

Reinforcement Learning

drogą do sztucznej inteligencji

Data science to dziedzina nauki, której celem jest uzyskanie jak najwięcej korzyści z danych. Modele potrafią przewidzieć suszę i załamanie rynku, są w stanie przeanalizować strukturę sprzedaży i odnaleźć efektywne kanały logistyczne. Za pomocą nowoczesnych technik mogą być identyfikowane zagrożenia i demaskowane oszustwa. Warto jednak zadać pytanie co tak naprawdę zmieniło się wraz z pojawieniem się data science? Modele matematyczne służące do prognozowania i optymalizacji istniały przecież wcześniej. Wcześniej danych było jednak mniej i trudniej je było pozyskać. Komputery były wolniejsze i trudniejsze w obsłudze. Niewątpliwie rozwój informatyki spowodował wzrost ilości danych. Tanie komputery oraz nowe oprogramowanie przyczyniły się do spopularyzowania modelowania i przetwarzania danych. Mimo wszystko, to co obecnie obserwujemy było już wcześniej tylko w mniejszej skali i bardziej prymitywnej formie. To tak jak byśmy mówili, że 50 lat temu było mniej samochodów i mniej dróg, a teraz to się zmieniło, mamy nowe samochody, dużo dróg i normy bezpieczeństwa. Jednak daleko nam do latania własnymi pojazdami powietrznymi.

Gdybyśmy jednak zatrzymali się przy tym porównaniu, moglibyśmy w kontekście data science powiedzieć, że kiedyś było mało samochodów, a dzisiaj zamiast jeździć samochodami wszyscy latamy pojazdami pilotowanymi przez autonomicznych komputerowych pilotów.

A więc jednak w data science pojawił się postęp nie tylko w obszarze ilości i powszechności, czyli prostej ewolucyjnej modernizacji. Stańliśmy jako ludzkość na progu bardzo stromych schodów biegnących do góry. Nie wiemy gdzie nas one zaprowadzą, niedługo nic już nie będzie takie jak dotychczas.

Człowiek kontra maszyna

Pierwszym modelem, który wygrał partię szachów ze światowym mistrzem Garri Kasparowem był algorytm zainstalowany w superkomputerze IBM Deep Blue. Stało się to 10 lutego 1996 roku. Partia ta stała się znana jako zwycięstwo komputera nad człowiekiem, mimo że Kasparow wygrał trzy kolejne partie i dwie zremisował, ostatecznie wygrywając z Deep Blue wynikiem 4:2¹. Liczba kombinacji szachowych jest nazywana liczbą shannona i wynosi 10^{120} . Claude Shannon wyprowadził tę liczbę by udowodnić, że komputer nie jest w stanie przeanalizować każdego możliwego ruchu, ponieważ przeliczenie jednego wariantu ruchu na mikrosekundę „zająłoby mu czas do końca wszechświata by dokończyć obliczenia”². Przy budowie systemu rozproszonego Deep Blue zastosowano metodę prawdopodobieństwa zdarzeń opisaną w łańcuchach Markowa. Na podstawie obecnej sytuacji komputer eliminował wszystkie ustawienia niemożliwe i pobierał historyczne ustawienia szachowe z przeszłości jako koncepcje, które następnie wybierał. W rozgrywce szachowej pierwszy ruch (tzw. otwarcie) ograniczony jest do kilkunastu możliwości, kolejny ruch daje kolejne ograniczone możliwości. Model oparto więc na drzewach decyzyjnych bazujących na historycznych rozgrywkach. Dzięki temu zyskano istotne ograniczenie wariantów, dlatego system grał w szachy stosunkowo płynnie.

Gra, która dotąd wydawała się niemożliwa do zaprogramowania to chińskie Go³. W przypadku Go jest odwrotnie

niż w szachach – pierwszy ruch to kilka-set możliwości ($19 \times 19 = 361$). W trakcie gry liczba możliwości spada aż do momentu pierwszego zbitia – wtedy ponownie rośnie. Jest więc w jakimś stopniu niemożliwe korzystanie z wcześniejszych doświadczeń, a liczba wariantów posunąć rośnie wykładniczo.

