

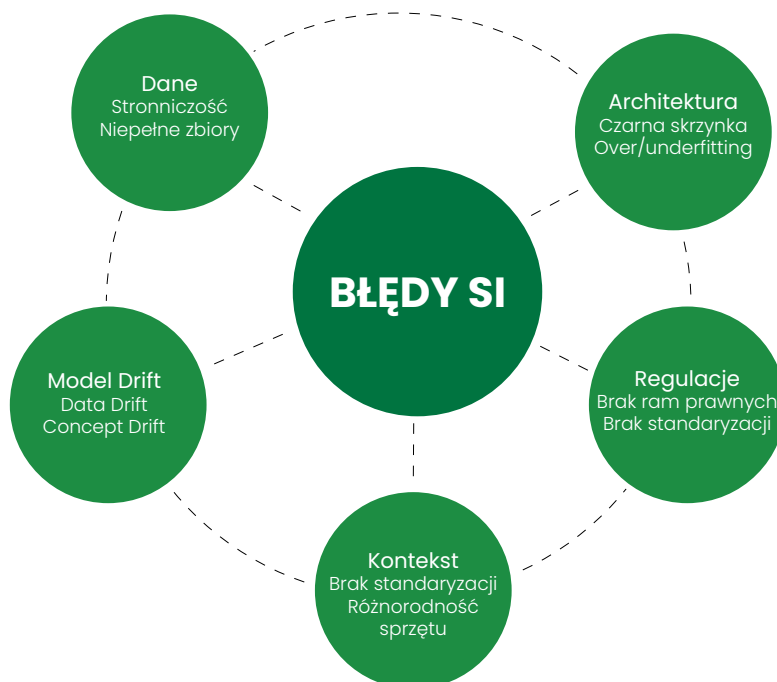
CZY SI SIĘ MYLI? O OGRANICZENIACH ALGORYTMÓW W MEDYCYNIE WETERYNARYJNEJ

Maciej Gad¹, Monika Szymańska²

¹PolyversAI, ²Libraxis AI

Sztuczna inteligencja nie nadchodzi – ona już tu jest. Ciąco, ale stanowczo zmienia codzienność klinik weterynaryjnych. W ciągu ostatniej dekady przeszła drogę od prototypów i systemów eksperckich do rzeczywistych narzędzi wspierających diagnostykę, terapię i dokumentację. Modele głębokiego uczenia analizują dziś obrazy RTG, przewidują przebieg chorób, interpretują zachowania pacjentów, generują automatyczne raporty (4, 8). Medycyna weterynaryjna staje się coraz bardziej cyfrowa – i coraz bardziej złożona (3). To jednak nie będzie opowieść o technologii, która wszystko rozwiązuje. Wręcz przeciwnie – zaczynamy tam, gdzie narracje o SI zwykle się kończą: gdy algorytm zawodzi, nie rozumie kontekstu, a jego „pewność” okazuje się iluzją. Tam, gdzie kończy się liczbowy świat modelu, a zaczyna realna, nieprzewidywalna rzeczywistość pacjenta – decyzję podejmuje nie maszyna, lecz lekarz. Z całym ciężarem odpowiedzialności i sztuką dostrzegania tego, czego nie da się przeliczyć (19). Maszyna jeszcze długo nie zastąpi lekarza, niewątpliwie jednak może go wspierać. Jest to możliwe, jeśli lekarz pozna jej naturę i dostrzeże ją jako największą zdobycz techniczną tego wieku, spek-

Ryc. 1. Źródła błędów SI w diagnostyce weterynaryjnej.



Schemat przedstawia główne źródła potencjalnych błędów w systemach SI stosowanych w diagnostyce weterynaryjnej: dane, architekturę, kontekst, model drift oraz brak ram regulacyjnych. Każde z tych źródeł może prowadzić do odmiennych błędów diagnostycznych.



takularnie zmieniającą obraz świata, niezależnie od obszaru adaptacji.

Nie jest to zatem tekst przeciwko sztucznej inteligencji – to tekst za mądrzejszym z nią sojuszem. To też apel: sztuczna inteligencja nie jest ani zagrożeniem, ani wyrocznią. SI to asystent – wspierający, szybki, precyzyjny, jednak tylko wtedy, gdy działa pod kontrolą lekarza.

Źródła błędów w algorytmach

Sztuczna inteligencja nie zawodzi przypadkowo. Jej pomyłki nie są ani irracjonalne, ani tajemnicze – choć często bywają zaskakujące. Każda decyzja algorytmu jest odbiciem danych, na których został wychowany oraz architektury, która go ukształtowała. To nie magia – to matematyka.

