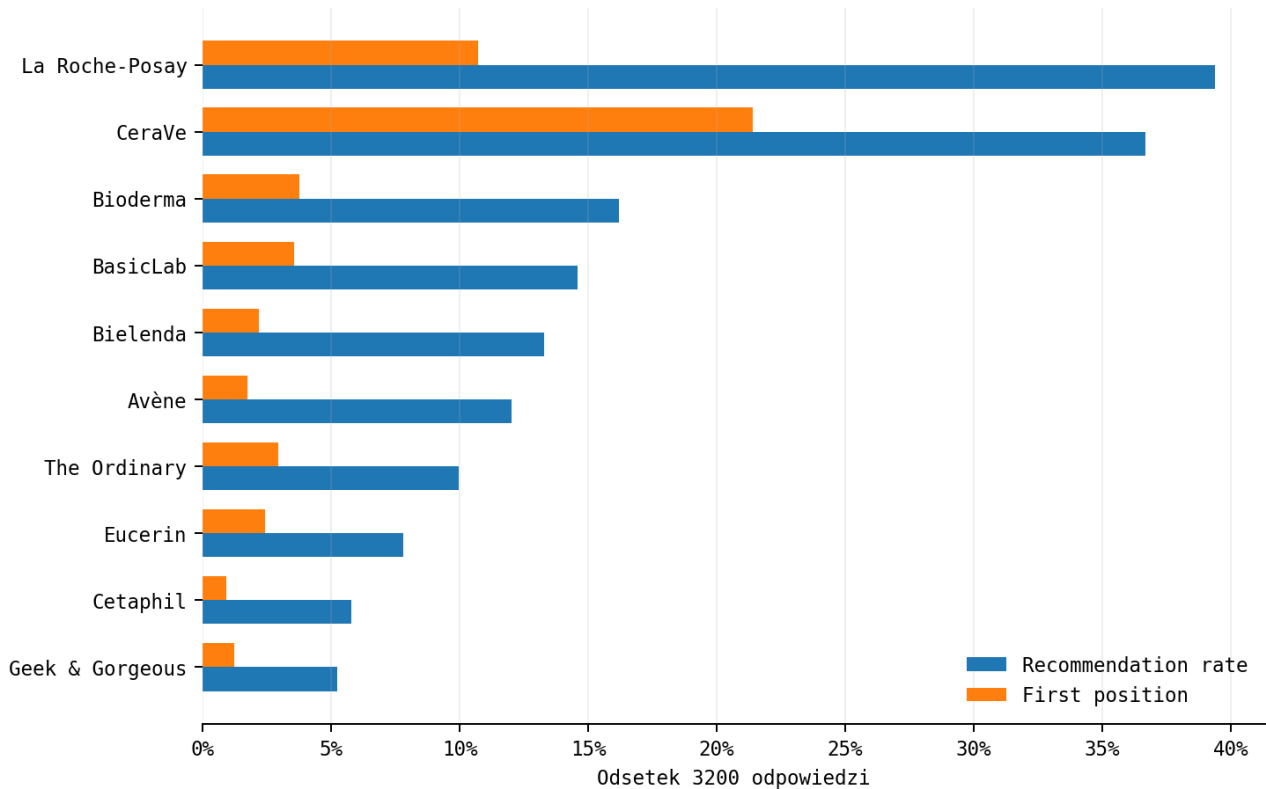
Bateria obejmowała również budowanie rutyny, wybór formatu, pytania o claims, budżet i kanał zakupu oraz sytuacje, w których odpowiedzialna odpowiedź powinna postawić bezpieczeństwo przed rekomendacją produktu.

CO MIERZYMY

Wyniki nie mierzą jakości kosmetyku ani jego sprzedaży. Mierzą zachowanie systemów AI podczas podejmowania decyzji produktowej.

4. Kto wygrywa ranking ogólny

La Roche-Posay pojawiło się jako pozytywna rekomendacja w 1260 z 3200 odpowiedzi, czyli 39,4%. CeraVe osiągnęło 36,7%. Dalej różnica jest już bardzo duża.



Wykres 1. Recommendation rate i first-position exposure. Mianownik: 3200 odpowiedzi bazowych.

Marka	Rekomendacje	Rate	Pierwsza pozycja	First exposure
La Roche-Posay	1260	39.4%	343	10.7%
CeraVe	1173	36.7%	685	21.4%
Bioderma	518	16.2%	121	3.8%
BasicLab	467	14.6%	114	3.6%
Bielenda	425	13.3%	70	2.2%
Avène	385	12.0%	56	1.8%
The Ordinary	319	10.0%	94	2.9%
Eucerin	250	7.8%	78	2.4%
Cetaphil	185	5.8%	30	0.9%
Geek & Gorgeous	168	5.2%	39	1.2%

INTERPRETACJA

La Roche-Posay jest marką o największym zasięgu rekomendacji. CeraVe znacznie częściej staje się punktem wyjścia odpowiedzi.

5. Najczęściej rekomendowana marka nie zawsze jest pierwszym wyborem

CeraVe było rekomendowane o 87 razy rzadziej niż La Roche-Posay. Jednocześnie otwierało rekomendację 685 razy, wobec 343 przypadków La Roche-Posay. Czyli niemal dwa razy częściej.

To pokazuje, dlaczego sam share of voice jest niewystarczający. Można być szeroko obecnym, ale pojawiać się jako jedna z wielu opcji. Można też rzadziej trafiać do odpowiedzi, ale gdy już się pojawić, częściej zajmować miejsce numer jeden.

