
Aleksandra Skrzypiec

Szkoła Doktorska Nauk Społecznych, Uniwersytet Jagielloński

ORCID: 0000-0001-9731-9111

Trendy generowane przy pomocy sztucznej inteligencji: O smogu informacyjnym w mediach społecznościowych na podstawie wybranych przykładów „AI spamu”

Wprowadzenie

Celem niniejszego artykułu o charakterze teoretyczno-przeglądowym jest przyjrzenie się nowym formom dezinformacji, misinformacji oraz malinformacji – składającym się na zjawisko smogu informacyjnego – rozprzestrzeniającym się w mediach społecznościowych za sprawą materiałów audiowizualnych generowanych przez sztuczną inteligencję. Problematyka dotycząca obecności w *social mediach* nieprawdziwych, wprowadzających w błąd treści, zyskała szczególny rozgłos w świecie polskiego internetu z początkiem 2025 r., za sprawą memu opublikowanego przez twórców facebookowej strony *Polska w dużych dawkach*. Zdjęcie przedstawiające – jak mogłoby się wydawać – kobietę pracującą w piekarni oraz stojącą obok niej gigantyczną chałkę w kształcie konia zyskało ponad 11 tysięcy reakcji. Fotografii towarzyszył opis „Ta kobieta upiekła chałkonia, ale nikt jej nie pogratulował”, nawiązujący do pojawiających się wówczas w mediach społecznościowych na masową skalę postów z generycznymi dopiskami „X zrobił Y, ale nikt mu nie pogratulował, bo...”. Ku zaskoczeniu samych

administratorów fanpage'a niewinny żart stał się przyczynkiem do głębszych refleksji na temat bezpieczeństwa w sieci, bowiem tysiące użytkowników platformy Facebook uwierzyło, iż opublikowane zdjęcie przedstawia efekt rzeczywistych działań. Pod materiałem pojawiły się gratulacje skierowane w stronę kobiety – w szczególności swoją aprobatę wyrażały osoby starsze. „Chałkoń” został również zweryfikowany przez Państwowy Instytut Badawczy NASK pod kątem oddziaływania wpływów rosyjskich na mające się odbyć w 2025 r. w Polsce wybory prezydenckie (Donald 2025). Autorzy mema niepełna dwa dni po udostępnieniu postu zamieścili apel, w którym pochyli się nad skutkami własnych działań, ostrzegając jednocześnie przed podobnego rodzaju materiałami, które można napotkać w mediach społecznościowych (Bojakowski 2025). Sam motyw „chałkonii” zyskał bardzo dużą popularność w polskim internecie – sprowokował nie tylko liczne reakcje ze strony internautów, ale także znalazł się w centrum zainteresowania marek takich jak m.in. Żabka, Dino, Kaufland, Ikea, mBank, Dr. Oetker, które podążając za panującym w sieci trendem, uczyniły go bohaterem prowadzonych w czasie rzeczywistym działań promocyjnych (tzw. *real-time marketing*).

„Afera chałkoniowa” w istocie wskazuje na piętrzące się w dobie rozwoju generatywnej AI wyzwania, obejmujące całą – zarówno materialną, jak i pozamaterialną – infrastrukturę platform społecznościowych, tj. środowisko technologiczne, a także przestrzeń społecznych interakcji oraz znaczeń symboliczno-kulturowych. Główny cel niniejszego artykułu o charakterze teoretyczno-przeglądowym stanowi odpowiedź na pytanie badawcze, z jakimi nowymi formami smogu informacyjnego rozprzestrzeniającymi się za sprawą materiałów audiowizualnych generowanych przy użyciu GenAI mają styczność użytkownicy *social mediów*. W odniesieniu do zaproponowanej problematyki postawiono również dwa pytania pomocnicze: (1) jakie są społeczno-kulturowe konsekwencje panujących w przestrzeni platform społecznościowych trendów? oraz (2) które z grup społecznych są na nie szczególnie narażone?

W tekście wskazuje się na konkretne zjawiska, takie jak polityczna dezinformacja w dobie generatywnej sztucznej inteligencji (m.in. strategia *flooding the zone*), *AI slop*, *brainrot* oraz *boomertrap*, które nie zostały dotychczas szczegółowo opisane na gruncie rodzimej oraz obcojęzycznej literatury naukowej, co na ten moment oznacza również brak osadzonych w dyscyplinie nauki o komunikacji społecznej i mediów konceptualizacji wyżej wymienionych terminów. Wobec tego autorka w znaczącym stopniu powołuje się m.in. na polskie i anglojęzyczne materiały dziennikarskie, tj. artykuły informacyjno-publicystyczne o charakterze pozanaukowym. Wraz z omawianymi przykładami udostępnionych w mediach społecznościowych postów należy traktować je jako materiał empiryczny aniżeli klasycznie rozumianą literaturę teoretyczną. W związku z tym na wykorzystane

w artykule metody badawcze składają się przede wszystkim: (1) **przegląd literatury** naukowej odnoszącej się do terminu smogu informacyjnego, zjawiska dezinformacji oraz innych, pokrewnych konstruktów teoretycznych, (2) **jakościowa analiza dyskursu** oparta na krajowych oraz zagranicznych materiałach dziennikarskich i *fact-checkingowych*, opisujących przykłady syntetycznych treści o potencjalnie szkodliwym charakterze, (3) **krótkie studium przypadków** konkretnych postów – bazujących na obrazach lub filmikach wygenerowanych przez sztuczną inteligencję – publikowanych na platformach społecznościowych. Z jednej strony zastosowana strategia metodologiczna – tj. przegląd literatury wraz z analizą znaczącej liczby materiałów medialnych – stanowi istotne ograniczenie podejmowanych rozważań, z drugiej jednak pozwala: (1) lepiej scharakteryzować nowe zjawiska, (2) w syntetyczny sposób przedstawić czytelnikowi możliwe przyczyny ich występowania, a także (3) omówić najistotniejsze – choć jeszcze nie w pełni rozpoznane – konsekwencje. Biorąc pod uwagę kwestie metodologiczne, należy również wskazać, iż włączone do analizy materiały medialne – zarówno dziennikarskie, jak i posty z mediów społecznościowych – musiały spełniać następujące kryteria: (1) zawierać/odnosić się do popularnych od około 2023 roku trendów – tj. od momentu, w którym GAI zyskała szerszy rozgłos wśród opinii publicznej za sprawą udostępnienia chatu GPT – wykreowanych w *social mediach* przy pomocy GenAI, (2) opisywać specyfikę popularnych w podanym okresie zjawisk lub (3) omawiać ich społeczno-kulturowe konsekwencje, zyskujące dopiero na znaczeniu w obszarze naukowych analiz.

Smog informacyjny w mediach społecznościowych

Trzecia dekada XXI w. to kolejna z tzw. „wiosen” w rozwoju sztucznej inteligencji – w szczególności uczenia maszynowego (ang. *machine learning*), uczenia głębokiego (ang. *deep learning*) oraz przetwarzania języka naturalnego (ang. *Natural Language Processing*), kluczowych dla działania, skutecznej implementacji oraz upowszechniania rozwiązań z zakresu generatywnej sztucznej inteligencji (Hagos, Battle i Rawat 2024). W obliczu zachodzących przeobrażeń wynikających z dokonującego się rozwoju technologicznego oczywista wydaje się konstatacja, iż internet stanowi przestrzeń do niemal nieograniczonej partycypacji w procesach wytwarzania, wymiany i transmisji treści zróżnicowanych pod względem nadawanych im form oraz pełnionych funkcji. Zwrot w tym kierunku nastąpił wraz z nadejściem epoki Web 2.0, która w historii komunikacji społecznej, mediów, w tym technologii informacyjno-komunikacyjnych zdążyła zapisać się jako moment, w którym za sprawą pojawienia się *World Wide Web* doszło do rozmycia wyraźnych granic pomiędzy nadawcą a odbiorcą komunikatów. Proces demokratyzacji

dostępu do informacji umożliwił biernym jak dotąd odbiorcom: (1) wypełnianie internetowej przestrzeni komunikacyjnej własnymi treściami, przy jednoczesnej możliwości ich konsumpcji za sprawą szerokiej oferty dostępnych materiałów medialnych, a także (2) realizację własnych przedsięwzięć oraz inicjatyw w skali globalnej (Sarowski 2017: 34–35). Od tamtej pory w znaczący sposób – a w niektórych sferach nawet diametralny – ewoluowały zarówno narzędzia technologiczne, jak i praktyki medialne, dając początek erze Web 3.0, która stopniowo ulega przeobrażeniom w kierunku Web 4.0.

Przemiany zachodzące w obszarze wytwarzania, dystrybucji oraz konsumpcji treści – choć konieczne i nieuchronne – stały się źródłem nowych zagrożeń. Ryszard Tadeusiewicz już w 1999 r., kiedy w Polsce zaczynało budować społeczeństwo informacyjne, opisywał zjawisko smogu informacyjnego, opierając się na analogicznym przykładzie smogu stanowiącego mieszaninę powietrza i zanieczyszczeń. Smog występujący w atmosferze to skutek uboczny:

[...] prymitywnego i nieuporządkowanego procesu spalania byle czego, byle gdzie i byle jak – dostarczającego niezbędnej energii dla rozmaitych procesów wytwórczych, burzliwie i chaotycznie rozwijanych na początku poprzedniej rewolucji przemysłowej (Tadeusiewicz 1999: 99).