Źródła błędów modeli SI stosowanych w medycynie weterynaryjnej można pogrupować w trzy obszary: dane, architektura i kontekst (ryc. 1). W każdej z nich kryją się mechanizmy, które, jeśli zostaną niezauważone, mogą prowadzić do błędnych diagnoz, niewłaściwego leczenia i decyzji obarczonych ryzykiem.

Dane: „To, czym karmimy model, karmi jego błędy”

Dane uczące to fundament każdego modelu SI. Jeśli są niepełne, stronicze al-

bo źle jakości – to nie jest błąd systemu, to błąd w jego wychowaniu. Zasada „garbage in, garbage out” nie zna wyjątków. W tej sekcji omówimy tę cechę na przykładzie zjawiska stroniczości danych (nierównowagi w próbie, z ang. – data bias) (2).

Algorytmy uczą się tylko tego, co im pokażemy – i dokładnie tak, jak im to pokażemy. Na przykład: jeśli większość danych pochodzi z nowoczesnych, dobrze wyposażonych klinik w dużych miastach, model nauczy się rozpoznawać choroby głównie w takich warunkach. Pacjenci ze schronisk, lecznic wiejskich czy mniejszych placówek mogą być diagnozowani mniej trafnie – nie dlatego, że są mniej ważni, lecz dlatego, że nie byli dostatecznie reprezentowani w zbiorze uczącym. To zdefiniowane wcześniej zjawisko, czyli data bias, prowadzi do tworzenia modeli dostrojonych do dominujących warunków – tych najczęściej reprezentowanych w danych treningowych. W efekcie model działa dobrze tam, gdzie ma doświadczenie, ale traci skuteczność w przypadku rzadszych gatunków, nietypowych środowisk czy schorzeń spoza głównego nurtu danych (2).

W medycynie weterynaryjnej problem pogłębia się już na poziomie liczby przypadków (3). W porównaniu do medycyny ludzkiej, dane kliniczne to często

Does AI Make Mistakes? On The Limitations of Algorithms in Veterinary Medicine

Artificial intelligence (AI) is increasingly applied in veterinary medicine, offering powerful diagnostic and decision-support tools. However, AI algorithms are not infallible – they have inherent limitations and can make mistakes with serious clinical consequences. This article, a continuation of “Artificial intelligence in veterinary medicine – from basics to practice”, explores when and why AI errs in veterinary contexts. We analyze constraints related to data quality and model architecture, highlight examples of algorithmic failures in clinical conditions, and discuss issues of responsibility, ethics, and systemic barriers to AI implementation. Understanding these limitations is crucial for the safe, ethical, and effective integration of AI as an aid – rather than a replacement – in veterinary practice.

Keywords: artificial intelligence; veterinary medicine; limitations; ethics; clinical decision-making.

mozaika niepełnych historii chorób, opisów tworzonych przez lekarzy o różnym poziomie doświadczenia. Model nie wie, kto był ekspertem, a kto rezydentem – traktuje wszystkie etykiety jednakowo.

Dodatkowo, wiele algorytmów trenowanych jest na danych z jednego źródła – jednej kliniki, jednego aparatu diagnostycznego, jednej populacji. Taki model może świetnie radzić sobie z przypadkami podobnymi do tych, które już widział. Jednak gdy trafi na pacjenta, którego objawy nie pasują do żadnego znanego wzorca – zaczyna zgadywać (2). To zjawisko, znane jako halucynacja modelu, może być nieszkodliwe, ale bywa też niebezpieczne (1).

Dlatego tak ważne jest, by modele trenować na zróżnicowanych danych, obejmujących różne gatunki, środowiska, aparaty i formy dokumentacji. Tylko wtedy każdy pacjent, niezależnie od pochodzenia, rasy czy miejsca badania, będzie miał szansę na równie dobrą diagnozę.

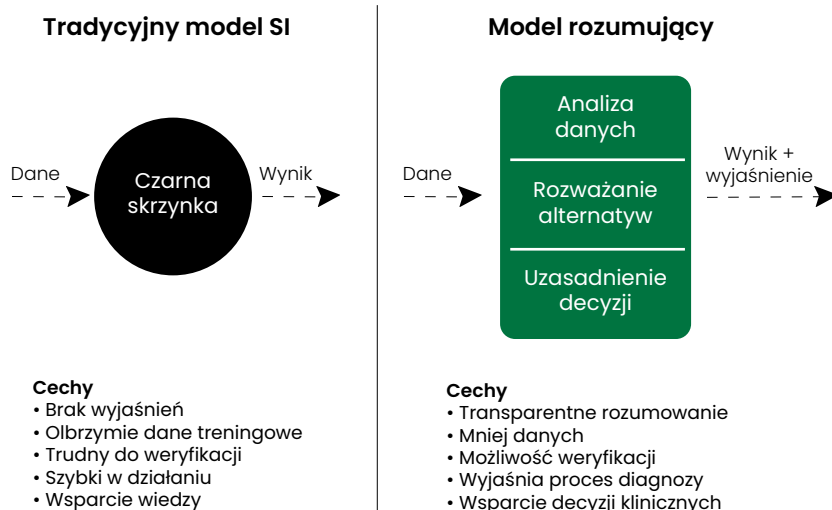
Architektura: „Co widzi model, ale nie potrafi połączyć?”