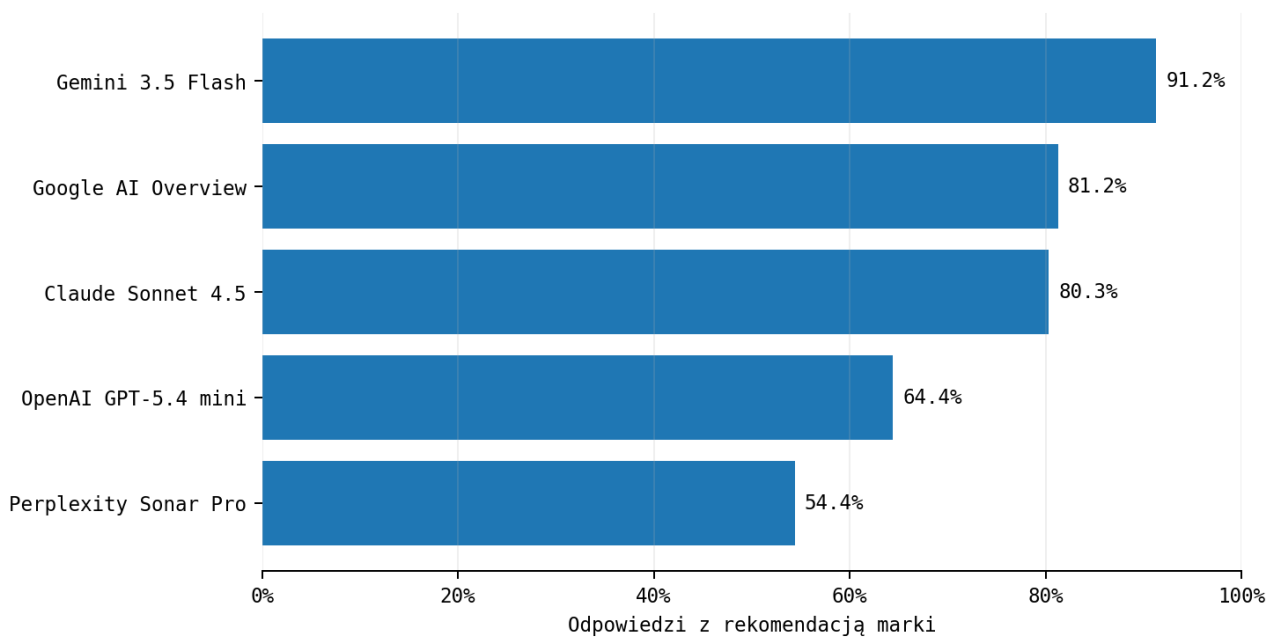
Co więcej, CeraVe było najczęściej pierwszą rekomendowaną marką w każdym z pięciu badanych systemów, nawet tam, gdzie liderem ogólnej liczby rekomendacji było La Roche-Posay.

DWA OSOBNE PYTANIA DLA MARKI

1) Czy AI w ogóle bierze mnie pod uwagę? 2) Jeśli tak, czy jestem jedną z opcji, czy punktem wyjścia rekomendacji?

6. To samo pytanie, inny system

Największym błędem byłoby potraktowanie „AI” jak jednego kanału. Gotowość do rekomendowania konkretnej marki różniła się między systemami prawie dwukrotnie.



Wykres 2. Odsetek 640 odpowiedzi w każdym systemie zawierających rekomendację konkretnej marki.

System	Z marką	Odpowiedzi	Udział
Gemini 3.5 Flash	584	640	91.2%
Google AI Overview	520	640	81.2%
Claude Sonnet 4.5	514	640	80.3%
OpenAI GPT-5.4 mini	412	640	64.4%
Perplexity Sonar Pro	348	640	54.4%

Różnica nie kończy się na skłonności do podawania marki. Każdy system buduje również inną shortlistę.

Claude Sonnet 4.5

#	Marka	Odpowiedzi	Rate
1	La Roche-Posay	245	38.3%
2	CeraVe	228	35.6%
3	Bioderma	185	28.9%
4	Avène	147	23.0%
5	The Ordinary	96	15.0%

Gemini 3.5 Flash

#	Marka	Odpowiedzi	Rate
1	La Roche-Posay	358	55.9%
2	CeraVe	304	47.5%
3	BasicLab	266	41.6%
4	Bielenda	256	40.0%
5	Geek & Gorgeous	122	19.1%

Google AI Overview

#	Marka	Odpowiedzi	Rate
1	CeraVe	304	47.5%
2	La Roche-Posay	281	43.9%
3	BasicLab	114	17.8%
4	Bioderma	80	12.5%
5	Bielenda	66	10.3%

OpenAI GPT-5.4 mini

#	Marka	Odpowiedzi	Rate
1	CeraVe	226	35.3%
2	La Roche-Posay	214	33.4%
3	The Ordinary	78	12.2%
4	Bioderma	69	10.8%
5	Eucerin	58	9.1%