W podobny zdaniem Tadeusiewicza sposób produkowane są treści w cyberprzestrzeni. Smog informacyjny składa się z rozrzuconych, rozdrobnionych i nieuporządkowanych informacji, które tak, jak kropelki wody:

[...] tworzą informacyjną mgłę, która oślepia, dusi, utrudnia orientację, pozbawia sens dotarcia bezpiecznie do spokojnego portu rzetelnej wiedzy – a w przypadku osób mało krytycznych i mało doświadczonych – nawet łatwo wyprowadza na manowce pseudoprawd i paranoi (Tadeusiewicz 1999: 100).

Choć zaprezentowane rozważania sformułowano ponad ćwierć wieku temu, pozostają one nadal aktualne. Można zaryzykować również stwierdzenie, iż w obliczu rozwoju AI nabierają coraz większego znaczenia. Tadeusiewicz w krytycznym tonie wypowiadał się o ciemniej stronie internetu, który jego zdaniem za sprawą materiałów o charakterze rasistowskim, pornograficznym, antyhumanistycznym, a także pseudonaukowym odpowiada za szerzenie się mowy nienawiści, radykalizację poglądów politycznych i religijnych, hamowanie w ludziach refleksyjności, empatii oraz tolerancji (Tadeusiewicz 1999: 100–101).

Postępy w dziedzinie sztucznej inteligencji prowadzą nie tylko do intensyfikacji opisywanych problemów, ale także przyczyniają się do powstawiania nowych zagrożeń. Zmienia się bowiem rola AI w procesie

komunikowania – biorąc pod uwagę narzędzia takie jak chatboty, voiceboty oraz roboty społeczne, wyraźnie dostrzegalne jest wyjście poza role przekazników treści w kierunku pozaludzkich nadawców oraz odbiorców w komunikacji zarówno na poziomie społecznym, jak i masowym. Wobec tego twórcami materiałów o charakterze dezinformacyjnym mogą być nie tylko ludzie, ale także narzędzia sztucznej inteligencji generujące „wysokiej jakości, kontekstowo trafne treści, niemal nieodróżnialne od tych stworzonych przez człowieka” (Banh i Strobel 2023: 62 [tłum. własne]). Przy pomocy modeli dyfuzyjnych i dużych modeli językowych tworzą one wysokiej jakości obrazy, proste teksty, sensowną prozę oraz poezję (Epstein i Hertzmann 2023: 1110). Dostępność coraz bardziej zaawansowanych technologii, w tym łatwość ich obsługi powodują przyrost syntetycznych treści, wzmacniających podziały społeczne, przemoc ideologiczną oraz represje polityczne, wobec czego zakłada się, że w jeszcze większym stopniu będą one manipulować jednostkami, szkodzić gospodarce i eskalować konflikty (World Economic Forum 2024: 18). Materiały GenAI coraz trudniej odróżnić od tych, które „wychodzą spod ludzkiej ręki” ze względu na coraz rzadsze halucynacje i nieścisłości. Obawy związane z ich wykorzystywaniem w mediach dotyczą również: (1) treści o charakterze wizualnym, ukazujących sfabrykowane i nieprawdziwe wydarzenia w sposób realistyczny, przez co coraz większym wyzwaniem staje się oddzielenie prawdy od fikcji; (2) rzetelnego oraz przejrzystego informowania o udostępnianiu zweryfikowanych obrazów; (3) praw autorskich do materiałów powstałych na podstawie istniejących dotychczas zdjęć, którym prawa te przysługują (Sohyun i in. 2024: 570).

Smog informacyjny zakłóca proces informowania, który powinien służyć:

[...] powiadomieniu społeczeństwa lub określonych zbiorowości w sposób zbiektywizowany, systematyczny i konkretny, za pomocą środków masowego przekazu o bieżących lub prognozowanych wydarzeniach, mających istotne znaczenie [...] lub wzbudzających szczególne zainteresowanie odbiorców (Januszko-Szakiel 2010: 210).

Informowanie umożliwia oddzielenie informacji od opinii, a także faktów od komentarzy, tworząc tym samym przestrzeń do kształtowania własnych interpretacji – w tym podejmowania świadomych wyborów w zakresie przyjmowanych postaw politycznych. Informacja jako jeden z elementów komunikowania powinna minimalizować odczuwany przez jednostkę stan niewiedzy i niedookreśloności (Januszko-Szakiel 2010: 210). Magdalena Szpunar pisze, iż „w dobie kultury cyfrowej, która niewątpliwie poszerza nasz dostęp do informacji, nasze zdolności do ich przyswojenia bynajmniej się nie powiększyły” (Szpunar 2017: 29) – w konsekwencji jednostki

narażone na kontakt ze smogiem informacyjnym mogą doświadczać syndromu braku odporności na informację – kulturowego *anti-information deficiency syndrome* (AIDS). Ewolucyjnie uwarunkowane deficyty w zakresie przetwarzania informacji – których w erze cyfrowej występuje nadmierna ilość – w połączeniu z niedostatecznym poziomem kompetencji medialnych uniemożliwiają rzetelną ocenę oraz selekcję przekazów docierających z różnych źródeł. W ten sposób podatność na bezkrytyczne przyswajanie materiałów o charakterze dezinformacyjnym wzrasta. Dezinformacja w wąskim ujęciu oznacza „nieprawdziwą lub nieprecyzyjną informację, która została wytworzona intencjonalnie w celu wprowadzenia w błąd” (Stasiuk-Krajewska 2023: 59). Składa się na nią szereg przemyślanych, zaplanowanych i usystematyzowanych kroków prowadzących przede wszystkim do destabilizacji politycznej lub społecznej – w tym wyrządzenia szkód innym aktorom ulokowanym na poziomie mikro-, mezo- oraz makrostruktur społecznych. Do metod dezinformacji zaliczają się m.in.: negacja faktów, odwrócenie faktów, mieszanie prawdy i kłamstwa, modyfikacja motywu i okoliczności, rozmycie, kamuflaż, interpretacja, ilustracja, generalizacja, nierówna lub równa reprezentacja (Januszko-Szakiel 2010: 212), a także – satyra (parodia), fałszywe (błędne) połączenia, treści wprowadzające w błąd, fałszywy kontekst, podszywanie się pod prawdziwe źródła, zmanipulowana lub sfabrykowana zawartość (Wardle 2017 za: Łódzki 2017: 23–24). Jak zostało wspomniane we wprowadzeniu, widocznym problemem w przestrzeni mediów społecznościowych stało się rozpowszechnianie treści w całości zmodyfikowanych i nieprawdziwych. Ponadto dezinformacja obecna w środowisku cyfrowym przejawia się również poprzez: (1) wykorzystywanie emocjonalnie nacechowanej retoryki, w tym niespójnych i wykluczających się argumentów; (2) posługiwanie się fałszywymi dychotomiami lub dylematami; (3) używanie argumentów *ad hominem*, uderzających wprost w człowieka, a nie konkretne poglądy oraz (4) przypisywanie winy kołowi ofiarnemu (Roozenbeek i in. 2022: 2). Poza dezinformacją wyróżnia się również pojęcia, takie jak:

- **misinformacja** – treści nieprawdziwe, tworzone i rozpowszechniane bez zamiaru spowodowania szkody. Za ich główne przyczyny uznaje się błędy, emocje, brak rzetelności – inaczej szybkie oraz nieprzemyślane działania, prowadzące do nieintencjonalnego generowania negatywnych efektów w sferze konsumpcji treści (m.in. poprzez udostępnianie niektórych memów internetowych);
- **malinformacja** – materiały prawdziwe, których celem jest wyrządzenie krzywdy. Są one podstawą dla prowadzenia działań dezinformacyjnych. Należą do nich m.in. prywatne wiadomości czy też treści dotyczące krępujących sytuacji osobistych. Szerzenie materiałów o wskazanym charakterze opiera się na podawaniu faktów oraz

intencjonalnym dodawaniu w ich świetle informacji mających wywołać chaos lub podziały społeczne (Instytut Kościuszki b.r.; NASK 2025a).

W dalszej części tekstu został zawarty opis, uzupełniony o przykłady wyróżnionych w następujący sposób syntetycznych, a więc nieprawdziwych materiałów składających się na smog informacyjny:

1. Dezinformacyjne o *stricte* szkodliwym charakterze, stanowiące poważne zagrożenie dla procesu komunikowania politycznego.
2. „Zaśmiecający” infosferę AI *slop* – przybierający niekiedy memiczno-rozrywkową formę określaną jako *brainrot*, w tym również *boomertrap*, inaczej „pułapki na boomerów”, niebezpieczne w szczególności z perspektywy starszych i mniej doświadczonych użytkowników platform społecznościowych.

Druga z wymienionych kategorii stanowi w istocie połączenie dezinformacji, misinformacji oraz malinformacji, bowiem szkodliwość materiałów zaliczanych do wskazanej grupy jest stopniowalna, a intencje przyświecające ich wytwarzaniu nie zawsze są jednoznaczne. Nie wszystkie treści stworzone przez generatywną sztuczną inteligencję są wykorzystywane do celowego upowszechniania dezinformacji, jednak mimo to – biorąc pod uwagę, iż część z nich jest celowo wytwarzana do osiągania korzyści finansowych – mogą one wyrządzać szkody. Skutki ich konsumpcji – oraz udostępniania dalej – również nie są jednakowe, ponieważ z pozoru niegroźne obrazy lub materiały audiowizualne negatywnie oddziałują na ogólny dobrostan poszczególnych grup społecznych, o czym krótko mowa w podpunkcie poświęconym treściom typu *brainrot*. Badawczym wyzwaniem staje się dokonanie rozłącznej i wyczerpującej klasyfikacji pojawiających się trendów internetowych – ściśle korespondującej z podziałem na dezinformację, misinformację oraz malinformację. W związku z tym w dalszej części artykułu wyróżniono zjawiska, wymienione w punkcie 1 oraz 2 niniejszego akapitu, określane mianem „AI spamu” – stanowiącym w ujęciu autorki zbiorczą kategorię pojęciową wpisującą się w metaforę smogu informacyjnego – nadrzędną względem syntetycznych materiałów dezinformacyjnych oraz treści typu *AI slop*, *brainrot* oraz *boomertrap* – do której zaliczają się wyprodukowane przez GenAI materiały w postaci zdjęć oraz filmów, mające na celu zwrócenie uwagi użytkowników, wprowadzenie ich w błąd lub wywołanie kontrowersji.