Drugim źródłem błędów SI jest konstrukcja modelu. Algorytmy głębokiego uczenia to potężne narzędzia, ale mają swoje ograniczenia. Mogą być zbyt płytkie, zbyt głębokie, zbyt sztywne albo zbyt elastyczne. Mogą uczyć się na pamięć – zamiast rozumieć. Zjawisko przetrenowania (ang. overfitting) oznacza, że model zna każdy piksel obrazu treningowego, ale nie rozpoznaje wzorca, jeśli zmieni się na przykład kontrast czy rozdzielczość. Nietrenowany model (ang. underfitting) może natomiast przeoczyć nawet wyraźne patologie, ponieważ nie wychwytuje zależności między objawami a rozpoznaniem.

Kolejnym wyzwaniem jest tzw. czarna skrzynka (ang. black box), czyli sytuacja, w której model generuje wynik, ale nie wiadomo, jaką drogą doszedł do decyzji. Brak przejrzystości mechanizmów decyzyjnych oznacza, że lekarz otrzymuje klasyfikację lub predykcję bez możliwości zrozumienia, które cechy pacjenta lub dane wejściowe miały kluczowe znaczenie. Jak podkreślają Bouchemla i wsp. (8), „najlepiej działające algorytmy są często najmniej przejrzyste, podczas gdy te dające się łatwo interpretować – jak drzewa decyzyjne – bywają mniej dokładne”. Ten paradoks między skutecznością a przejrzystością zmusza do ostrożności i odpowiedzialności – zwłaszcza, gdy stawką jest zdrowie lub życie pacjenta.

Niektóre podejścia próbują ten problem złagodzić, wprowadzając koncepcję wytłumaczalnej SI (ang. explainable AI,

Ryc. 2. Tradycyjna SI vs. Modele rozumujące.



Porównanie tradycyjnych modeli SI typu „czarna skrzynka” z modelami rozumującymi. Model rozumujący oferuje nie tylko wynik, ale również przejrzyste wyjaśnienie procesu decyzyjnego, co zwiększa zaufanie i możliwość weryfikacji diagnozy przez lekarza weterynarii.

XAI), czyli systemów, które nie tylko dają odpowiedź, ale też wskazują, dlaczego ją wygenerowały. Takie rozwiązania mogą zwiększać zaufanie i kontrolę nad działaniem algorytmu. Nadal jednak pozostają wyzwaniem technologicznym i nie zawsze osiągają poziom skuteczności klasycznych modeli „black box”. Znaczącą poprawę wytłumaczalności działania algorytmów przynoszą obecnie modele SI rozumujące (ang. reasoning, thinking), pokazujące użytkownikowi, nie tylko końcową odpowiedź, ale cały sposób dotarcia do niej (ryc. 2). Modele te zostaną omówione w dalszych częściach cyklu.

Warto pamiętać, że nawet ten sam model może zachowywać się inaczej w zależności od warunków technicznych wdrożenia. Zmiana wersji oprogramowania, sprzętu czy formatu danych wejściowych może znacząco wpłynąć na działanie SI – mimo że formalnie używamy tego samego systemu. Algorytmy nie funkcjonują w próżni. Są wrażliwe na każdy detal środowiska, w którym pracują. A w medycynie weterynaryjnej, gdzie różnorodność sprzętu, słownictwa i praktyk jest normą, takie niuanse mają realne konsekwencje kliniczne.

Kontekst: „Wszystko, o czym model nie wie, że nie wie”

Trzecim, często najtrudniejszym do uchwycenia źródłem błędów SI są czynniki zewnętrzne – czyli to, co dzieje

się poza modelem: niespójność opisów badań, brak standaryzacji, różnice sprzętowe i praktyki kliniczne. To nie są wady algorytmu – to granice jego rozumienia. W medycynie weterynaryjnej nie istnieją jednolite standardy – ani w protokołach badań, ani w dokumentacji. To, co w jednej lecznicy opisane jest jako „przewlekły kaszel”, gdzie indziej może być „kaszelem od dwóch tygodni z podejrzeniem zapalenia tchawicy”. Dla lekarza to naturalna różnica językowa – dla SI to dwa różne przypadki. Podobnie jest z obrazami. Jakość badania USG zależy od sprzętu, ustawienia pacjenta, doświadczenia operatora. Algorytm, który zna tylko jedną wersję rzeczywistości, może „zgubić się” w innej – nawet jeśli klinicznie przedstawia ten sam problem. Kluczowa bariera? Model nie rozpoznaje braków swojej wiedzy. Nie widzi luk, nie zna kontekstu – ale mimo to generuje wynik z pełnym przekonaniem. Dlatego kontekst – czyli to, co dzieje się pomiędzy danymi – nadal wymaga ludzkiego osądu.

