

Sztuczna inteligencja może bardzo się mylić

18 sierpnia 2020

Autorzy badań opublikowanych na łamach „PNAS” ostrzegają, że nie można ufać technikom obrazowania medycznego rekonstruowanym za pomocą sztucznej inteligencji. Międzynarodowy zespół naukowy pracujący pod kierunkiem Andersa Hansena z Uniwersytetu w Cambridge stwierdził, że narzędzia do głębokiego uczenia się, które rekonstruuja obrazy wysokiej jakości na podstawie szybkich skanów, tworzą liczne przekłamania i artefakty, które mogą wpływać na diagnozę.



Jak niejednokrotnie informowaliśmy, systemy sztucznej inteligencji są już na tyle zaawansowane, że równie dobrze jak radiolodzy, a [często i lepiej](#), potrafią opisywać zdjęcia RTG, obrazy tomografii komputerowej czy rezonansu magnetycznego. W związku z tym pojawił się pomysł, by SI zaprząć do rekonstrukcji obrazów.

Pomysł polega na tym, by wykonywać obrazowanie o niższej rozdzielczości, czyli pobierać dane z mniejszej liczby punktów, a następnie, by wytrenowane systemy algorytmy

sztucznej inteligencji rekonstruowały na tej postawie obraz o wysokiej rozdzielczości. W ten sposób można by zaoszczędzić czas i pieniądze potrzebny na wykonanie badania. Wykorzystywane tutaj algorytmy były trenowane na dużej bazie danych obrazów wysokiej jakości, co stanowi znaczne odejście od klasycznych technik rekonstrukcji bazujących na teoriach matematycznych.

Okazuje się jednak, że takie systemy SI mają poważne problemy. Mogą one bowiem przegapić niewielkie zmiany strukturalne, takie jak małe guzy nowotworowe, podczas gdy niewielkie, niemal niewidoczne zakłócenia spowodowane np. poruszeniem się pacjenta, mogą zostać odtworzone jako poważne artefakty na obrazie wyjściowym.

Zespół w skład którego weszli Vegard Antun z Uniwersytetu w Oslo, Francesco Renna z Uniwersytetu w Porto, Clarice Poon z Uniwersytetu w Bath, Ben Adcock z Simon Fraser University oraz wspomniany już Anders Hansen, przetestował sześć sieci neuronowych, wykorzystywanych do rekonstrukcji obrazów tomografii i rezonansu. Sieciom zaprezentowano dane odpowiadają trzem potencjalnym problemom, które mogą się pojawić: niewielkim zakłóceniom, niewielkim zmianom strukturalnym oraz zmianom w próbkowaniu w porównaniu z danymi, na których system był trenowany.

„Wykazaliśmy, że niewielkie zakłócenia, których nie widać gołym okiem, mogą nagle stać się poważnym artefaktem, który pojawia się na obrazie, albo coś zostaje przez nie usunięte. Dostajemy więc fałszywie pozytywne i fałszywie negatywne dane” – wyjaśnia Hansen.

Uczni, chcą sprawdzić zdolność systemu do wykrycia niewielkich zmian, dodali do skanów niewielkie litery i symbole z kart do gry. Tylko jedna z sieci była w stanie je prawidłowo zrekonstruować. Pozostałe sieci albo pokazały w tym miejscu niewyraźny obraz, albo usunęły te dodatki.

Okazało się też, że tylko jedna sieć neuronowa radziła sobie ze zwiększaniem tempa skanowania i tworzyła lepszej jakości obrazy niż wynikałoby to z otrzymanych przez nią danych wejściowych. Druga z sieci nie była w stanie poprawić jakości obrazów i pokazywała skany niskiej jakości, a trzy inne rekonstruowały obrazy w gorszej jakości niż otrzymały do obróbki. Ostatni z systemów nie pozwalał na zwiększenie szybkości skanowania.

Hansen stwierdza też, że badacze muszą zacząć testować stabilność takich systemów. „Wówczas przekonają się, że wiele takich systemów jest niestabilnych”. Jednak największym problemem jest fakt, że nie potrafimy w sposób matematyczny zrozumieć, jak działają tego typu systemy. „Są one dla nas tajemnicą. Jeśli ich porządnie nie przetestujemy, możemy otrzymać katastrofalnie złe wyniki”.

Na szczęście takie systemy nie są jeszcze wykorzystywane w praktyce klinicznej. Zespół Hansena stworzył odpowiednie testy do ich sprawdzenia. Uczni mówią, że nie chcą, by takie systemy zostały dopuszczone do użycia jeśli nie przejdą szczegółowych testów.

Autorstwo: Mariusz Błoński

Zdjęcie: [Parentingupstream](#) (CC0)

Na podstawie: [Physicsworld.com](#), [PNAS.org](#)

Źródło: [KopalniaWiedzy.pl](#)