Przykłady politycznej dezinformacji tworzonej z użyciem GenAI

Sztuczna inteligencja – w szczególności duże modele językowe – umożliwiają generowanie materiałów o charakterze dezinformacyjnym na masową skalę, pochłaniając przy tym coraz mniej czasu i kosztów. W ten sposób kampanie

dezinformacyjne tworzone przy pomocy AI mogą sprzyjać zarówno wprowadzaniu w błąd opinii publiczną, kształtowaniu poglądów na ważne w dyskursie publicznym tematy, jak i tworzeniu fałszywych opinii większości (Weidinger i in. 2022: 218).

Biorąc pod uwagę percepcję treści o charakterze audiowizualnym, należy za Cristianem Vaccari i Andrew Chadwickem zwrócić uwagę, iż zarówno obraz, jak i wideo mają większą siłę oddziaływania niż tekst – samemu procesowi ich przetwarzania towarzyszy mniejszy wysiłek poznawczy (Vaccari i Chadwick 2020: 2). Materiał prezentowany w formie obrazu jest bliższy codziennym ludzkim doświadczeniom. Znajome sytuacje wywołują w związku z tym „efekt prawdziwości”. Obraz, łatwiejszy do przetworzenia i zrozumienia od tekstu, staje się bardziej wiarygodny dla odbiorcy. Nadmierna konsumpcja sfabrykowanych materiałów może doprowadzić do poczucia niepewności, a w konsekwencji do spadku zaufania wobec treści umieszczanych w mediach społecznościowych – w tym również do ich unikania. Obecność materiałów o charakterze dezinformacyjnym w przestrzeni *social mediów* może również przyczyniać się do kształtowania kultury obywatelskiej – kreowania niewłaściwych norm oraz zachowań, zwiększając przy tym: (1) obojętność użytkowników względem odpowiedzialnego dzielenia się informacjami i (2) przekonanie, że w internecie wszystko „uchodzi na sucho” (Vaccari i Chadwick 2020: 9).

We wskazanym kontekście warto w szczególności zwrócić uwagę na dezinformację o charakterze politycznym, która niesie za sobą poważne konsekwencje, ponieważ oddziałuje na przebieg wyborów i prowadzi do niepokojów społecznych – w tym niestabilności politycznej, załamania zaufania do instytucji oraz mediów, podżegania do nietolerancji i przemocy, naruszania bezpieczeństwa psychologicznego – umacniając równocześnie pozycje nieodpowiedzialnych i autorytarnych liderów (Walker, Schiff i Schiff 2024: 23053–23054). Szkody związane z obecnością materiałów dezinformacyjnych w mediach – w tym odpowiedzi zamieszczanych w kontrze do nich – „wynikają z wszechobecnych czynników działających poprzez złożone łańcuchy przyczynowe, z efektami sprzężenia zwrotnego i heterogenicznymi efektami utrudniającymi analizę” (Walker, Schiff i Schiff 2024: 23054 [tłum. własne]).

Popularna w ostatnim czasie strategia dezinformacyjna została określona przez amerykańskiego doradcę politycznego oraz dziennikarza Steve’a Bannona jako *flooding the zone*. Polega ona na intencjonalnym „zalewaniu” infosfery materiałami, które – biorąc pod uwagę cechę, jaką jest wiarygodność – wzbudzają niepewność wśród odbiorców (*The Media Manipulation Casebook* b.r.). W przeciwieństwie do tradycyjnej propagandy, której celem jest rozpowszechnianie, umacnianie i podtrzymywanie jednej, spójnej linii narracyjnej – *flooding the zone* opiera się na udostępnianiu mediom bardzo dużej liczby sprzecznych informacji – co nie tylko utrudnia weryfikację ich

prawdziwości, ale przede wszystkim odbiera chęć do poszukiwania prawdy (Illing 2020). Zsolt Kapelner z Uniwersytetu w Oslo wskazuje, iż wraz z nastaniem drugiej prezydentury Donalda Trumpa w Stanach Zjednoczonych światowe media zaczęły donosić o szeregu wydawanych dekrétów wykonawczych, a także o kontrowersyjnych oświadczeniach, takich jak deklaracja chęci przejęcia kontroli nad Grenlandią – mających służyć w istocie wywołaniu zamieszania wśród mediów, opozycji oraz społeczeństwa (Kapelner 2025). Strategia *flooding the zone* zdaniem Kapelnera jest nie tylko antydemokratyczna, ale i cyniczna, ponieważ wykorzystuje biologicznie uwarunkowaną ludzką słabość w postaci ograniczonych możliwości poznawczych (Kapelner 2025). Aleksandra Przegalińska do działań składających się na opisywany rodzaj dezinformacji zalicza również udostępniane przez sztab Donalda Trumpa obrazy, wygenerowane przy pomocy AI, na których prezydent USA został przedstawiony jako papież lub postać z sagi *Gwiezdnych wojen* – co świadczy o znaczeniu GenAI w działaniach umacniających chaos informacyjny (Przegalińska 2025). *Flooding the zone* to w amerykańskim wydaniu serwowana mediom i społeczeństwu mieszanka sprzecznych informacji, abstrakcyjnych obrazów i filmików wygenerowanych przez sztuczną inteligencję – takich jak chociażby kontrowersyjne wideo znajdujące się w serwisie YouTube pod nazwą *Trump Gaza*, udostępnione przez samego prezydenta USA prawdopodobnie w celu umocnienia inkorporowanego społeczeństwu wizji przekształcenia Strefy Gazy w amerykańską riwierę (Kreuer i Salem 2025).

Przykładem materiału z polskiej sceny politycznej, który podobnie jak „chałkoń” zyskał zasięg w mediach społecznościowych i tym samym zdołał wprowadzić wielu użytkowników w błąd, było zdjęcie udostępnione przez polityków Prawa i Sprawiedliwości – Agnieszkę Sosin oraz Jana Warzechę – w trakcie kampanii prezydenckiej w roku 2025. Materiał mający stanowić wyraz poparcia dla Karola Nawrockiego – kandydata ubiegającego się wówczas o stanowisko prezydenta – przedstawiał młode, uśmiechnięte kobiety pozujące w letnich sukienkach (mimo iż miesiącem publikacji obydwóch postów był marzec) z napisem „NAWROCKI 2025”, trzymanym w rękach. Zarówno Agnieszka Sosin, jak i Jan Warzecha nie podali źródła pochodzenia zdjęcia, jednak kardynalne błędy w postaci m.in. nienaturalnie ułożonych rąk, braku palców, zdeformowanych dłoni oraz letnich ubrań wskazywały w tym przypadku na użycie AI. Pomimo wymienionych niedociągnięć pod postami Sosin i Warzechy pojawiło się zarówno mnóstwo przychylnych komentarzy od osób sympatyzujących z kandydatem, jak i wulgarnych opinii na temat kobiet, zamieszczonych przez przeciwników Karola Nawrockiego (Demagog 2025a). Jedynie nieliczni użytkownicy zwracali uwagę, że obraz został wygenerowany przez sztuczną inteligencję. Podobny chaos informacyjny w polskich mediach społecznościowych wywołało sfabrykowane zdjęcie przedstawiające przywódców europejskich państw podczas spotkania

z Donaldem Trumpem w Białym Domu (Demagog 2025b), czy też materiał, w którym Karol Nawrocki wymienia pocałunek ze swoją doradczynią podczas wręczenia nominacji po objęciu prezydentury (Demagog 2025c).

Opisane dotychczas sposoby rozpowszechniania dezinformacji przy pomocy strategii *flooding the zone*, czy też wykorzystywanie nieprawdziwych, stworzonych z użyciem GenAI obrazów, intencjonalnie wprowadzających w błąd odbiorców to przykłady działań ukierunkowanych na wprowadzanie chaosu informacyjnego do sfery komunikowania politycznego, a tym samym na zakłócanie właściwego procesu informowania poprzez tworzenie smogu informacyjnego, składającego się z niskiej jakości materiałów. Zasadnym wydaje się stwierdzenie, iż całość tych działań może służyć wywoływaniu wśród odbiorców poczucia dezorientacji, spowodowanego przez trudności w odróżnianiu fałszywych od prawdziwych reprezentacji otaczającej rzeczywistości społecznej. W dalszej części tekstu przedstawiono inne przykłady materiałów, z którymi spotykają się użytkownicy mediów społecznościowych, intensyfikujących chaos informacyjny i przyczyniających się jednocześnie do obniżania jakości napotykanych w cyberprzestrzeni treści.