SI jako wsparcie edukacji – potencjał i ryzyko

Wirtualna sala sekcyjna

Choć zastosowania SI w medycynie weterynaryjnej koncentrują się głównie wokół diagnostyki i dokumentacji, coraz więcej uwagi poświęca się także edukacji – zwłaszcza w obszarze anatomii.

Rzeczywistość wirtualna (ang. virtual reality, VR) i rozszerzona (ang. augmented reality, AR) oraz cyfrowa anatomia to narzędzia, które rewolucjonizują sposób, w jaki studenci uczą się złożonych struktur anatomicznych (11). Umożliwiają one wielokrotne, nieliniowe eksplorowanie struktur anatomicznych, bez zapachu formaliny i potrzeby trzymania skalpela. Warto jednak pamiętać, że technologia może zachwycać wizualnie, ale nie wykształci ręki chirurgicznej ani nie nauczy wycucia tkanek.

SI jako tutor i partner w edukacji klinicznej

Rola sztucznej inteligencji w edukacji weterynaryjnej wykracza dziś poza wizualizacje anatomiczne. Duże modele językowe (ang. large language models, LLMs), wspierają proces kształcenia poprzez generowanie pytań egzaminacyjnych, symulowanie przypadków klinicznych i tłumaczenie trudnych pojęć. Z jednej strony pozwala to na efektywne ćwiczenie rzadkich scenariuszy klinicznych i oszczędność czasu wykładowców, z drugiej – niesie ryzyko spłycenia refleksji, jeśli interakcja z algorytmem nie jest odpowiednio nadzorowana. Jak ostrzegają badacze (14), zbyt duże poleganie na gotowych odpowiedziach może osłabiać krytyczne myślenie studentów. Kluczowe jest zatem projektowanie interakcji człowiek – SI tak, by technologia wspierała rozwój kompetencji, zamiast je zastępować.

Granice i ryzyko „spłycenia” myślenia

Nadmierne zaufanie do algorytmów może paradoksalnie osłabić umiejętność samodzielnego analizowania przypadków. Jak podkreśla Sobkowich (19), nadmierne poleganie na gotowych odpowiedziach generowanych przez algorytmy może prowadzić nie tylko do błędów poznawczych, ale też do erozji samodzielnego myślenia klinicznego. Mniej doświadczeni użytkownicy mogą traktować odpowiedzi modelu jako „niepodważalny fakt”. Gibson i wsp. (14) alarmują, że – szczególnie w początkowej fazie nauki – studenci powinni uczyć się analizować dane, a nie tylko przyjmować wnioski. W przeciwnym razie istnieje ryzyko spłycenia procesu myślowego, redukującego złożoną diagnozę do kilku sugestii generowanych przez algorytm.

Dryf modelu: kiedy algorytm przestaje być aktualny

Jedną z mniej oczywistych, a jednocześnie bardzo istotnych cech systemów SI jest ich podatność na utratę trafności

diagnostycznej w czasie – zjawisko znane jako dryf modelu (ang. model drift). Oznacza to, że nawet dobrze wytrenowany model, który wczoraj osiągał 90 % trafności w rozpoznawaniu konkretnego schorzenia, dziś może działać gorzej – bo zmieniły się wytyczne albo zmodyfikowano sposób opisywania objawów. Dlatego wdrożenie SI w klinice nie może być jednorazowym aktem instalacji oprogramowania, lecz musi być procesem ciągłego monitorowania, aktualizacji i walidacji (2).

Model drift można podzielić na:

- Dryf danych (ang. data drift), gdy zmieniają się dane wejściowe (np. inne rasy, nowe formaty zdjęć, nowe urządzenia),
- Dryf koncepcji (ang. concept drift), gdy zmienia się związek między objawem a diagnozą (np. zmiana wytycznych klinicznych),
- Starzenie się modelu (ang. model aging), gdy dochodzi do „utraty świeżości” systemu.

Bez systematycznego aktualizowania modelu jego trafność może stopniowo spadać – aż do momentu, w którym przestaje być adekwatny do zmieniających się warunków klinicznych i zaczyna generować błędy, których wcześniej nie obserwowano.

Wszystko to pokazuje, że nawet najlepszy algorytm może się mylić, jeśli zostanie pozbawiony kontekstu, wśród danych, których nie zna. Dlatego kluczowe jest, by nie traktować SI jako uniwersalnego narzędzia, ale jako system wrażliwy na otoczenie, który działa najlepiej wtedy, gdy informacje są zgodne z jego wcześniejszym doświadczeniem – a to doświadczenie trzeba zbudować świadomie i różnorodnie.