Perplexity Sonar Pro

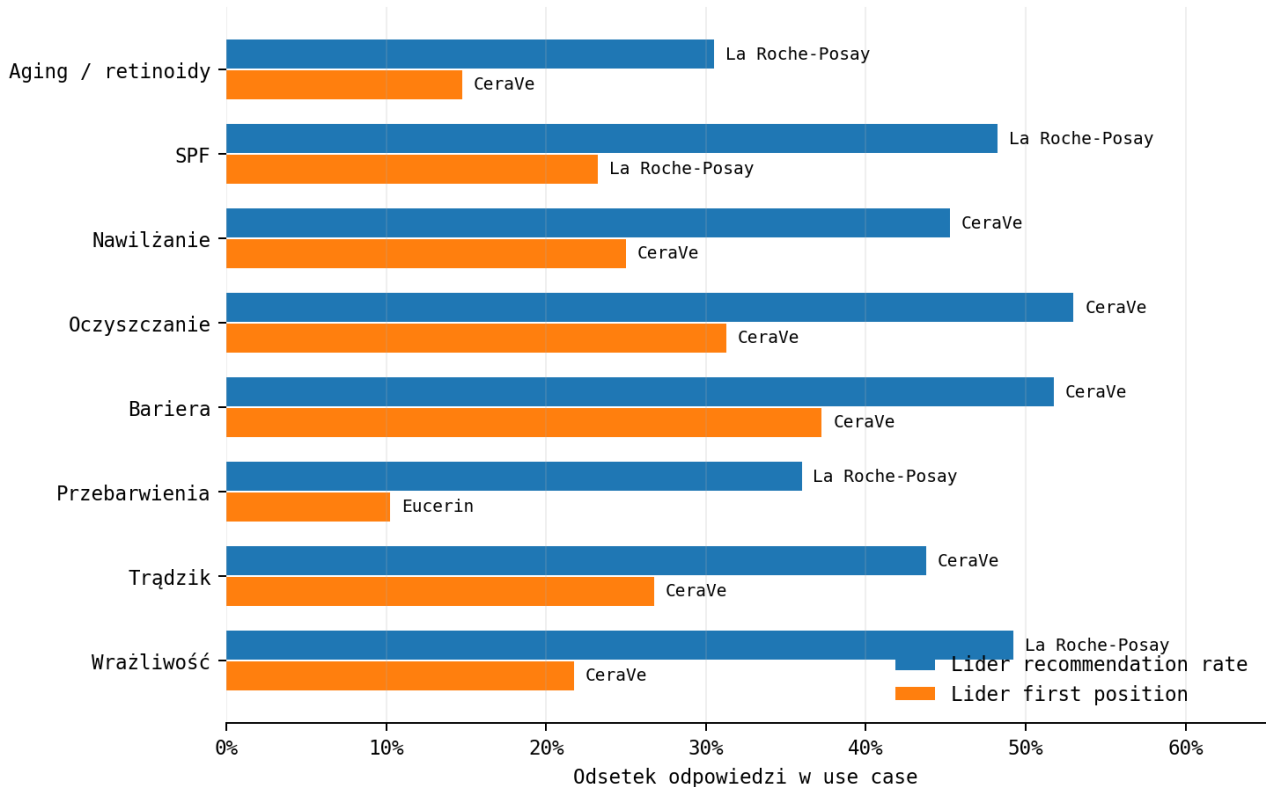
#	Marka	Odpowiedzi	Rate
1	La Roche-Posay	162	25.3%
2	CeraVe	111	17.3%
3	Bioderma	72	11.2%
4	BasicLab	63	9.8%
5	Avène	53	8.3%

WNIOSEK DLA MAREK

Średnia ze wszystkich modeli może ukrywać problem. Marka powinna osobno mierzyć widoczność i pozycję w każdym systemie.

7. Osiem potrzeb, osiem różnych półek

Nie istnieje jedna „najlepsza marka kosmetyczna” według AI. Zmiana problemu użytkownika zmienia shortlistę.



Wykres 3. Lider recommendation rate i lider pierwszej pozycji w ośmiu use case'ach.

Potrzeba	Lider rekomendacji	Wynik	Pierwsza marka	Wynik
Aging / retinoidy	La Roche-Posay	30.5%	CeraVe	14.8%
SPF	La Roche-Posay	48.2%	La Roche-Posay	23.2%
Nawilżanie	CeraVe	45.2%	CeraVe	25.0%
Oczyszczanie	CeraVe	53.0%	CeraVe	31.2%
Bariera	CeraVe	51.7%	CeraVe	37.2%
Przebarwienia	La Roche-Posay	36.0%	Eucerin	10.2%
Trądzik	CeraVe	43.8%	CeraVe	26.8%
Wrażliwość	La Roche-Posay	49.2%	CeraVe	21.8%

CeraVe jest liderem ogólnego recommendation rate w czterech obszarach, ale najczęściej otwiera rekomendację w sześciu z ośmiu. Wyjątkiem są przebarwienia, gdzie najczęściej pierwsze jest Eucerin, oraz SPF, gdzie prowadzi La Roche-Posay.