Rodzaje generatywnych treści w mediach społecznościowych

Renée DiResta i Josh A. Goldstein na podstawie analizy wyświetlonych setki milionów razy treści¹ znajdujących się na 125 facebookowych stronach publikujących obrazy generowane przez AI ustalili, iż: (1) stosują one taktyki *clickbaitowe*, przez które użytkownicy trafiają na strony zewnętrzne w postaci farm niskiej jakości treści – sprzedające niekiedy nieistniejące produkty, usiłujące przejąć cudze profile lub wyłudzić dane, (2) dokonując oszustw, starają się uchodzić za wiarygodne w oczach odbiorców, dlatego angażują się w dyskusje z nimi i gromadzą nieautentycznych obserwatorów, (3) algorytmy Facebooka same wyświetlają generyczne materiały, niezależnie od tego, czy konkretne konto zostało wcześniej zaobserwowane przez użytkowników, (4) komentarze umieszczane pod postami wskazują, iż większość odbiorców nie jest świadoma, skąd w rzeczywistości pochodzą publikowane materiały (DiResta i Goldstein 2024). DiResta i Goldstein wskazują również, iż osobami dokonującymi oszustw przy użyciu nowoczesnych technologii są ich najwcześniejsi użytkownicy, wykorzystujący lukę w okresie przejściowym, między momentem, w którym konkretne narzędzie zaczyna dostarczać

¹ Jeden z tego typu postów w trzecim kwartale 2023 zyskał 40 mln wyświetleń i 1,9 mln interakcji, co ulokowało go wśród 20 najpopularniejszych publikacji na Facebooku w tamtym czasie.

nowe możliwości w zakresie podejmowania nieuczciwych praktyk, a ich rozpoznaniem i przeciwdziałaniem skoncentrowanym na ograniczeniu negatywnych wpływów (DiResta i Goldstein 2024). W dalszej części niniejszej sekcji artykułu opisane zostały zjawiska bazujące na wymienionych powyżej praktykach, określane również przez autorkę trendami, obecnymi od około 2023 roku w mediach społecznościowych.

AI slop – monetyzowanie niskiej jakości materiałów

Pochodzący z języka angielskiego wyraz *slop* w tłumaczeniu na polski oznacza „pomyje”, „zlewki”, „sentymentalne bzdury” oraz „kicz” (słownik języka angielskiego Diki). *AI slop* to nieoficjalny – popularny w świecie internetu – termin, którym określa się nowoczesny spam generowany przez sztuczną inteligencję z minimalną pomocą człowieka. Są nim niskiej jakości treści przedstawiające celebrytów, polityków, przesyczone fantastyką krajobrazy oraz antropomorficzne zwierzęta (Malik 2025). Choć większość tego rodzaju obrazów zdaje się nie mieć większego sensu, to znajdują się wśród nich takie, które ze względu na swój propagandowy charakter są społecznie szkodliwe, a nawet niebezpieczne – zważywszy na możliwości technologiczne w zakresie kreowania fikcyjnych wydarzeń. Wśród tego rodzaju materiałów znajdują się m.in. filmiki z liberałami obalanymi przez administrację Donalda Trumpa (Malik 2025). *AI slop* może również prowadzić do pogłębiania się tendencji polaryzacyjnych ze względu na tematykę niektórych „śmieciovych” publikacji, które wśród skrajnie konserwatywnej części społeczeństwa mogą wywoływać tęsknotę za „dawnymi czasami”. Mowa tu w szczególności o nadmiernej ekspozycji na obrazy przedstawiające tradycyjny model białej rodziny – w szczególności gospodyń (ang. *tradwives*) oddanych obowiązkom domowym i rodzinnym – co w dobie umacniających się podziałów politycznych stwarza ryzyko powstawania coraz większych nierówności na tle rasowym i płciowym (Malik 2025).

AI slop, jak zostało wcześniej wspomniane, to nie tylko treści o charakterze polityczno-propagandowym, ale także generowane masowo – bezużyteczne, przypadkowe, chaotyczne i tandetne – zdjęcia oraz filmiki. Przykuwają one uwagę odbiorców, wywołując przy tym niesmak, niepokój, a czasem nawet zaskoczenie. Arwa Mahdawi z „The Guardian” przytacza w tym kontekście przykład „krewetkowego Jezusa” i zwraca uwagę na politykę Facebooka, umożliwiającą algorytmom rozpowszechnianie podobnych materiałów (Mahdawi 2025). Ta sama sytuacja ma miejsce w przypadku grafik, przy których widnieją generyczne opisy „X zrobił Y, ale nikt mu nie pogratulował, bo...” – przypadkowe zetknięcie się z *AI slop* skutkuje wyświetlaniem jeszcze większej liczby postów o powtarzalnym schemacie. Jak udało się

ustalić internetowemu portalowi 404 Media, Facebook za sprawą programu bonusowego dla twórców wynagradzał osoby publikujące „śmieciowe” materiały, z kolei do ich tworzenia zachęcali youtube’owi twórcy. Jeden z nich to Gyan Abhishek, który uczył swoich widzów promptowania z użyciem fraz takich jak „poor people thin body” (Tangermann 2024). Internetowy twórca jednocześnie podkreślał, iż „the Indian audience is very emotional. After seeing photos like this, they like, comment and share them. So you too should create a page like this, upload photos and make money through performance bonus” (Koebler 2024). Meta w swoim stanowisku przekazanym 404 Media zapewniała, iż większość spamu tworzonego przy użyciu GenAI nie naruszała zasad platformy, a sam program wynagradzający twórców działał zgodnie ze swoimi założeniami, o ile zasięg postów nie był wzmacniany przez boty.

W kontekście upowszechniania się *AI slop* kluczowe okazują się: (1) działanie algorytmów wyświetlających użytkownikom treści pochodzące z niezaobserwowanych wcześniej kont, (2) programy dedykowane monetyzacji oraz (3) zwiększająca się liczba osób oferujących za niewielką opłatą naukę generowania „AI spamu” – przy czym biorąc pod uwagę fakt, iż całkowity zysk pochodzących z tego rodzaju działalności w większości przypadków nie przekracza kilku centów, oznacza to, że podejmują się jej osoby z państw takich jak Indie, Tajlandia, Indonezja i Pakistan, gdzie nawet drobne kwoty mają relatywnie wysoką wartość (The Guardian 2025). Podsumowując niniejszy wątek, należy również wskazać, iż opisywane zjawisko prowadzi do ubożenia środowiska informacyjnego, ponieważ, „śmieciowe” treści zmniejszają widoczność materiałów edukacyjnych, dziennikarskich oraz rozrywkowych (Konopczyński 2025). Co więcej, obrazy i filmy generowane przez AI stały się narzędziami w rękach pracowników sztabów wyborczych, co potwierdziła rywalizacja Kamali Harris i Donalda Trumpa w trakcie wyborów prezydenckich w Stanach Zjednoczonych (Konopczyński 2025). Sięganie po tego rodzaju praktyki umożliwia szerzenie się dezinformacji na coraz większą skalę, prowadząc jednocześnie do jej milczącej akceptacji. Produkcja „śmieciowych” grafik oraz materiałów wideo obciąża również środowisko naturalne, bowiem działanie multimodalnych modeli pochłania i tworzy dalsze zapotrzebowanie na dużą ilość energii (Konopczyński 2025).

Brainrot – memiczno-rozrywkowa forma „AI spamu”

Brainrot to kolejna podkategoria *AI slop*, która bawi użytkowników mediów społecznościowych, wywołuje duże zaangażowanie w ich odbiór, a nawet zachęca do produkcji podobnych treści. *Brainrot* stanowi znak rozpoznawczy cyberkultury pokolenia alfa (Zinkiewicz 2025) i generacji „zetek” (Oxford University Press 2024). Szczególną popularnością cieszy się *italian*

brainrot – seria *viralowo* rozprzestrzeniających się obrazów i materiałów wideo określanych jako abstrakcyjne, a niekiedy wręcz surrealistyczne. Przyciągają uwagę m.in. ze względu na: (1) melodyjne, wpadające w ucho nazwy – przykładowo *Bombardiro Crocodilo*, *Tralalero Tralala*, *Tung Tung Tung Sahur*, *Ballerina Cappuccina* lub *Lirili Larila*, (2) wygląd, będący mieszanką atrybutów antropomorficznych, zoomorficznych w połączeniu z cechami obiektów takich jak filiżanka, kawałek drewna, arbuz i inne, (3) podkład głosowy podobny do wypowiedzi formułowanych w języku włoskim, w istocie stanowiący „zlepek” przypadkowych słów, (4) towarzyszące niekiedy remiksy popularnych utworów. Język charakterystyczny dla treści typu *brainrot* wykracza poza sferę internetu i zmienia sposób codziennej komunikacji. Jest to proces zachodzący w sposób nieświadomy – zaczyna się od obserwacji panujących trendów, następnie dochodzi do ich rozumienia, w tym do ironicznego lub prześmiewczego posługiwania się określonymi wyrażeniami i w konsekwencji stosowania ich w codziennej mowie (Bolek 2024). Sam termin *brainrot* w tłumaczeniu na język polski oznacza „zgniliznę mózgu” i odnosi się do skutków nadmiernej konsumpcji bezsensownych, trywialnych i niewymagających materiałów znajdujących się w internecie (Oxford University Press 2024). Jednym z tematów dyskusji dotyczącej negatywnych konsekwencji przyswajania opisywanych treści przy pomocy mediów społecznościowych, w szczególności TikToka, jest zdrowie psychiczne dzieci i młodzieży – grup, które już teraz przyznają, że wskutek długiego oglądania materiałów typu *brainrot* odczuwają problemy z koncentracją, a także z zasypianiem (Zinkiewicz 2025). Na problem coraz gorszej jakości snu w kontekście osób młodych zwracają uwagę Magdalena Hodalska oraz Łukasz Buksa, wskazując, iż:

aktualnie – w dobie intensywnego życia wielozadaniowego – końcowa część dnia wydaje się być dla wielu młodych ludzi jedynym czasem, w którym, leżąc w łóżku ze smartfonem, mogą doświadczyć chwilowego pocucia braku obowiązków (Hodalska i Buksa 2025: 241)

Warto w tym kontekście warto rozważyć, czy przegłądanie treści typu *brainrot* stanowi dla dzieci oraz nastolatków ucieczkę przed stresami codziennego życia.