Przykłady błędów w praktyce weterynaryjnej

Oponiaki i błędna klasyfikacja przez Inception-V3

W badaniu dotyczącym różnicowania histopatologicznego stopnia złośliwości oponiaków za pomocą głębokiego uczenia maszynowego, model Inception-V3 mylnie zaklasyfikował 1 na 38 guzów WHO Grade II-III oraz 7 na 79 guzów WHO Grade I na podstawie map współczynnika dyfuzji pozornej (ADC). Pokazuje to, że nawet obiecujące modele SI mogą generować wyniki fałszywie negatywne (przeoczenie poważniejszego przypadku) oraz fałszywie dodatnie (zaklasyfikowanie łagodnego przypadku jako złośliwego) (6).

Różnice w diagnozie: lekarz kontra SI

W badaniu oceniającym zgodność diagnoz SI i radiologa w kierunku kardio-genego obrzęku płuc (CPE) u psów (n=481), uzyskano 444 zgodne rozpoznania (CPE+ lub CPE-), ale aż 37 przypadków wykazało rozbieżności. W 33 sytuacjach SI „widziała” obrzęk, którego nie potwierdził radiolog, a w 4 – nie rozpoznała go tam, gdzie lekarz go stwierdził (18). Pokazuje to, że mimo wysokiej skuteczności, algorytmy diagnostyczne mogą popełniać błędy wymagające ludzkiej korekty – zwłaszcza, gdy konsekwencje kliniczne są istotne.

Pałapki diagnostyki opartej o opisy tekstowe

W eksperymencie diagnostycznym ChatGPT nieprawidłowo ocenił wszystkie pięć przypadków paroksyzmalnej dyskinety u psów, bazując jedynie na historii choroby. Nawet po dołączeniu wyników badań klinicznych, model nie postawił poprawnego rozpoznania (1). Dopiero po zmianie języka dokumentacji i ponownym przedstawieniu przypadków, jedno rozpoznanie zostało trafnie wskazane. To pokazuje, że generatywne modele językowe – choć obiecujące – mogą być podatne na błędy interpretacyjne, szczególnie w przypadku schorzeń rzadkich lub trudnych do opisu tekstowo.

Potencjał w analizie radiologicznej

Warto zauważyć, że skuteczność systemów SI bywa nie tylko porównywalna z wynikami lekarzy, ale w niektórych przypadkach – istotnie wyższa. W badaniu Boissady i wsp. (18), obejmującym ponad 22 000 zdjęć RTG klatki piersiowej psów i kotów, algorytm głębokiego uczenia wykazał niższy współczynnik błędów niż zarówno lekarze weterynarii, jak i lekarze wspomagani przez ten sam system (odpowiednio: 10,7 % vs. 16,8 % i 17,2 %). SI okazała się szczególnie skuteczna w rozpoznawaniu powiększenia sylwetki serca oraz wzorca oskrzelowego. Autorzy podkreślają jednak, że system nie stawiał diagnoz – identyfikował jedynie zmiany obrazowe, które wymagały dalszej interpretacji klinicznej przez człowieka.

Błędy interpretacyjne w testach diagnostycznych

System Vetscan Imagyst, oceniany pod kątem wykrywania cyst Giardia duodenalis u psów, mimo wysokiej czułości i swoistości (88,4 % i 98,1 %), wykazał istotne ograniczenia. W 1,4 % przypadków

generował fałszywie dodatnie wyniki, a liczba wykrytych cyst znacząco się różniła między powtórzeniami tego samego testu (CV: 67 %). Autorzy uznali, że narzędzie może być użyteczne jakościowo, ale jego zastosowanie ilościowe wymaga dużej ostrożności i nadzoru (17).

Ograniczenia technologiczne i systemowe: gdy infrastruktura nie nadąża za ambicją

Choć możliwości sztucznej inteligencji wydają się coraz szersze, wiele barier nadal wynika nie z algorytmów jako takich, lecz z otoczenia, w którym te algorytmy działają. Technologie SI potrzebują określonych warunków – danych, sprzętu, standaryzacji i ludzi rozumiejących ich naturę. Ich brak może szybko ujawnić ograniczenia całego systemu.