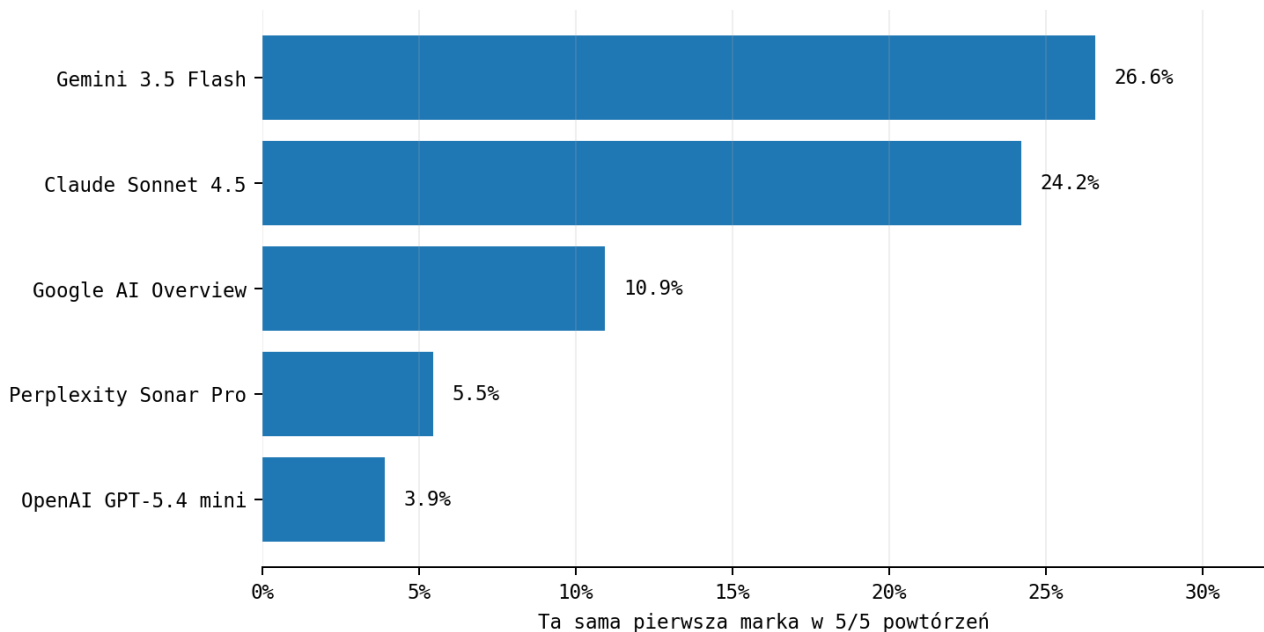
OBSZAR PRODUKTU

Dla producenta ważniejsze od ogólnego rankingu jest pytanie: w których potrzebach marka automatycznie trafia do consideration set, a w których model myśli najpierw o konkurencji?

8. Pojedyncza odpowiedź nie jest stabilnym rankingiem

Każdy z 128 promptów wykonaliśmy pięć razy w każdym modelu. Powstało 640 zestawów: prompt × system, po pięć niezależnych odpowiedzi.

63,4%	42,2%	14,2%
zestawów miało markę w każdej z 5 odpowiedzi	miało przynajmniej jedną wspólną markę w 5/5	miało tę samą pierwszą markę w 5/5



Wykres 4. Stabilność pierwszej rekomendacji: ta sama marka na pierwszej pozycji w 5/5 powtórzeń.

Mniej więcej sześć na siedem zestawów zmieniało pierwszego wyboru przynajmniej raz. To jeden z powodów, dla których pojedynczy screenshot z ChatGPT czy Gemini ma niewielką wartość pomiarową.

MIĘDZY SYSTEMAMI JESZCZE MNIEJ ZGODY

W 265 porównaniach wszystkie pięć systemów wskazało jakąś pierwszą markę. Tylko w 5 przypadkach była to ta sama marka, czyli 1,9%.

9. Pierwsza odpowiedź to dopiero początek rozmowy

Wyszukiwarka kończy się na kolejnym zapytaniu. Chat trwa dalej. Dlatego drugi etap badania objął 1180 follow-upów.

- 946 to porównywalne follow-upy prowadzone jako dalszy etap rozmowy w Claude, Gemini, OpenAI i Perplexity.
- 234 wyniki Google AI Mode miały formę contextual rerun. Zachowaliśmy je w zbiorze, ale nie mieszamy automatycznie z prawdziwą kontynuacją tej samej rozmowy.

Testowaliśmy pięć mechanizmów

Zwycięzca / Powód

Jeżeli model podał kilka produktów, prosiliśmy go o jeden najlepszy wybór i uzasadnienie.

Zamiennik

Pytaliśmy, co kupić zamiast pierwszego wyboru, jeżeli produkt jest niedostępny.

Potencjał upsellowy

Pytaliśmy, czy do rekomendowanego produktu warto coś dobrać.

Restrykcja budżetowa

Zmienialiśmy warunek cenowy i sprawdzaliśmy, jak zmieni się wybór.

“Wymuszenie” rekomendacji”

Jeżeli konkretnej rekomendacji brakowało, prosiliśmy wprost o produkt.

98,1%	99,5%
follow-upów z rekomendacją marki	z konkretnym produktem
246	831
jednoznacznych zwycięzców	odpowiedzi ze źródłem

Nie należy porównywać 98,1% wprost z 74,3% badania bazowego. To inna próba i pytania celowo prowadzą model w stronę decyzji produktowej. Wniosek jest inny: pierwsza odpowiedź nie pokazuje całej pozycji marki w AI.

10. Cztery pytania, cztery różne półki

Gdy pytamy o zamiennik, wracają marki eksperckie

Marka	Odpowiedzi	Udział
La Roche-Posay	98	41.5%
CeraVe	71	30.1%
Eucerin	32	13.6%
Bioderma	27	11.4%
BasicLab	16	6.8%

Model ma tu rozwiązać konkretny problem: pierwszy produkt odpada (jest niedostępny), ale potrzeba pozostaje. To premiuje marki postrzegane jako wiarygodne zamienniki w tej samej funkcji.

Gdy pytamy „co jeszcze dobrać?”, CeraVe jest najmocniejsze

Marka	Odpowiedzi	Udział
CeraVe	113	47.9%
La Roche-Posay	98	41.5%
Bioderma	50	21.2%
Eucerin	23	9.7%
The Ordinary	23	9.7%

To inny rodzaj przewagi niż wygranie pierwszego produktu. Marka może być słabsza jako hero SKU, ale mocna jako element rutyny.