Boomertrap – pułapki na boomerów

Obrazy generowane przez sztuczną inteligencję przyciągają również uwagę osób starszych, które stają się celem oszustw internetowych dokonywanych przy użyciu platform społecznościowych. W kontekście *boomertrap*

szczególną uwagę zwraca się na pokolenie *baby boomers*, urodzone w latach 1946–1964, stereotypowo postrzegane jako mniej doświadczone w zakresie użytkowania narzędzi technologicznych. Choć w istocie wskazuje się, iż osoby te wyróżnia inny styl przetwarzania informacji, a w związku z tym mniejsza chęć i sprawność w zakresie podążania za rozwojem technologicznym, to pandemia COVID-19 zmieniła nastawienie *baby boomers* do technologii, a przede wszystkim wymusiła zapoznanie się z jej rodzajami (Yan i Chan 2024: 91). Ponadto osoby urodzone w latach 1946–1964 kontynuują korzystanie z określonego narzędzia, jeżeli dostrzegają w tym korzyści – a ponadto mimo ograniczonej wiedzy na temat użytkowania technologii są w stanie szybko ją przyswoić w wyniku odpowiedniego przeszkolenia (Yan i Chan 2025: 91). Magdalena Hodalska i Łukasz Buksa zwracają uwagę, iż kompetencje cyfrowe *baby boomers* określanych także srebrnym pokoleniem (ang. *silver generation*) kształtowały się w dobie dynamicznego rozwoju technologicznego – który niejako wymusił dostosowanie się do panujących warunków – przez co *boomerzy* wciąż znajdują się w trybie tzw. nadążania (Hodalska i Buksa 2025: 25). Choć stają się oni coraz bardziej aktywną i świadomą korzyści społecznych, ekonomicznych i zdrowotnych grupą użytkowników internetu, to znaczna część *silver generation* nie korzysta z tego narzędzia w ten sam sposób, co inni użytkownicy. Do przyczyn tego stanu zalicza się braki w zakresie zaawansowanych kompetencji, motywacji, odpowiedniego wykształcenia, doświadczenia zawodowego, dostępności, a także koszty technologii (Hodalska i Buksa 2025: 25).

Boomertrap (ang. *trap* – pułapka) to inaczej pułapki w formie obrazków lub filmików wygenerowanych przez sztuczną inteligencję – „zastawiane” na platformach społecznościowych (w szczególności Facebooku) na *boomerów*. Ich celem jest zwiększenie zasięgów kont publikujących w większości „AI spam” poprzez zgromadzenie jak największej liczby reakcji, komentarzy oraz udostępnień. Aby stało się to możliwe, należy skutecznie przyciągnąć uwagę odbiorcy, z tego względu wykorzystuje się ludzkie skłonności do okazywania empatii oraz współczucia w sytuacji, gdy inne osoby doświadczają krzywdy. Dlatego zawartość *boomertrap* odwołuje się do tematów emocjonalnych i kontrowersyjnych takich jak religia, polityka, samotność, bieda oraz choroby. Materiały tego typu zamieszczane są zarówno na profilach o tematyce motywacyjnej, religijnej, jak i sentymentalnej – w towarzystwie innych postów przedstawiających m.in. motywy biblijne, zwierzęta, złote myśli i inne dobrze znane starszym użytkownikom mediów społecznościowych aspekty życia codziennego (m.in. robótki ręczne, gotowanie, remonty, praca). Często są to grafiki o kuriozalnym i mało autentycznym charakterze, którym towarzyszą generyczne opisy (niekiedy z oczywistymi błędami składniowymi): „Dziś bierzemy ślub, ale nikt nam nie błogosławi, bo wychowaliśmy się w domu dziecka!”, „Jestem dzieckiem z biednej rodziny, dziś są moje

urodziny, nie ma tortu, tylko ziemniaki, nikt nie życzy mi dobrze”, „Moja mama ma 80 lat. Jeśli nie przeszkadza Ci, że jest stara, po prostu się z nią przywitaj”, „Dziś mam sto dziesięć lat Mieszkam na wsi. Sama upiekłam to ciasto Chciałam je zjeść z dziećmi, ale nie ma ich już ze mną ...”. Budowanie zasięgu za pomocą zaangażowania emocjonalnego użytkowników ma służyć uwiarygodnieniu przekazu. W ten sposób łatwiejsza staje się realizacja oszustw finansowych w sieci, a także wyłudzenie danych po nawiązaniu kontaktu z osobami, które pozostawiły reakcje pod materiałami. Udostępnianie generatywnego spamu służy również rozpoznawaniu preferencji odbiorców na podstawie zgromadzonych na ich profilach danych, co z kolei umożliwia kierowanie kampanii dezinformacyjnych do najbardziej podatnych na manipulację jednostek (NASK 2025b). Dodatkowo zwiększanie zaangażowania użytkowników w odbiór treści wiąże się ze wzrostem wartości funpage’a jako towaru, stopniowo przygotowywanego do sprzedaży odbywającej się najczęściej na rynku w *dark webie*, gdzie sprzedawcami są zwykle osoby z Algierii, Indonezji, Wietnamu (Kuśmierek 2024). Po zawarciu transakcji charakter publikowanych materiałów często ulega zmianie na treści reklamowe lub dezinformacyjne.

Ze względu na: (1) niejednoznaczne – choć najczęściej motywowane potrzebą osiągnięcia korzyści finansowych lub chęcią wywołania chaosu informacyjnego – intencje twórców „AI spamu”, (2) skupiający uwagę użytkowników – niekiedy w wyniku emocjonalnej manipulacji – charakter treści, angażujący ich w interakcje z konkretnymi materiałami oraz dalsze udostępnianie, (3) abstrakcyjne, a zarazem surrealistyczne formy niektórych obrazów oraz filmików, a także (4) zniekształcone, fałszywe lub częściowo nieprawdziwe odzwierciedlenia sytuacji dziejących się w rzeczywistości społecznej – należy uznać, że materiały typu *AI slop*, *brainrot* oraz *boomertrap* zawierają elementy dezinformacji, misinformacji oraz malinformacji. Oznacza to, iż intensyfikujące smog informacyjny trendy oparte na upowszechnianiu w mediach społecznościowych treści memiczno-rozrywkowych oraz tworzeniu fałszywych lub zmanipulowanych materiałów mogą być wykorzystywane przez różne grupy osób – w tym m.in. te podzielone światopoglądowo lub też pochodzące z biedniejszych rejonów świata – do wprowadzania w błąd opinii publicznej, szerzenia teorii spiskowych i *fake newsów*, monetyzowania treści, wyłudzenia danych, a nawet wyrządzania szkód o charakterze finansowym.

Znaczenie świadomości medialnej w dobie smogu informacyjnego

Przedstawione w niniejszym tekście działania mogą prowadzić do: (1) promowania teorii spiskowych lub *fake newsów*, (2) wzrostu polaryzacji społecznej

i tworzenia się osobnych rzeczywistości informacyjnych, (3) obniżenia poziomu dyskursu publicznego i kulturowego, (4) pogorszenia jakości przekazu w *social mediach*, w tym postrzegania platform społecznościowych, jako mniej wartościowych źródeł informacji i rozrywki, (5) ograniczenia widoczności przydatnych treści (Knowski 2024). Kolejna z tzw. „wiosen AI” uświadamia potrzebę skutecznej edukacji medialnej umożliwiającej jednostkom: (1) krytyczną ocenę oddziaływania mediów na ich życie, (2) rozpoznawanie technik, którymi posługują się twórcy przekazów medialnych, (3) analizę dekodowanych materiałów, w tym ukrytych przekazów, (4) ocenę ich wartości i jakości, (5) korzystanie z mediów do celów edukacyjnych, (6) tworzenie własnych treści, (7) samodzielne wyrażanie opinii na temat instytucjonalnej działalności medialnej, w tym domaganie się jej poprawy, (8) znajomość i stosowanie zasad etyki mediów oraz (9) wykorzystywanie posiadanej wiedzy w codziennym życiu (Lee 2010). Zdaniem Agnieszki Ogonowskiej w formalną edukację medialną powinny wpisywać się zagadnienia związane z: cyberbezpieczeństwem, prawem własności intelektualnej, zarządzaniem wizerunkiem oraz danymi wrażliwymi, netykietą, etyką w nowych mediach, krytycznym myśleniem, przeciwdziałaniem dezinformacji i manipulacji w mediach (Ogonowska 2022: 20).

W obliczu zwiększającej się liczby materiałów dezinformacyjnych w mediach społecznościowych przydatnym sposobem przeciwdziałania ich oddziaływaniu może okazać się zwracanie uwagi na zawartość kont – ich nazwy, tematykę zamieszczanych postów, detale zdjęć – a w razie potrzeby korzystanie z opcji „pokaż mniej”, „ukryj post/spam” i innych dostępnych na platformach funkcjonalności. Niebagatelne znaczenie ma również tłumaczenie osobom starszym, czym jest *boomertrap*. Ponadto DiResta i Goldstein (2024) wskazują na:

- ograniczanie widoczności materiałów dezinformacyjnych przez firmy obsługujące media społecznościowe – inwestowanie w infrastrukturę wykrywającą oszustwa, testowanie rozwiązań technicznych, a także etykietowanie treści generowanych przez AI;
- minimalizowanie oddziaływania dezinformacyjnych materiałów docierających do odbiorców poprzez zaangażowane działania mediów w edukację użytkowników platform społecznościowych – m.in. tworzenie kampanii społecznych. Przydatna może okazać się metoda inokulacji, polegająca na „wszczepianiu” treści dezinformacyjnych w celu zbudowania odporności odbiorców. Jej skuteczność potwierdzają badania Rozenbeeka, które wykazały, iż „oglądanie krótkich filmów z inokulacją poprawia zdolność ludzi do identyfikowania technik manipulacji powszechnie stosowanych w dezinformacji online, zarówno w warunkach laboratoryjnych, jak i w rzeczywistym środowisku, gdzie ekspozycja na dezinformację jest powszechna” (Roozenbeek i in. 2022: 5).