Problemy integracji i kompatybilności sprzętowej

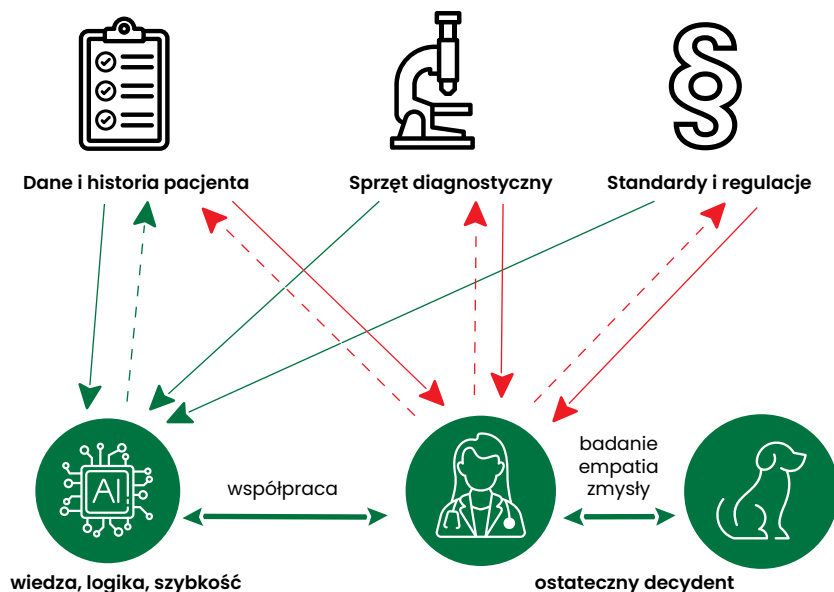
Dodatkowo, wiele systemów SI nie jest przystosowanych do integracji z funkcjonującymi już w klinikach interoperacyjnymi systemami PIMS (Practice Information Management System – system zarządzania lecznicą) i PACS (Picture Archiving and Communication System – system archiwizacji obrazów medycznych). Brak kompatybilności między oprogramowaniem skutkuje koniecznością kopiowania wyników lub pracy w równoległych systemach, co nie tylko spowalnia diagnostykę, ale także zwiększa ryzyko błędów.

Zaplecze techniczne: fundament skutecznego wdrożenia

Systemy SI, aby funkcjonować poprawnie, potrzebują odpowiedniego zaplecza: serwerów, pamięci, szybkich łącz, zabezpieczeń, a także ludzi, którzy potrafią je wdrożyć i obsłużyć. W wielu lecznicach brakuje nie tylko takich zasobów, ale również odpowiednio przeszkolonego personelu technicznego. To ograniczenie ma bardzo praktyczny wymiar: nawet najlepszy model może być bezużyteczny, jeśli nie zostanie poprawnie wdrożony, zintegrowany lub zaktualizowany.

Z tych powodów wiele obecnych wdrożeń SI w medycynie weterynaryjnej ogranicza się do bardzo wąskich zastosowań – np. wspomaganie oceny konkretnych struktur w obrazach RTG – i rzadko ma charakter systemowy. Dla rozwoju SI w medycynie weterynaryjnej niezbędne jest nie tylko projektowanie lepszych algorytmów, ale również budowanie środowiska, które pozwoli im

Ryc. 3. Ekosystem współpracy człowiek – SI w weterynarii.



Ekosystem współpracy między lekarzem weterynarii a sztuczną inteligencją, gdzie lekarz pozostaje centralnym decydem.
Przedstawiono wzajemne powiązania między lekarzem, systemem AI, danymi pacjenta, infrastrukturą oraz regulacjami prawnymi i etycznymi.

działać: ustandaryzowanej dokumentacji, interoperacyjnych systemów PIMS, otwartych baz danych i odpowiedniego wsparcia technicznego.

Brak ram prawnych – czyli kto pilnuje SI w medycynie weterynaryjnej?

W odróżnieniu od systemów stosowanych w medycynie ludzkiej, narzędzia sztucznej inteligencji wykorzystywane w medycynie weterynaryjnej nie podlegają obecnie formalnym procedurom dopuszczenia do użytku ani nadzorowi instytucji regulacyjnych. Jak podkreślają Cohen i Gordon (13), brak odpowiednich regulacji i standardów w zakresie walidacji narzędzi SI może prowadzić do poważnych konsekwencji klinicznych.

Również Bellamy (7) zwraca uwagę, że w przeciwieństwie do urządzeń medycznych stosowanych u ludzi, które muszą uzyskać aprobatę odpowiednich organów (np. FDA czy Health Canada), w medycynie weterynaryjnej nie istnieją analogiczne procedury. W konsekwencji, algorytmy wspomagające decyzje diagnostyczne i terapeutyczne mogą być wdrażane bez niezależnej walidacji, bez nadzoru oraz bez obowiązku ujawniania źródeł danych i metodyki ich opracowania.