Istnieje znana maksyma: „Jeżeli ktoś uważa, że szachy to król gier, to Go jest jego cesarzem. Do niedawna uważano, że niemożliwe jest stworzenie algorytmu Go, który byłby w stanie pokonywać arcymistrzów tej gry.

W 2014 roku Mark Zuckerberg, twórca Facebooku, wspomniął, że zatrudnił specjalny zespół naukowców zajmujących się trenowaniem AI do gry w Go. Podobne sygnały wysłał zarząd Google. Zaczęła się rywalizacja. W obu firmach zastosowano nowatorską metodę nauki maszyn o nazwie „uczenie przez wzmacnianie” (*reinforced learning*). Głębokie sieci neuronowe zaczęły miesiącami grać ze sobą w Go. Każda trwająca kilka minut rozgrywka powoli poprawiała zdolności maszyn.

Firma Google zaprosiła do gry z systemem AlphaGo trzykrotnego Mistrza Europy (2013, 2014, 2015), Fana Hui, który doskonalił się w Go od 12 roku życia. W październiku 2015 roku przegrał on z AlphaGo wynikiem 0:5. Fan Hui był pod ogromnym wrażeniem. Stwierdził: *bardzo mocy i stabilny program. Gdyby nikt mi nie powiedział, że AlphaGo jest komputerem, myślałbym, że to nieco dziwne, ale bardzo mocny gracz*⁴.

Modele i ich prostota

Trudno nazwać model matematyczny sztuczną inteligencją. Najprostszym modelem jest średnia ważona krocząca. Wynikiem takiego modelu jest jakieś

¹ Wikipedia Wolna Encyklopedia, https://pl.wikipedia.org/wiki/Deep_Blue

² Szachy (home.agh.edu.pl) https://home.agh.edu.pl/~zobmat/2019/2_nagroda_2/math.html

³ „Padł kolejny bastion, w którym królował człowiek – komputer wygrał w go”; Hubert Taler 29 stycznia 2016; <https://spidersweb.pl/2016/01/komputer-wygral-w-go.html>

⁴ „Komputer pokonał mistrza Go (a to nie koniec)”, Wojciech Kulik, www.benchmark.pl <https://www.benchmark.pl/aktualnosci/komputer-pokonala-mistrza-go-a-to-nie-koniec.html>

uogólnienie, wartość średnia z pewnego okresu. Używając takiego modelu możemy dowiedzieć się jaka będzie prawdopodobna wartość (średni poziom cen, opadów, narodzin) w określonym punkcie w przyszłości. Bazując na rozkładach gęstości prawdopodobieństwa szeregów czasowych możemy przewidywać rozmaite zjawiska, ich skalę i trend.

Z czasem pojawiły się pierwsze modele regresji, które były macierzowym rozwinięciem prostych modeli bazujących na średnich. Potem pojawiły się modele oparte na zjawisku entropii mające zdolność wykrywania i wykorzystywania zjawiska porządkowania chaosu w układzie zamkniętym.

Prawdziwym przełomem było wynalezienie sieci neuronowych. To modele naśladujące procesy myślowe zachodzące w mózgach istot żyjących. Bardzo długo trwała stagnacja w rozwoju tych modeli. Po wynalezieniu sieci jedno- i dwuwarstwowej badacze nie potrafili zrobić kolejnego kroku. Pierwsze sieci neuronowe działały jak filtr w odkurzaczu. Przez ich neurony przelatywały informacje, zniekształcane przez wagi. Poprzez zwiększanie i zmniejszanie znaczenia poszczególnych informacji sieci neuronowe dostarczały właściwych odpowiedzi, poprawnych wyników. Nie różniły się one jednak wiele od najprostszych modeli regresji liniowych. Tak jak powietrze wdmuchiwanie do organów kościelnych przekształcało się w piękny dźwięk. Nie możemy jednak powiedzieć, że organy kościelne są bardzo skomplikowanym i wyrafinowanym instrumentem. Mechanizm sprężyn, rur i trybików działał tak jak przed wiekami lecz teraz to wszystko przybrało formę cyfrową.