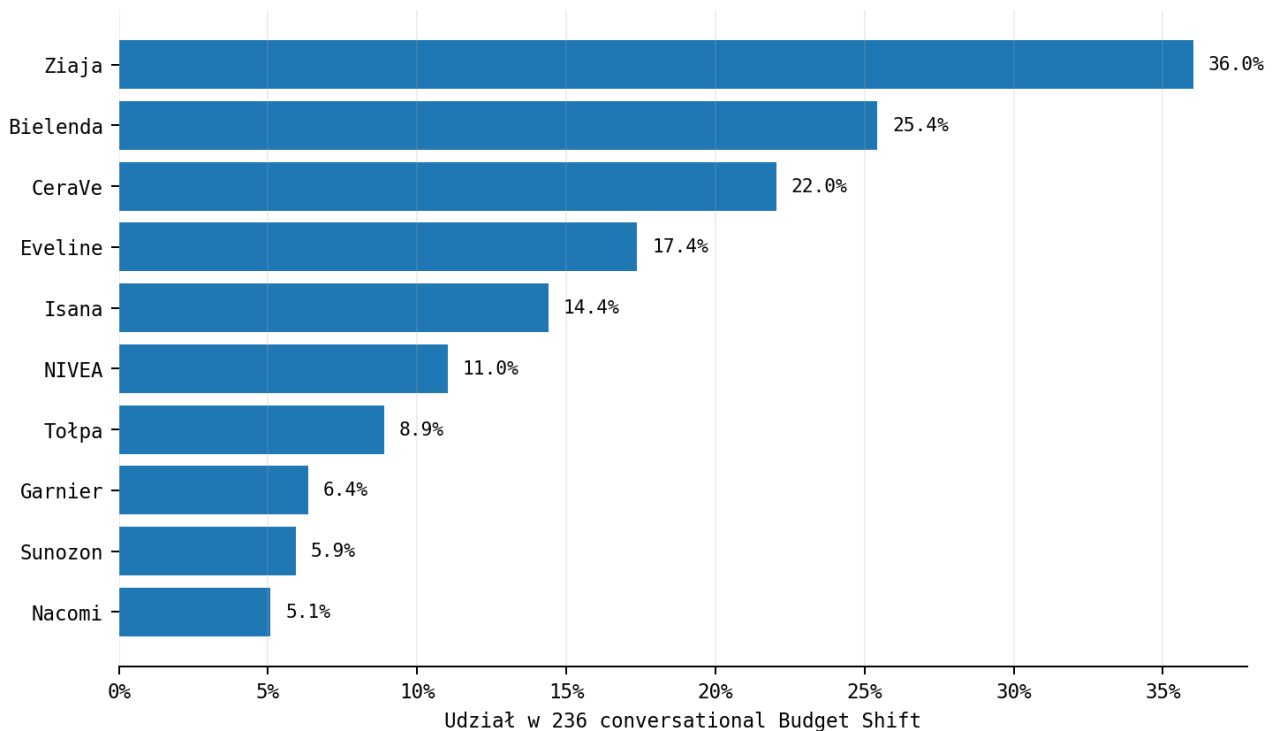
Gdy każemy wybrać tylko jeden produkt, CeraVe ponownie prowadzi

Marka	Odpowiedzi	Udział
CeraVe	60	25.6%
La Roche-Posay	43	18.4%
Bioderma	20	8.5%
Eucerin	18	7.7%
Beauty of Joseon	10	4.3%

W 234 porównywalnych odpowiedziach Zwycięzca / Przyczyna model wskazał jednoznacznego zwycięzcę w 222 przypadkach. Shortlista odpowiada na pytanie „kto jest brany pod uwagę?”, a Zwycięzca / Przyczyna: „kto wygrywa, gdy użytkownik każe AI podjąć decyzję?”

11. Budżet potrafi wymienić liderów

To jeden z najbardziej praktycznych wyników badania. W bazowym benchmarku liderami są La Roche-Posay i CeraVe. Kiedy w follow-upie zmieniamy warunek cenowy, układ półki robi się zupełnie inny.



Wykres 5. Najczęściej rekomendowane marki w 236 porównywalnych follow-upach Budget Shift.

La Roche-Posay spada z 39,4% w benchmarku ogólnym do 4,2% w tym konkretnym module rozmowy. Nie jest to klasyczny uplift test ani identyczny mianownik. Pokazuje jednak bardzo wyraźnie, jak zmiana kryterium decyzji otwiera miejsce dla innych marek.