Bellamy (7) podkreśla, że taka sytuacja niesie istotne ryzyko zarówno dla leka-

rzy, jak i pacjentów. Autor postuluje stworzenie procedur zatwierdzania modeli SI dedykowanych medycynie weterynaryjnej, uwzględniających m.in. opis danych treningowych i testowych, parametry skuteczności (czułość, swoistość, PPV, NPV) oraz obowiązkowy nadzór postwdrożeniowy. Co więcej, przypomina, że systemy o charakterze ciągłego uczenia się (continuous learning) powinny być wykorzystywane wyłącznie w warunkach badawczych. W codziennej praktyce klinicznej dopuszczalne powinny być jedynie modele „zamknięte”, stabilne i powtarzalne (7).

Dodatkowe zagrożenie wiąże się z potencjalnym naruszeniem prywatności danych. Jak zauważają Chu i wsp. (12), brak wbudowanych mechanizmów ochronnych może skutkować przypadkowym przesyłaniem do modeli SI wrażliwych danych pacjentów lub ich opiekunów np. poprzez nieświadome umieszczenie ich w treści zapytania. Tego typu luki w zabezpieczeniach stanowią poważne wyzwanie dla odpowiedzialnego stosowania SI w praktyce.

Brak formalnych ram prawnych dla sztucznej inteligencji w medycynie weterynaryjnej nie oznacza, że temat może zostać pominięty. Wręcz przeciwnie – to dobry moment, by środowisko zawodo-

we rozpoczęło dyskusję o standardach wdrażania i użytkowania SI. Refleksja nad tym, jak weryfikować skuteczność narzędzi, jak szkolić użytkowników i jak dzielić odpowiedzialność, może pozwolić uniknąć chaosu regulacyjnego w przyszłości. Przemysłane podejście do nowych technologii to nie tylko kwestia bezpieczeństwa pacjentów, ale też szansa na ugruntowanie roli lekarza weterynarii jako świadomego lidera w świecie nowoczesnej medycyny.

Podsumowanie

W świetle przedstawionych analiz jasne staje się, że sztuczna inteligencja w medycynie weterynaryjnej niesie ze sobą ogromny potencjał, ale też istotne ryzyka. Zrozumienie źródeł błędów technologicznych, świadomość ograniczeń algorytmów oraz kompetentne wykorzystanie SI przez lekarzy to warunek bezpiecznego i skutecznego wdrażania tych narzędzi w praktyce klinicznej. Współpraca człowieka i maszyny może być niezwykle efektywna (ryc. 3), ale tylko wtedy, gdy fundamentem relacji będzie docieklivość, świadomość ograniczeń i gotowość do powiedzenia: „Sprawdzam”.

Ten artykuł jest próbą pokazania, że choć SI potrafi wiele, to nie jest narzędziem usprawniającym wszystko i w każdych okolicznościach. Modele mogą się mylić. Mogą opierać się na niepełnych lub błędnych danych, ignorować kontekst, halucynować wyniki (1). Sztuczna inteligencja może wspierać decyzje, ale ich ciężar zawsze spoczywa na człowieku. Dlatego tak ważne są krytyczne myślenie, odpowiednie szkolenia i budowanie transparentnych rozwiązań. Tylko wtedy technologia staje się realnym wsparciem – a nie źródłem błędów i frustracji. Technologia przyspiesza decyzje. Lekarz jednak nadaje kierunek.

Profesjonalny sceptycyzm, świadomość ograniczeń oraz gotowość do kwestionowania wyników stanowią dziś kluczowe filary bezpiecznego stosowania sztucznej inteligencji w praktyce klinicznej. Algorytmy mogą wspierać diagnostykę, przyspieszać dokumentację, porządkować dane, ale nie przejmą odpowiedzialności. Nie zastąpią ludzkiego osądu, doświadczenia ani relacji z pacjentem i jego opiekunem. Dlatego decyzja nie może należeć do algorytmu – musi należeć do lekarza. Do tego, który rzeczywiście widzi pacjenta. Choć pozornie ten tekst dotyczył błędów, w istocie mówił o czymś głębszym: o zaufaniu. Do siebie. Do wiedzy. Do doświadczenia. Do empatii.



ADOBE STOCK

Także do narzędzi, których działania nie rozumiemy w pełni, ale które decydujemy się stosować z pełną świadomością ryzyka i odpowiedzialności. Bo technologia nie musi być doskonała. Wystarczy, że wiemy, kiedy jej zaufać – a kiedy zapytać jeszcze raz. ●