Zjawisko rekurencji

Ktoś wreszcie wpadł na pomysł, aby wewnątrz sieci neuronowej zastosować ujemne sprzężenie zwrotne, zwane rekurencją (z łac. *recurrere* – przybiec z powrotem). Informacja wdmuchiwana jak powietrze do organów zamiast wylatywać po drugiej stronie piszczałki instrumentu zaczynała krążyć w skończonych pętlach. W modelu pojawiło się urządzenie, które zaczęło zmieniać charakterystykę modelu od środka. To tak jak byśmy w organach mieli urządzenie

mechaniczne, które zmieniało od środka budowę piszczałek, aby po dłuższej nauce wydawały one dźwięk oczekiwany przez kompozytora. Wiem, że brzmi to niedorzecznie, a to był dopiero początek „znaków”.

Jak więc działa rekurencyjna sieć neuronowa? Składa się ona z wielu warstw neuronów, z pośród których warstwy wewnątrz mają wbudowany system autokorekty błędów. Informacja przetwarzana jest wielokrotnie aż do uzyskania zadowalając wyników na wyjściu. Zjawisko to przypomina naukę, jest jednak czymś więcej. Jest to proces, który znamy z ewolucji. To rodzaj nauki międzypokoleniowej. Kolejne pokolenia uczą się na bazie wiedzy i doświadczenia swoich poprzedników.

Convolutional neural network

Wydawałoby się, że coś co może w nieskończoność uczyć się na swoich błędach, dojdzie do niebotycznych możliwości poznawczych. Takie zdolności w istocie zostały osiągnięte i kuriozalnie okazały się największym przełomem rekurencyjnych sieci neuronowych. Okazało się mianowicie, że modele te ucząc się w skończonych pętlach z danych treningowych dochodziły do wręcz niewyobrażalnego poziomu dokładności naśladowania procesów. Niestety, gdy dawano modelom nowe, nieznane wcześniej dane, stawały się one bezradne nie radząc sobie w ogóle z ich interpretacją.

Okazało się, że nauka sieci rekurencyjnej bez odpowiedniej konfiguracji prowadzi do braku generalizowania zjawiska. Model zamiast szukać wzorców, współzależności i powiązań uczył się dużej części danych na pamięć. Tą częścią danych był szum, czyli dane przypadkowe, niemające związku z procesem.

Opisane tu zjawisko nosi nazwę przeuczenia i stanowi podstawowy problem pracy z zaawansowanymi modelami *deep learning* i *machine learning*. Z biegiem czasu badacze nauczyli się radzić sobie z tym zjawiskiem, niestety budowa złożonych sieci neuronowych nie jest umiejętnością łatwą. Jest to bardzo żmudny proces, który trudno zautomatyzować. Wymaga on wysokich kwalifikacji i doświadczenia. Biznes nie może czekać na rzadkich ekspertów,

modelowanie potrzebne jest pilnie. Nowatorskie badania z obszaru *deep learning* pochłaniają dużo czasu i zasobów.

Ensemble Learning

Pomysł rekurencyjnych sieci neuronowych zapoczątkował nowy sposób myślenia o zastosowaniu modeli z obszaru *machine learning*. Obszar ten jest nieco prostszy i bardziej praktyczny niż *deep learning*. Do tej pory większość modeli działa na zasadzie algorytmu, który był uruchamiany i w sposób jednorazowy starał się nauczyć, wyszukać zależności, zrozumieć i ostatecznie stworzyć równanie opisujące badany proces.