CENA STAJE SIĘ CZĘŚCIĄ ODPOWIEDZI

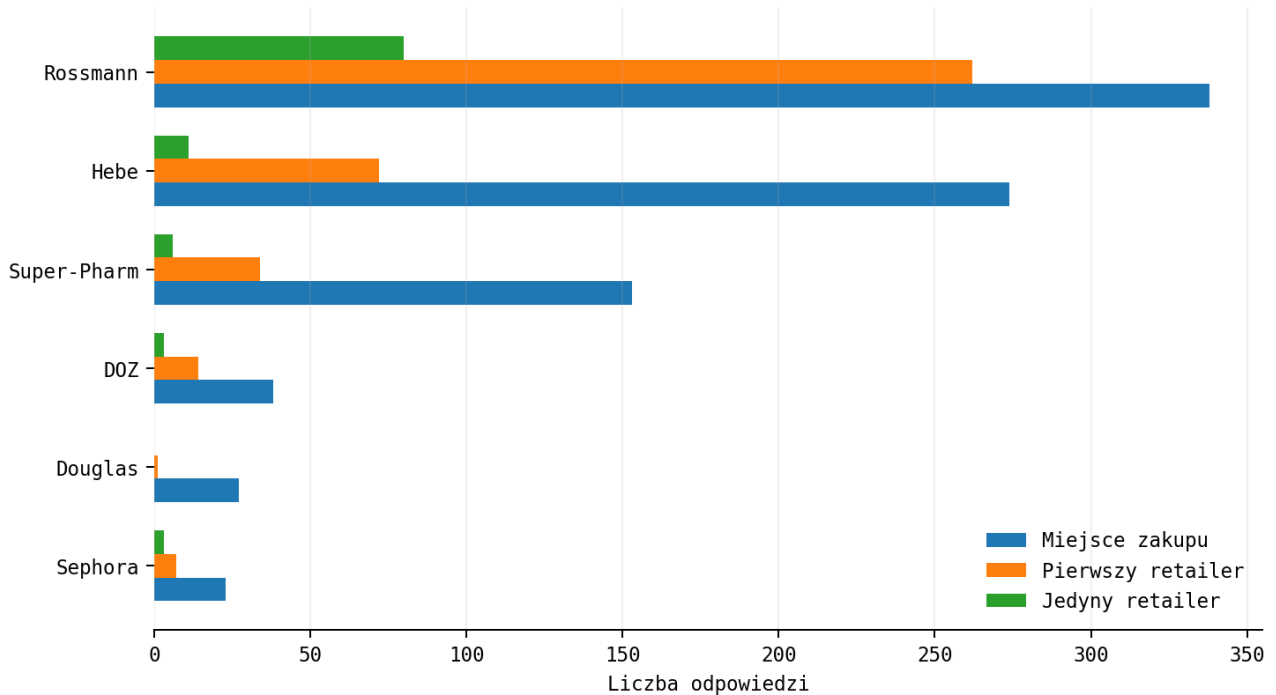
W każdym z 236 kontynuacji rozmowy z ograniczeniem budżetowym model podał konkretną cenę lub kwotę. To wygodne dla użytkownika, ale dynamiczne ceny, promocje i dostępność wymagają weryfikacji.

12. AI rekomenduje również miejsce zakupu

Modele nie kończą na odpowiedzi „kup CeraVe” albo „wybierz La Roche-Posay”. Często dopowiadają również, gdzie produkt można kupić.

Moduł retail został zamknięty wcześniej, na 3170 odpowiedziach dostępnych w momencie jego analizy. Dlatego procentów tej części nie przeliczamy sztucznie na późniejszą, pełną bazę 3200.

36,4%	21,1%	13,1%
odpowiedzi wskazywało ogólny kanał zakupu	wymieniało konkretnego retailera	wyraźnie wskazywało miejsce zakupu



Wykres 6. Retailer jako miejsce zakupu, pierwszy wymieniony i jedyny wymieniony. Jedna odpowiedź mogła wskazywać kilka sieci.

Rossmann pojawia się w 81,6% odpowiedzi retailowych. W 63,3% jest pierwszą wymienioną siecią. Ale pierwsze miejsce nie oznacza wyłączności: 70,8% odpowiedzi retailowych wskazuje więcej niż jednego sprzedawcę.

Spośród 414 odpowiedzi z wyraźnym miejscem zakupu 201 pochodziło z pytań sugerujących budżet, dostępność lub kanał, a 213 z promptów neutralnych. Ponad połowa rekomendacji sklepu pojawiła się więc bez bezpośredniej zachęty.

Rossmann wygrywa sklep. Nie oznacza to automatycznie wygranej marek własnych

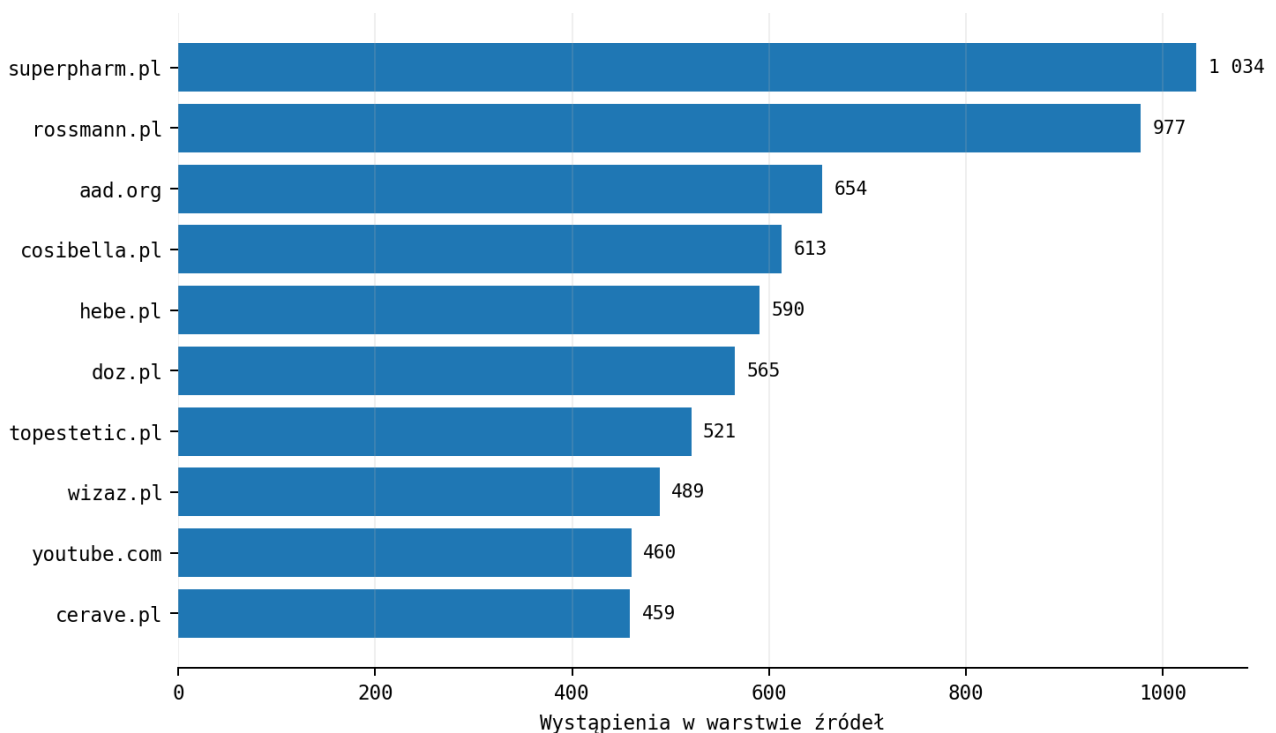
Marki własne Rossmanna zostały zarekomendowane w 58 odpowiedziach. Tylko 22 z 338 odpowiedzi kierujących użytkownika do Rossmanna rekomendowały jednocześnie jego markę własną, czyli około 6,5%.