- analizę przez badaczy nauk społecznych rzeczywistego wykorzystania i wpływu dezinformacyjnych treści na użytkowników mediów społecznościowych, konieczną dla określenia priorytetowych interwencji oraz działań edukacyjnych służących walce ze zmanipulowanymi informacjami.

Wraz z rozwojem GenAI coraz trudniej będzie odróżnić sfabrykowane materiały od tych, które przedstawiają prawdziwe wydarzenia i osoby. Dowód na to stanowi m.in. wideo z Willem Smithem w roli głównej, które *viralowo* rozprzestrzeniło się w internecie. Na nagraniu znajdującym się w serwisie YouTube pod nazwą *Will Smith Eating a spaghetti in 2025* popularny aktor degustuje makaron i jak mogłoby się wydawać, całość wygląda jak ujęcia nakręcone podczas pobytu w nadmorskim kurorcie – jednak w rzeczywistości krótki filmik został całkowicie wygenerowany komputerowo. Niespełna dwa lata wcześniej w sieci ukazał się podobny materiał, jednak stał się on wówczas internetowym memem, pełnym groteskowych i kardynalnych błędów popełnionych przez sztuczną inteligencję. W tamtym czasie nie przypuszczano, iż AI dorówna człowiekowi w zakresie tworzenia realistycznych materiałów – dziś, jak się jednak okazuje, potrafi:

renderować fotorealistycznych ludzi wykonujących zarówno prozaiczne, jak i fantastyczne czynności. Od topornych gagów wizualnych przeszliśmy do filmów krótkometrażowych. Cyfrowe awatary nie tylko jedzą spaghetti, ale także je gotują, serwują i występują w reklamach restauracji na poziomie gwiazdek Michelin – bez konieczności stawania przed kamerą (Inamdar 2025 [tłum. własne]).

W kontekście fotorealistycznych postaci – których problematyka nie została uwzględniona w rozważaniach wyznaczających główną oś zawartego w artykule wyводу – warto wspomnieć o tzw. wirtualnych influencerach obecnych w mediach społecznościowych w formie awatarów. Awatar to animacja stworzona przy pomocy trójwymiarowej grafiki (3D), przedstawiająca ludzi oraz pozaludzkie reprezentacje w środowisku cyfrowym (European Commission 2025). W dobie *influencer marketingu* – który opiera się na współpracy marek z twórcami treści, wyróżniającymi się dużą liczbą obserwatorów w mediach społecznościowych – coraz popularniejszą praktyką staje się tworzenie wirtualnych person pełniących w świecie wirtualnym te same role, co ludzcy influencerzy. Internetowe sławy za sprawą swoich zasięgów zyskują nie tylko status lidera opinii, ale także zaufanego eksperta (Gomes, Marques i Dias 2022). Ich pozycja wykorzystywana jest przez marki do promowania produktów i usług, aby zbudować określony wizerunek firmy i zwiększyć wyniki sprzedażowe. Istotne problemy pojawiają się, gdy influencerzy dopuszczają się działań uchodzących za skandaliczne lub

powszechnie nieakceptowalne, co naraża przedsiębiorstwa na realne ryzyko strat. Dlatego też w obliczu istniejących zagrożeń związanych z nieprzewidywalnością ludzkich zachowań coraz popularniejsze staje się tworzenie przy pomocy sztucznej inteligencji zdigitalizowanych person podlegających pełnej kontroli. Wirtualni influencerzy w formie awatarów przedstawiani są jako:

osoby żyjące życiem podobnym do życia influencerów, cieszące się przyjaciółmi, hobby, zobowiązaniami [...]. Wirtualni influencerzy są jak bohaterowie seriali telewizyjnych lub mangi, którzy fascynują swoich odbiorców, którym ożywiają swoje przygody [...]. Ludzie mogą oglądać ich teledyski lub koncerty, oglądać ich śniadania przed wyjściem na imprezę, chodzić na przymiarki lub uczestniczyć w pokazach mody (Allal-Chérif, Puertas i Carracedo 2024: 123113 [tłum. własne]).

Z opisywanych możliwości, oferowanych przez technologię AI korzystają marki skupione wokół handlu detalicznego, branży luksusowej, modowej, kosmetycznej, motoryzacyjnej oraz turystycznej – w tym m.in. Calvin Klein, Dior, Prada, Samsung, Adidas, czy BMW.

Jednym z przykładów wirtualnej osoby jest Lil Miquela, której konto na Instagramie śledzi ponad 2,4 mln obserwujących. Sam awatar został stworzony na podobieństwo młodej kobiety o indiańsko-brazylijskim pochodzeniu przez amerykańską agencję interaktywną Brud. Lil Miquela zdaje się wykonywać wszystkie te same aktywności, co prawdziwe influencerki – m.in. śpiewa, nagrywa teledyski, bierze udział w sesjach zdjęciowych i kampaniach reklamowych, udziela wywiadów, spędza czas ze znajomymi, umawia się na randki, podróżuje, dzieli się swoimi preferencjami kulinarnymi, odtwarza popularne trendy internetowe – z kolei w 2018 r. w towarzystwie najpopularniejszych w skali globu osób ze świata polityki i show biznesu znalazła się w rankingu dwudziestu pięciu najbardziej wpływowych osobowości internetowych magazynu „Time” (Fortuna 2022). Obecność wirtualnych influencerów w świecie mediów społecznościowych – a także ich zaangażowanie w działania marketingowe uznanych marek – wskazują na powolne zacieranie się granic między światem offline i online. W szczególności osoby młode – dla których internet stanowi medium towarzyszące od wczesnych momentów życia – mogą nie dostrzegać wyraźnych linii podziału. Biorąc pod uwagę brak ujednoliconych kodeksów etycznych odnoszących się do kwestii wykorzystywania wirtualnych person w marketingu internetowym, pojawiają się obawy o przejrzystość tego rodzaju działalności, ryzyka związane z manipulowaniem faktami – w tym o dokonywanie nadużyć – a także promowanie nierealistycznych standardów piękna (Moustakas i in. 2020). Zarówno świadomość istnienia zaprezentowanego zjawiska, jak i ogólna orientacja w zakresie treści obecnych w mediach społecznościowych są potrzebne do samodzielnego wyznaczania granic między prawdą a cyfrową fikcją.

Podsumowanie

W niniejszym tekście o charakterze teoretyczno-przeglądowym zaprezentowano nowe trendy obecne w świecie mediów społecznościowych – przyczyniające się do intensyfikacji zjawiska smogu informacyjnego – których powstawaniu sprzyja rozwój generatywnej sztucznej inteligencji. Ukazano tym samym, jak wygląda: (1) dezinformacja polityczna z wykorzystaniem GenAI, w tym strategia *flooding the zone*, a także (2) *AI slop*, *brainrot* oraz *boomertrap*, składające się w obydwóch przypadkach na pojęcie „AI spamu”. Jednocześnie opisano ich konkretne przykłady, określono charakter, jaki mogą przybierać treści o społecznie szkodliwym charakterze, uwzględniając przy tym grupy odbiorców w szczególny sposób narażone na negatywne skutki konsumpcji tego rodzaju materiałów obecnych w *social mediach*. „AI spam” to szeroka kategoria, do której można zaklasyfikować materiały zawierające elementy dezinformacji, misinformacji oraz malinformacji, w tym przede wszystkim: (1) niskiej jakości treści generowane intencjonalnie celem wywołania chaosu w sferze komunikowania politycznego lub dla osiągnięcia korzyści finansowych – poprzez wykorzystanie abstrakcyjnej estetyki lub emocjonalnej manipulacji, przyciągających uwagę użytkowników, a także (2) memiczno-rozrywkowe formy, *viralowo* tworzone i udostępniane przez użytkowników mediów społecznościowych. Wyzwania związane z obecnością materiałów dezinformacyjnych w przestrzeni mediów społecznościowych wynikają przede wszystkim z: (1) coraz łatwiejszego dostępu do narzędzi generatywnej sztucznej inteligencji, (2) polityki platform społecznościowych i działania algorytmów zwiększających widoczność „AI śmieci”, a także (3) udostępniania ich przez polityków oraz bliskich im współpracowników, wyrażających tym samym ciche przyzwolenie na wykorzystywanie „generatywnego spamu” w wyścigach politycznych. Zaprezentowane rozważania dowodzą również o znaczeniu oraz aktualności pojęcia smogu informacyjnego, który – jak zostało wspomniane – nie tylko oddala użytkowników od prawdy, ale również zniechęca do jej poszukiwań. Jak zostało wskazane w tekście, przeciwdziałanie szkodliwym trendom generowanym przez sztuczną inteligencję – a tym samym dezinformacji – wymaga zaangażowania wielu aktorów, począwszy od użytkowników, administratorów platform społecznościowych, aż po media, organizacje pozarządowe, badaczy i politycznych decydentów. W podejmowanych działaniach instytucjonalnych i pozainstytucjonalnych szczególną rolę powinna odgrywać edukacja medialna promująca krytyczne myślenie, które:

w czasach chaosu informacyjnego i dezinformacji, polaryzacji społecznej i dekonstrukcji norm społecznych dyskusji mogłoby stanowić przeciwwagę dla działań nadawców próbujących wprowadzić do dyskursu treści o charakterze

rozrywkowym, zagrażającym jakości nie tylko debaty publicznej, ale w ogóle procesom społecznym (Kaczmarek-Śliwińska 2023: 70).