Ważnym krokiem był pomysł nauki grupowej (*ensemble learning*). Postanowiono zastosować bardzo wiele identycznych modeli, które uczyłyby się jednocześnie tych samych danych, aby na końcu zebrać, podsumować i uzyskać średni wspólny wynik. Według podstawowej zasady nauczania grupowego, jeden, nawet najlepszy, model zawsze będzie gorszy niż tłum słabych modeli. Aby grupowy wynik kilkuset identycznych modeli był dobry modele te muszą się od siebie różnić. Ponieważ identyczne modele nie mogą się od siebie różnić zastosowano różne zestawy losowo dobranych danych dla tych modeli. Metoda ta nazywa się *bagging*. Losowo też zaczęto dobierać zmienne objaśniające (*feature randomness*). W ten sposób w ramach jednego modelu o nazwie *random forest* powstała populacja ogromnej ilości drzew decyzyjnych różniących się od siebie posiadanymi informacjami. Modele *random forest* zyskały ogromną popularność, często przewyższając zdolności prognostyczne innych modeli.

Boosting Machine Learning, czyli modele ze wzmocnieniem

Ludzkość osiągnęła swoją wiedzę i zaawansowanie technologiczne dzięki pokoleniom badaczy i naukowców. Dzięki umiejętności pisania kolejne generacje naukowców nie musiały wymyślać wszystkiego od nowa. Badacze mogli skorzystać z wiedzy, wyciągnąć wnioski z sukcesów i porażek poniesionych przez wcześniejsze pokolenia.

Metoda uczenia zespołowego zastosowana w przytoczonym wcześniej modelu *random forest* bardziej przypomina

proces wyborczy. Wielka liczba podobnych do siebie obywateli, z pośród których każdy ma własne losowo zebrane informacje, oddaje głos i w wyniku wspólnego głosowania powstaje wynik. Wygrywa to rozwiązanie, które zyskało większą liczbę głosów⁵.

Czy procesy wyborcze przyczyniały się do rozwoju naukowo-technicznego? Jak wspominałem, zostały wynalezione sieci neuronowe, które w nieskończonym procesie wewnętrznej nauki dochodziły do bardzo wysokich poziomów dokładności interpretacji procesów.

Ktoś postanowił ten sam mechanizm zastosować w modelu drzewa decyzyjnego. Zamiast grupować wielką ilość drzew decyzyjnych w jednym szeregu i wymagać od nich jednocześnie odpowiedź stworzono jedno drzewo decyzyjne, które po przeprowadzeniu procesu uczenia oceniało własną naukę, a następnie tworzyło kolejne drzewo, które uczyło się na błędach swojego poprzednika. Dokonania kolejnych drzew analizowane są przez funkcję straty (*loss function*), czyli praktycznie ten sam mechanizm, który stosują opisane wcześniej konwolucyjne sieci neuronowe. Powstał model, który znacznie przewyższył dotychczasowe rozwiązania przyczyniając się do kolejnego wielkiego kroku w rozwoju sztucznej inteligencji.

Czy już istnieje sztuczna inteligencja?

Obecnie sztuczna inteligencja rozwija się właśnie w kierunku konwolucyjnych i rekurencyjnych sieci neuronowych

⁵ Należy wskazać, że modele random forest nie są przeznaczone wyłącznie do zadań klasyfikacyjnych lecz również doskonale sprawdzają się w procesie regresji.

oraz w obszarze modeli opartych o naukę ze wzmocnieniem. Jest to intuicyjne rozwiązanie bazujące na ewolucyjnym międzypokoleniowym charakterze postępów nauki. Co będzie dalej? Niewątpliwie metody te są tylko krokiem milowym na drodze do powstania autonomicznych inteligentnych sztucznych bytów. Istnieje jeszcze wiele znanych obecnie lecz niewykorzystanych kierunków rozwoju. Są to między innymi metody łączące programowanie operacyjne oraz naukę ze wzmocnieniem w obrębie modularnych sieci neuronowe (*modular neural networks*). Rozwijana jest również technologia komputerów kwantowych, na bazie której mogą powstać modele kwantowe o niezwykłych możliwościach poznawczych.