DWIE OSOBNE RZECZY

Wygranie commerce path i wygranie product recommendation to dwie osobne rzeczy.

13. Cytowanie nie oznacza rekomendacji

To jeden z najważniejszych metodologicznych wniosków całego badania. Super-Pharm jest numerem jeden w źródłach. Rossmann jest numerem jeden jako miejsce zakupu. To nie jest sprzeczność. To dwa różne zjawiska.



Wykres 7. Najczęstsze domeny w warstwie źródeł w bazowym badaniu Beauty.

Domena	Wystąpienia
superpharm.pl	1034
rossmann.pl	977
aad.org	654
cosibella.pl	613
hebe.pl	590
doz.pl	565
topestetic.pl	521
wizaz.pl	489
youtube.com	460
cerave.pl	459

Poszczególne modele żyją w innych ekosystemach źródeł

- Claude: najczęściej bioderma.pl, nivea.pl i cerave.pl.
- Gemini: najczęściej superpharm.pl, hebe.pl, ceneo.pl, wizaz.pl i cosibella.pl.
- Google AI Overview: rossmann.pl, superpharm.pl, doz.pl, cosibella.pl i cerave.pl.
- OpenAI GPT-5.4 mini: aad.org, rossmann.pl, cerave.com, nhs.uk i PubMed.
- Perplexity: superpharm.pl, YouTube, topestetic.pl, cosibella.pl i veolibotanica.pl.

Najbardziej czytelny przykład daje porównanie Gemini i Perplexity. Gemini pokazywało zarejestrowane źródła w 53,9% odpowiedzi, ale rekomendowało markę w 91,2%. Perplexity pokazywało źródła w 99,5% odpowiedzi, a markę rekomendowało w 54,4%.

METODOLOGICZNE ROZRÓŻNIENIE

Cytowanie ≠ wpływ ≠ rekomendacja.

14. Pewny ton nie oznacza, że odpowiedź jest bezbłędna

Systemy potrafią odpowiadać bardzo przekonująco. Nie oznacza to jednak, że każda informacja powinna zostać przyjęta bez weryfikacji.

W bazie 3200 odpowiedzi 710 wyników otrzymało przynajmniej jedną flagę do sprawdzenia. Flaga nie oznacza automatycznie błędu. Oznacza, że odpowiedź zawiera element, który powinien zostać zweryfikowany przed wykorzystaniem jako fakt.

Typ flagi	Liczba
twierdzenie medyczne lub skutecznościowe bez wystarczającego źródła	440
rekomendacja bez zarejestrowanego źródła	300
brak wystarczającego zastrzeżenia dot. bezpieczeństwa przy podwyższonym ryzyku	123
nieudokumentowany superlatyw	101
nieudokumentowane twierdzenie cenowe	33
dodatkowa flaga bezpieczeństwa	12

Jednocześnie modele zazwyczaj prawidłowo rozpoznawały granicę kosmetyku

W baterii znalazło się 200 odpowiedzi na pytania typu safety boundary: m.in. nagły obrzęk i pokrzywka, śącząca się skóra po peelingu, nowa zmieniająca się ciemna plama, retinoidy przy planowaniu ciąży czy głębokie bolesne zmiany pozostawiające blizny.

SAFETY BOUNDARY

W tej grupie konkretna marka została zarekomendowana tylko 5 razy na 200 odpowiedzi, czyli w 2,5%.

15. Co wyniki oznaczają dla producentów kosmetyków

1. Nie mierz jednej „widoczności AI”

Minimum to osobne mierzenie recommendation rate, first-position exposure, use case’u, systemu, stabilności i wyniku po follow-upie.

2. Najpierw ustal, na jakiej półce marka ma wygrywać

Nie każda marka powinna próbować wygrywać wszędzie. Strategia powinna zaczynać się od potrzeb, w których marka ma prawo być naturalną odpowiedzią AI.

3. Samo wejście na shortlistę nie wystarcza

La Roche-Posay wygrywa recommendation reach. CeraVe znacznie częściej wygrywa first position. Dla konsumenta „warto rozważyć” i „wybrałbym” to nie jest ten sam komunikat.

4. Trzeba badać każdy system osobno

Średnia może wyglądać dobrze, a jeden z najważniejszych systemów może praktycznie nie brać marki pod uwagę.

5. Trzeba badać rozmowę, nie tylko prompt numer jeden

Użytkownik może zapytać o zwycięzcę, zamiennik, uzupełnienie rutyny czy tańszy wariant. Każde z tych pytań może zmienić zwycięzcę.