Należy jednocześnie przypuszczać, iż wyzwań związanych z rozprzestrzenianiem się materiałów o charakterze dezinformacyjnym będzie przybywać wraz z postępami dokonującymi się w obszarze zaawansowanych technologii. W związku z tym komunikolodzy oraz medioznawcy będą odgrywać coraz większą rolę zarówno w procesie identyfikacji nowych zjawisk związanych z przyrostem treści generowanych przez AI na platformach społecznościowych, jak i ich konsekwencji. Niniejszy artykuł stanowi jeden z kroków we wskazanym kierunku, mający na celu uświadomienie czytelnikom, jak z pozoru niewinne trendy obecne w mediach społecznościowych mogą kształtować postrzeganie istotnych w dyskursie publicznym tematów.

Bibliografia

- Allal-Chérif Oihab, Puertas Rosa, Carracedo Patricia. 2024. „Intelligent influencer marketing: how AI powered virtual influencers outperform human influencers”. *Technological Forecast Social Change* t. 2000. 123113. <https://doi.org/10.1016/j.techfore.2023.123113> (dostęp: 04.11.2025).
- Banh Leonardo, Strobel Gero. 2023. „Generative artificial intelligence”. *Electron Markets* nr 33(63). <https://doi.org/10.1007/s12525-023-00680-1> (dostęp: 04.11.2025).
- Bojakowski Jakub. 2025. Fenomen chałkonia. „Okazał się bestią, której potęgi nie doceniliśmy”. <https://www.wprost.pl/kraj/11907430/chalkon-hitem-sieci-ostrzezenie-przed-oszustwem-ai.html> (dostęp: 04.11.2025).
- Bolek Maria. 2024. Czym jest brainrot? Co znaczy skibidi?. <https://www.o-jezyku.pl/2024/11/13/czym-jest-brainrot-co-znaczy-skibidi/> (dostęp: 04.11.2025).
- Demagog. 2025a. Kobiety popierające Nawrockiego? Zdjęcie wygenerowała sztuczna inteligencja. <https://demagog.org.pl/wypowiedzi/kobiety-popierajace-nawrockiego-zdjecie-wygenerowala-sztuczna-inteligencja/> (dostęp: 04.11.2025).
- Demagog. 2025b. To zdjęcie przywódców ze szczytu w Białym Domu? Uważaj na obrazek AI. https://demagog.org.pl/fake_news/to-zdjecie-przywodcow-ze-szczytu-w-bialym-domu-uwazaj-na-obrazek-ai/ (dostęp: 04.11.2025).
- Demagog. 2025c. Prezydent Nawrocki całuje swoją doradczynię? To przeróbka stworzona przez AI. https://demagog.org.pl/fake_news/prezydent-nawrocki-caluje-swoja-doradczynnie-to-przerobka-stworzona-przez-ai/ (dostęp: 04.11.2025).
- Diki. Slop. <https://www.diki.pl/slownik-angielskiego?q=slop> (dostęp: 04.11.2025).
- DiResta Renée, Goldstein Josh. A. 2024. „How spammers and scammers leverage AI-generated images on Facebook for audience growth”. *Harvard Kennedy School (HKS) Misinformation Review* nr 5(4). <https://doi.org/10.48550/arXiv.2403.12838> (dostęp: 04.11.2025).

- Donald. 2025. Chałkoń został sprawdzony przez służby, czy nie jest próbą wpływu na wybory. <https://www.donald.pl/artykuly/XzWmwp6n/chalkon-zostal-sprawdzony-przez-sluzby-czy-nie-jest-proba-wplywu-na-wybory> (dostęp: 04.11.2025).
- Epstein Ziv, Hertzmann Aaron. 2023. „Investigators of Human Creativity, Art and the science of generative AI. Understanding shifts in creative work will help guide AI’s impact on the media ecosystem”. *Science* nr 380(6650). <https://doi.org/10.1126/science.adh4451> (dostęp: 04.11.2025).
- European Commission. 2025. Avatars. <https://digital-strategy.ec.europa.eu/en/policies/virtual-worlds-avatars> (dostęp: 04.11.2025).
- Fortuna Piotr. 2022. „Miquela versus Real Life”. *Teorie i Praktyki Kultury Wizualnej* nr 33. <https://doi.org/10.36854/widok/2022.33.2581> (dostęp: 04.11.2025).
- Gomes Marina Alexandra, Marques Susana, Dias Alvaro Lopes. 2022. „The impact of digital influencers’ characteristics on purchase intention of fashion products”. *Journal of Global Fashion Marketing* nr 13(3). 187–204. <https://doi.org/10.1080/020932685.2022.2039263> (dostęp: 04.11.2025).
- Desta Haileselassie Hagos, Battle Rick, Rawat Danda B. 2024. „Recent Advances in Generative AI and Large Language Models: Current Status, Challenges, and Perspectives”. *IEEE Transactions on Artificial Intelligence* nr 5(12). 5873–5893. <https://doi.org/10.48550/arXiv.2407.14962> (dostęp: 04.11.2025).
- Hodalska Magdalena, Buksa Łukasz. 2025. Przewinięci: Smartfon w polskiej codzienności. Przeglądanie telefonu w świetle badań empirycznych. Kraków.
- Illing Sean. 2020. „Flood the zone with shit”: How misinformation overwhelmed our democracy. <https://www.vox.com/policy-and-politics/2020/1/16/20991816/impeachment-trial-trump-bannon-misinformation> (dostęp: 04.11.2025).
- Inamdar Anil. 2025. Will Smith Eating Spaghetti: Why We Keep Underestimating AI. <https://medium.com/@anilnamdar/will-smith-eating-spaghetti-why-we-keep-underestimating-ai-beb65092ff85> (dostęp: 04.11.2025).
- Instytut Kościuszki. b.r. Dezinformacja, misinformation, malinformation i fake news: Rozszyfrowując zagrożenia informacyjne. Gracia Sumariva Reyes. <https://ik.org.pl/2023/11/24/dezinformacja-misinformation-malinformation-i-fake-news-rozszyfrowujac-zagrozenia-informacyjne-gracia-sumariva-reyes/> (dostęp: 04.11.2025).
- Januszko-Szakiel Aneta. 2010. Dezinformacja jako narzędzie medialnej manipulacji świadomością. W: *Manipulacja. Pedagogiczno-społeczne aspekty*. Joanna Akseman (red.). Kraków. 209–216.
- Kaczmarek-Śliwińska Monika. 2023. Edukacja medialna w dobie obecnej przestrzeni medialnej. W: *Edukacja medialna – zasady funkcjonowania w świecie nowych mediów*. Klaudia Cymanow-Sosin (red.). Kraków. 63–73.
- Kapelner Zsolt. 2025. „Flooding the zone” and the politics of attention. <https://justice-everywhere.org/general/flooding-the-zone-and-the-politics-of-attention/> (dostęp: 04.11.2025).
- Knowski Wiktor. 2024. Co to jest AI Slop? Ściek, w którym utoniemy my i nasza kultura. <https://innpoland.pl/209038,ai-slop-nie-tylko-meczy-ale-i-zagraza> (dostęp: 04.11.2025).

- Koebler Jason. 2024. Where Facebook's AI Slop Comes From. <https://www.404media.co/where-facebooks-ai-slop-comes-from/> (dostęp: 04.11.2025).
- Konopczyński Filip. 2025. Syntetyczny Potop: Deepfake, AI slop i koniec wspólnej rzeczywistości. <https://krytykapolityczna.pl/swiat/wplyw-generatywnej-ai-na-media-deepfake-dezinformacja-energia/> (dostęp: 04.11.2025).
- Krever Mick, Salem Mostafa. 2025. „Trump Gaza is finally here!”: US president promotes Gaza plan in AI video”. <https://edition.cnn.com/2025/02/26/world/trump-promotes-gaza-plan-ai-video-intl> (dostęp: 04.11.2025).
- Kuśmierk Malwina. 2024. Czym są AI slopy i boomertrapy – jak je rozpoznać? Wyjaśniamy. <https://spidersweb.pl/2024/07/ai-slopy-boomertrapy-poradnik-co-to.html> (dostęp: 04.11.2025).
- Lee Alice Y.L. 2010. „Media Education: Definitions, Approaches and Development around the Globe”. *New Horizons in Education* nr 58(3). <https://files.eric.ed.gov/fulltext/EJ966655.pdf> (dostęp: 04.11.2025).
- Łódzki Bartłomiej. 2017. „Fake news» – dezinformacja w mediach internetowych i formy jej zwalczania w przestrzeni międzynarodowej”. *Polityka i Społeczeństwo* nr 4(15). 19–30. <https://doi.org/10.15584/polispol.2017.4.2> (dostęp: 04.11.2025).
- Mahdawi Arwa. 2025. AI-generated „slop” is slowly killing the internet, so why is nobody trying to stop it?. <https://www.theguardian.com/global/commentis-free/2025/jan/08/ai-generated-slop-slowly-killing-internet-nobody-trying-to-stop-it> (dostęp: 04.11.2025).
- Malik Nesrine. 2025. With „AI slop” distorting our reality, the world is sleepwalking into disaster. <https://www.theguardian.com/commentisfree/2025/apr/21/ai-slop-artificial-intelligence-social-media> (dostęp: 04.11.2025).
- Moustakas Evangelos, Lamba Nishtha, Mahmoud Dina, Ranganathan Chandrasekaran. 2020. „Blurring lines between fiction and reality: Perspectives of experts on marketing effectiveness of virtual influencers”. *International Conference on Cyber Security and Protection of Digital Services (Cyber Security)*. 1–6. <https://doi.org/10.1109/CyberSecurity49315.2020.9138861> (dostęp: 04.11.2025).
- NASK. 2025a. Dezinformacja, misinformacja i malinformacja – czym się różnią?. <https://www.nask.pl/magazyn/dezinformacja-misinformacja-i-malinformacja-czym-sie-roznia> (dostęp: 21.02.2026).
- NASK. 2025b. Chałkon z farmy lajków. Humor w social mediach i jego ciemne oblicze. <https://www.nask.pl/aktualnosci/chalkon-z-farmy-lajkow-humor-w-social-mediach-i-jego-ciemne-oblicze> (dostęp: 04.11.2025).
- Ogonowska Agnieszka. 2022. „Kompetencje medialne i cyfrowe dzieci w wieku 3–16 lat. W stronę medioznawstwa stosowanego”. *Studia De Cultura* nr 14(4). 17–29. <https://doi.org/10.24917/20837275.14.4.2> (dostęp: 04.11.2025).
- Oxford University Press. 2024. „Brain rot” named Oxford Word of the Year 2024. <https://corp.oup.com/news/brain-rot-named-oxford-word-of-the-year-2024/> (dostęp: 04.11.2025).
- Przegalińska Aleksandra. 2025. Kojarzycie termin «flooding the zone»? Bo jesteśmy właśnie świadkami pełnoskalowego zastosowania tej taktyki w globalnej polityce... <https://www.linkedin.com/posts/aleksandra-przegalinska>