Piśmiennictwo

- Abani S., De Decker S., Tipold A., Nessler J. N., Volk H. A.: Can ChatGPT diagnose my collapsing dog? „Front. Vet. Sci.”, 2023, 10, 1245168.
- Albadrani B. A., Abdel-Raheem M. A., Al-Farwachi M. I.: Artificial Intelligence I Veterinary Care: A Review of Applications for Animal Health. „Egypt. J. Vet. Sci.”, 2024, 55 (6), 1725–1736, doi: 10.21608/EJVS.2024.260989.1769.
- Akinsulie O. C., Idris I., Aliyu V. A., Shahzad S., Banwo O. G., Ogunleye S. C., Olorunshola M., Okedoyin D. O., Ugwu C., Oladapo I. P., Gbadegoye J. O., Akande Q. A., Babawale P., Rostami S., Soetan K. O.: The potential application of artificial intelligence in veterinary clinical practice and biomedical research. „Front. Vet. Sci.”, 2024, 11, 1347550, doi: 10.3389/fvets.2024.1347550.
- Appleby R. B., Basran P. S.: Artificial intelligence in veterinary medicine: a bibliometric study. „JAVMA”, 2022, 260 (8), 819, doi: 10.2460/javma.22.03.0093.
- Aubreville M., Bertram C. A., Marzahl C., Gurtner C., Dettwiler M., Schmidt A., et al.: Deep learning algorithms out-perform veterinary pathologists in detecting the mitotically most active tumor region. „Sci. Rep.”, 2020, 10, 16447, doi: 10.1038/s41598-020-73246-2.
- Banzato T., Causin F., Della Puppa A., Cester G., Mazzi L., Zotti A.: Accuracy of Deep Learning to Differentiate the Histopathological Grading of Meningiomas on MR Images: A Preliminary Study. „J. Magn. Reson. Imaging”, 2019, n. d., n. d, doi: 10.1002/jmri.26723.
- Bellamy J. E. C.: Artificial intelligence in veterinary medicine requires regulation. „Can. Vet. J.”, 2023, 64, 968–970.
- Bouchemla F., Akchurin S. V., Akchurina I. V., Dylugher G. P., Latynina E. S., Grecheneva A. V.: Artificial intelligence feasibility in veterinary medicine: A systematic review. „Veterinary World”, 2023, 16 (10), 2143–2149, doi: 10.14202/vetworld.2023.2143-2149.
- Burti S., Zotti A., Banzato T.: Role of AI in diagnostic imaging error reduction. „Front. Vet. Sci.”, 2024, 11, 1437284, doi: 10.3389/fvets.2024.1437284.
- Şakir S. S., Ötünçtemur A.: Artificial Intelligence in Medicine. „Eur. Arch. Med. Res.”, 2018, 34 (Suppl. 1), S1–S3, doi: 10.5152/eamr.2018.43534.
- Choudhary O. P., Infant S. S., As V., Vickram, Chopra H., Manuta N.: Exploring the potential and limitations of artificial intelligence in animal anatomy. „Ann. Anat.”, 2025, 258, 152366, doi: 10.1016/j.aanat.2024.152366.
- Chu C. P.: ChatGPT in veterinary medicine: a practical guidance of generative artificial intelligence in clinics, education, and research. „Front. Vet. Sci.”, 2024, 11, 1395934, doi: 10.3389/fvets.2024.1395934.
- Cohen A., Gordon K.: First, do no harm. Ethical and legal issues of artificial intelligence and machine learning in veterinary radiology and radiation oncology. „Vet. Radiol. Ultrasound”, 2022, 63 (Suppl. 1), 840–850, doi: 10.1111/vru.13171.
- Gibson R., Ury D.: Artificial intelligence: reflection on risks and benefits in veterinary education and practice. „Transformative Dialogues: Teach. Learn. J.”, 2025, 18 (1), n. d., doi: 10.26209/td2025vol18iss11866.
- Jarynowski A.: Możliwości wykorzystywania sztucznej inteligencji (AI) w kontroli stanu zdrowotnego bydła i świń w Polsce oraz regionie. Preprint, October 2024, doi: 10.13140/RG.2.2.21595.58407.
- Kennedy N., Brodbelt D. C., Church D. B., O'Neill D. G.: Detecting false-positive disease references in veterinary clinical notes without manual annotations. „npj Digit. Med.”, 2019, 2, 33, doi: 10.1038/s41746-019-0108-y.
- Kanski S., Busch K., Hailmann R., Weber K.: Performance of the Vetscan Imagoyst in point-of-care detection of Giardia duodenalis in canine fecal samples. „J. Vet. Diagn. Invest.”, 2024, n. d., 1–8, doi: 10.1177/10406387241279177.
- Kim E., Fischetti E., Weltman J. G., Sreetharan P., Fox P. R.: Comparison of artificial intelligence to the veterinary radiologist's diagnosis of canine cardiogenic pulmonary edema. „Vet. Radiol. Ultrasound”, 2022, n. d., 1–7, doi: 10.1111/vru.13062.
- Sobkowich K. E.: Demystifying artificial intelligence for veterinary professionals: practical applications and future potential. „Am. J. Vet. Res.”, 2025, 86 (S1), S6–S12, doi: 10.2460/ajvr.24.09.0275.

Monika Szymańska,
e-mail: m.szymanska@libraxis.ai