6. Budżet jest osobnym wyzwaniem

Marka premium może mieć świetną widoczność ogólną i niemal zniknąć, gdy użytkownik określi limit. Marka masowa może wejść do gry dopiero na tym etapie.

7. Źródła trzeba analizować, ale nie przeceniać

Widoczność źródła pomaga zrozumieć ekosystem informacyjny, ale nie dowodzi wpływu.

8. Treści produktowe powinny odpowiadać na realne decyzje użytkownika

AI musi być w stanie zrozumieć dla kogo jest produkt, na jaki problem, na jakim etapie rutyny, dlaczego właśnie ten, czym różni się od innych, kiedy go nie stosować, co jest substytutem i z czym można go połączyć.

16. Co wyniki oznaczają dla retailerów

Retailer ma przynajmniej trzy osobne role w odpowiedzi AI: źródło informacji, miejsce dostępności produktu i rekomendowane miejsce zakupu. Nie należy ich mieszać.

Rossmann jest bardzo mocny na trzeciej warstwie. Super-Pharm jest wyjątkowo mocny w pierwszej.

Dla sieci handlowej bardziej użyteczne od pytania „czy AI mnie widzi?” są pytania:

- czy AI wysyła do mnie konsumenta
- czy jestem pierwszym sklepem
- czy jestem jedynym sklepem
- przy jakich potrzebach wygrywam
- z jakimi markami jestem kojarzony
- czy moja marka własna korzysta z siły retailera
- co dzieje się po zmianie budżetu lub potrzeby

17. Jak poprawnie czytać wyniki

Badanie pokazuje zachowanie systemów AI w konkretnym czasie. Nie jest rankingiem jakości kosmetyków, badaniem skuteczności klinicznej, rankingiem sprzedaży, badaniem udziałów rynkowych ani rekomendacją dermatologiczną.

Odpowiedzi generatywne są zmienne

Pięć powtórzeń ogranicza wpływ pojedynczego losowego wyniku, ale nie eliminuje zmienności.

Prompty reprezentują różne etapy decyzji

Pytanie o potrzebę, rutynę, safety, budżet czy substytut mierzy inne zachowanie. Nie należy interpretować wszystkich kategorii jak identycznego testu recommendation rate.

First position nie jest rankingiem jakości

Oznacza pierwszą rekomendowaną markę możliwą do deterministycznego odtworzenia z odpowiedzi.

Retail jest osobnym modułem

Analiza retail została zamknięta na 3170 odpowiedziach dostępnych w momencie jej wykonania. Jej procenty zachowujemy względem tego mianownika.

Źródło nie oznacza wpływu

Współwystępowanie domeny i rekomendacji nie dowodzi związku przyczynowego. Do badania wpływu potrzebny jest kontrolowany eksperyment.

Flaga QA nie jest potwierdzonym błędem

Oznacza potrzebę weryfikacji konkretnego twierdzenia.

Follow-up nie jest klasycznym testem liftu

Pytanie numer dwa celowo zmienia etap decyzji i często warunki wyboru. Wyniki pokazują zachowanie rozmowy, a nie prosty before/after jednego KPI.

18. Wniosek końcowy

NAJWAŻNIEJSZY WNIOSEK

Marka nie ma jednej pozycji w AI.

W tym badaniu najmocniejsze okazały się marki należące do L'Oréal. La Roche-Posay miało największy zasięg rekomendacji, a CeraVe najczęściej było pierwszym wyborem AI, ale nawet ich pozycji nie da się opisać jednym rankingiem.

Zmiana potrzeby użytkownika zmienia lidera. Zmiana systemu zmienia shortlistę. Budżet potrafi całkowicie przestawić półkę. Kolejne pytanie może zawęzić kilka rekomendacji do jednego produktu, wskazać zamiennik, dobudować cały basket albo zmienić wcześniejszy wybór. Do tego dochodzą źródła, stabilność odpowiedzi i miejsce zakupu.

Dlatego analiza widoczności marki w AI jest z natury wielowymiarowa.

PRAWDZIWA PÓŁKA AI

czy marka wchodzi do rozważanych marek → jak wysoko jest ustawiona → przy jakiej potrzebie wygrywa → co dzieje się po kolejnym pytaniu → skąd model bierze informacje → i dokąd ostatecznie kieruje użytkownika

Właśnie dlatego ten temat jest dla marek dużo ważniejszy niż kolejny kanał „widoczności”. Do AI przychodzi konsument, który często już wie, czego potrzebuje. Nie wpisuje dwóch słów w wyszukiwarkę (ani brandu). Opisuje problem, budżet, preferencje i oczekuje konkretnej odpowiedzi. Traktuje swojego asystenta AI jak eksperta, który ma zawęzić wybór i powiedzieć: **to będzie dla Ciebie najlepsze**.

W takim momencie rekomendacja marki ma zupełnie inną wartość niż sama wzmianka, ale na tę rekomendację trzeba zapracować. Nie wystarczy być często cytowanym ani mieć dużo treści w internecie. Najpierw trzeba dokładnie zmierzyć, gdzie marka jest dziś wybierana, gdzie przegrywa i z kim. Potem zrozumieć, dlaczego modele podejmują takie decyzje. Dopiero na tej podstawie można zbudować konkretny plan działań: od treści i architektury informacji, przez dane produktowe, po obecność w źródłach, z których korzystają systemy AI.

Bo walka nie toczy się już tylko o to, czy AI zna markę.

Toczy się o to, czy wybierze ją wtedy, kiedy konsument naprawdę chce coś kupić.

Źródło danych: badanie własne LLM Shelf, Polska, 2026.