- 5b17125_kojarzycie-termin-flooding-the-zone-bo-activity-7325163101857751040-vCB1/?originalSubdomain=pl (dostęp: 04.11.2025).
- Roozenbeek Jan, van der Linden Sander, Goldberg Beth, Rathje Steve, Lewandowsky Stephan. 2022. „Psychological inoculation improves resilience against misinformation on social media”. *Science Advances* nr 8(34). <https://doi.org/10.1126/sciadv.abo6254> (dostęp: 04.11.2025).
- Sarowski Łukasz. 2017. „Od Internetu Web 1.0 do Internetu Web 4.0 – ewolucja form przestrzeni komunikacyjnych w globalnej sieci”. *Rozprawy Społeczne* nr 11(1). 32–39. <https://rozprawyspoleczne.edu.pl/OD-INTERNETU-WEB-1-0-DO-INTERNETU-WEB-4-0-EWOLUCJA-FORM-nPRZESTRZENI-KOMUNIKACYJNYCH,110907,0,1.html> (dostęp: 04.11.2025).
- Sohyun Park, Someen Park, Jaehoon Kim, Kyungsik Han. 2024. „Exploring the Impact of AI-Generated Images on Political News Perception and Understanding”. *Companion of the 2024 Computer-Supported Cooperative Work and Social Computing (CSCW Companion '24)*. 565–571. <https://doi.org/10.1145/3678884.3681907> (dostęp: 04.11.2025).
- Stasiuk-Krajewska Karina. 2023. „Dezinformacja. Próba ujęcia dyskursywnego”. *Media Biznes Kultura* nr 1(14). 55–72. <https://czasopisma.bg.ug.edu.pl/index.php/MBK/article/view/9083> (dostęp: 04.11.2025).
- Szpunar Magdalena. 2017. *Imperializm kulturowy internetu*. Kraków.
- Tadeusiewicz Ryszard. 1999. „Smog informacyjny”. *Prace Komisji Zagrożeń Cywilizacyjnych* t. 2. 97–107. https://www.academia.edu/30020848/Smog_Informacyjny (dostęp: 04.11.2025).
- Tangermann Victor. 2024. *Investigation Finds Actual Source of All That AI Slop on Facebook*. <https://futurism.com/the-byte/source-ai-slop-facebook> (dostęp: 04.11.2025).
- The Guardian. 2025. John Oliver on AI slop: „Some of this stuff is potentially very dangerous”. <https://www.theguardian.com/tv-and-radio/2025/jun/23/john-oliver-last-week-tonight-recap-ai#:~:text=Oliver%20spoke%20about%20the%20many,up%20new%20disasters%2C%E2%80%9D%20he%20said> (04.11.2025).
- The Media Manipulation Casebook. b.r. Distributed amplification. <https://media-manipulation.org/definitions/distributed-amplification/> (dostęp: 04.11.2025).
- Vaccari Cristian, Chadwick Andrew. 2020. *Deepfakes and Disinformation: Exploring the Impact of Synthetic Political Video on Deception, Uncertainty, and Trust in News*. *Social Media + Society* nr 6(1). <https://doi.org/10.1177/2056305120903408> (dostęp: 04.11.2025).
- Walker Christina P., Schiff Daniel S., Schiff Kaylyn J. 2024. „Merging AI Incidents Research with Political Misinformation Research: Introducing the Political Deepfakes Incidents Database”. *IAAI-24, EAAI-24, AAAI-24 Student Abstracts, Undergraduate Consortium and Demonstrations* nr 38(21). 23053–23058. <https://doi.org/10.1609/aaai.v38i21.30349> (dostęp: 04.11.2025).
- Weidinger Laura i in. 2022. „Taxonomy of Risks posed by Language Models”. *2022 ACM Conference on Fairness, Accountability, and Transparency (FAccT '22)*. 214–229. <https://doi.org/10.1145/3531146.3533088> (dostęp: 04.11.2025).

- World Economic Forum. 2024. „The Global Risks Report 2024. 19th Edition”. IN-SIGHT REPORT. https://www3.weforum.org/docs/WEF_The_Global_Risks_Report_2024.pdf (dostęp: 04.11.2025).
- Yan Law K., Chan Janelle. 2024. „Understanding Baby Boomers’ Psychological Contradiction in Adopting Automatic Technology: An Application of Protection Motivation Theory”. *Journal of China Tourism Research* nr 21(1). 87–109. <https://doi.org/10.1080/19388160.2024.2311174> (dostęp: 04.11.2025).
- Zinkiewicz Marta. 2025. Ballerina Cappuccina i Tralalero Tralala. Tłumaczymy „brainrot” – trend, który pokochały nastolatki. <https://innpoland.pl/212714,italian-brainrot-dziwny-trend-pokolenia-alfa-ktory-podbil-tiktoka> (dostęp: 04.11.2025).

Streszczenie

Prezentowany artykuł o charakterze teoretyczno-przeglądowym ma na celu identyfikację popularnych trendów internetowych obecnych w przestrzeni mediów społecznościowych za sprawą generatywnej sztucznej inteligencji (ang. *Generative Artificial Intelligence* – GAI, GenAI) umożliwiającej tworzenie przy – z pozoru – niewielkim koszcie dużej liczby materiałów, które następnie *viralowo* rozprzestrzeniają się w internecie. Na przykładzie zjawisk takich jak m.in. *flooding the zone*, *AI slop*, *brainrot* oraz *boomertrap* zaprezentowano nowe formy dezinformacji, misinformacji i malinformacji, składające się na zjawisko smogu informacyjnego, z którymi mają styczność użytkownicy platform społecznościowych. Opisano jednocześnie ich konsekwencje o charakterze społeczno-kulturowym. W tekście przedstawiono również pokrewne trendy internetowe, wykazujące potencjał do kształtowania komunikatów obecnych w dyskursie politycznym, koncentrując się równocześnie na możliwych przyczynach ich występowania oraz potencjalnych skutkach obecności materiałów tworzonych przez sztuczną inteligencję w mediach społecznościowych.

Artificial Intelligence-generated trends: On information smog in social media based on selected examples of “AI Spam”

Summary

This theoretical review article aims to identify popular internet trends circulating on social media that have emerged through the use of generative artificial intelligence (GAI, GenAI). This technology makes it possible to produce a vast amount of content at a seemingly low cost, which then spreads virally online. Drawing on examples such as *flooding the zone*, *AI slop*, *brainrot*, and *boomertrap*, the article presents new forms of disinformation, misinformation, and malinformation that contribute to the phenomenon of information smog encountered by social media users. It also discusses their sociocultural consequences. In addition, the article highlights related internet trends that may shape messages circulating in political discourse, while also addressing the possible causes of their emergence and the potential consequences of AI-generated content in social media environments.

Słowa kluczowe: *AI slop*, *boomertrap*, *flooding the zone*, *brainrot*, dezinformacja

Keywords: *AI slop*, *boomertrap*, *flooding the zone*, *brainrot*, disinformation

Aleksandra Skrzypiec – doktorantka w Szkole Doktorskiej Nauk Społecznych UJ, kształcąca się w dyscyplinie nauki o komunikacji społecznej i mediach. Obszar badawczych zainteresowań autorki obejmuje zaawansowane technologie w komunikowaniu, ze szczególnym uwzględnieniem interakcji oraz relacji zachodzących na linii człowiek–maszyna. Na co dzień zawodowo związana z Urzędem Marszałkowskim Województwa Małopolskiego w Krakowie. W drugiej połowie 2025 ukazały się następujące publikacje autorki: „Korzystanie z generatywnej sztucznej inteligencji w codziennym życiu na przykładzie chatbota GPT” [w:] Praktyki medialne w zdigitalizowanym świecie; „Human-Machine Communication jako nowy obszar badań w nauce o komunikacji społecznej i mediach”. Media i Społeczeństwo nr 22(1/z. 2).