Jeżeli powstanie prawdziwa replikowalna sztuczna inteligencja na przełomowe wynalazki nie będziemy czekali dekady lecz dni lub nawet godziny. Uważa się, że pojawienie się sztucznej superinteligencji doprowadzi do rozwoju technologicznego jakiego nie da się teraz nawet wyobrazić.

Jesteśmy obecnie świadkami procesu, gdzie konkurencyjne super algorytmy prowadzą ze sobą niekończące się gry, uczą się identyfikować obrazy, zachowania i znajdują nieprawdopodobnie unikalne zjawiska. Zdolności poznawcze ludzkiego mózgu pozostały daleko w tyle za wyspecjalizowanymi systemami analitycznymi. Systemy te są zbyt nieporadne, aby mogły zmieniać swoje specjalizacje w uniwersalne umiejętności. Innymi słowy daleko im do tego, aby stały się tak wszechstronne jak ludzie czy szerzej zwierzęta. Pewnie to tylko kwestia czasu, ich rozwój trwa.

Wojciech Moszczyński



Wojciech Moszczyński – absolwent Katedry Ekonometrii i Statystyki Uniwersytetu Mikołaja Kopernika w Toruniu, specjalista z zakresu ekonometrii, *data science* oraz rachunkowości zarządczej. Specjalizuje się w optymalizacji procesów produkcyjnych i logistycznych. Prowadzi badania w obszarze rozwoju i zastosowania sztucznej inteligencji. Od lat zaangażowany w popularyzację ekonometrii oraz *data science* w środowiskach biznesowych

W Pakistanie powstaje największy młyn pszeniczny

Położony w Azji Pakistan należy do krajów o najwyższym przyroście naturalnym (2% rocznie), zajmując 6. miejsce na świecie pod względem liczby ludności, a prognozy wskazują, że w 2055 r. może zająć 3 pozycję. Jednocześnie roczne spożycie pszenicy na jednego mieszkańca należy do najwyższych na świecie (124 kg), a mąka pszenna jest podstawowym produktem w diecie żywieniowej Pakistanczyków, dostarczając 72% kalorii. Ponadto Pakistan jest znaczącym eksporterem mąki pszennej, głównie do sąsiedniego Afganistanu, na poziomie 700 tys. ton rocznie.

Choć w Pakistanie czynnych jest ok. 1000 młynów o zdolności przemiałowej pszenicy 80 tys. ton/dobę to większość z nich jest technicznie przestarzała. Pomimo rekordowych zbiorów pszenicy w Pakistanie, szacowanych w sezonie 2021-22 na 27 mln ton, ilość ta staje się niewystarczająca dla pokrycia wewnętrznego popytu, eksportu i utrzymania rezerw strategicznych. W przeszłości Pakistan był eksporterem pszenicy netto, obecnie importuje ok. 2 mln ton pszenicy, głównie z Rosji i Ukrainy, a w 2020 r. rząd zniósł podatek importowy na pszenicę.

W tych okolicznościach firma Volka Food International podjęła decyzję o budowie nowoczesnego młyna pszenicznego o wydajności 800 t/dobę. Firma powstała w 2008 r. i jest wiodącym producentem wyrobów cukierniczych, piekarskich, makaronowych, a także słodczy jak lizaki, guma do żucia, cukierki, czekolada, które eksportuje na rynki Azji, Środkowego Wschodu i Europy. Zatrudnia ogółem 1500 osób.

Nowy młyn pozwoli firmie zapewnić zaopatrzenie w mąki wysokiej jakości na własne potrzeby, a także dla odbiorców zewnętrznych. Młyn ma zostać oddany do użytku latem 2022 r. razem z bazą magazynową obejmującą 32 silosy o łącznej pojemności 80 tys. ton. (*World Grain*, styczeń 2022 r.)

Opracował Krzysztof Zawadzki