

BIOINFOWORLD

Nr 4 Grudzień/Styczeń

Czasopismo Koła Naukowego Bioinformatyki BIT

ISSN 2449 -5913

Wydanie online

Opiekun naukowy

dr Jacek Śmietański

Redakcja

Paulina Kania

Rafał Szkotak

Anna Szylak

Okładka

Kamil Malisz

Korekta

Marta Szkop

Skład

Rafał Szkotak

Wydawca

Koło Naukowe Bioinformatyki BIT UJ

Kontakt

bioinfoworld.redakcja@gmail.com

Sfinansowane przez



RADA KÓŁ NAUKOWYCH
UNIwersytetu Jagiellońskiego

Spis treści

Spis treści.....	3
Wykorzystanie metod sztucznej inteligencji w bioinformatyce (dr hab. Igor Podolak).....	5
Duże dane – duże możliwości? (Michał Foryt)	9
Zagrożenia i nadzieje związane ze stworzeniem i rozwojem sztucznej inteligencji (Aleksander Baranowski)	13
Uczenie przez wzmocnienie – od plastyczności synaptycznej do uczenia maszynowego (Agata Szlaga)	17
Randomizowane sieci neuronowe (Agnieszka Słowik)	24
Sprawozdanie z działań Koła (Paulina Kania)	29
Słowniczek dla biologów	31
Słowniczek dla informatyków	32

Wykorzystanie metod sztucznej inteligencji w bioinformatyce

Dr hab. Igor Podolak

Wydział Matematyki i Informatyki
Instytut Informatyki i Matematyki Komputerowej

Wszędzie tam gdzie pojawia się dużo danych, wśród których należy znaleźć zależności, pojawia się uczenie maszynowe. Dzieje się tak szczególnie wtedy, gdy znamy albo niewiele reguł opisujących dane, albo nie znamy ich wcale. Jedynym co ma naukowiec jest duża (najlepiej) liczba danych z przypisanymi etykietami lub też bez nich.

Uczenie maszynowe jest uczeniem z danych. Wyróżniamy dwa podstawowe podejścia ze względu na postać danych. Jeśli dla każdego przykładu istnieje jakaś jego etykieta opisująca prawidłową wartość tego przykładu, to będziemy mówić o uczeniu nadzorowanym. Jeśli natomiast przykłady nie mają etykiet, to problemem będzie podział na grupy (skupiska, klastry) przykładów o podobnych cechach, które będzie można później rozpoznać (tu pewnie zajmie się tym ekspert).

Chciałbym omówić poniżej dwa ciekawe podejścia, które wciąż są obiecujące. W problemach bioinformatycznych mamy do czynienia z różnymi podejściami. Przykładem może być grupowanie genotypów, tak by rozpoznać geny odpowiedzialne za dane choroby i wyszukać ich współzależności. Przykładem może być próba znalezienia genów, które są odpowiedzialne na przykład za raka piersi. Możliwym rozwiązaniem jest wykorzystanie tzw. sieci Kohonena, które grupują razem przykłady o podobnych cechach jednocześnie układając utworzone grupy (klastry) w taki sposób, że dwa klastry bardziej do siebie podobne znajdą się w rozwiązaniu bliżej siebie niż dwa klastry mniej podobne.

Pozwala to na wyszukanie istotnych cech rozróżniających przykłady z obu grup. Jeśli na przykładach (tutaj genotypach) z poszczególnych grup wykonamy statystyczne testy, to będziemy często w stanie znaleźć pojedyncze cechy (geny) w rzeczywistości rozróżniające je od innych.

W przypadku sieci Kohonena jest to dodatkowo wygodne ponieważ model jest w rzeczywistości mapowaniem bardzo wysoko wymiarowej przestrzeni wejściowej do przestrzeni o ustalonej topologii, którą zwykle jest krata. Pozwala to więc dodatkowo na wizualizację danych. Model Kohonena radzi sobie z częstym problemem uczenia maszynowego zwanym przekleństwem wymiarowości. Ujawnia się on zwykle w postaci bardzo wysoko wymiarowej przestrzeni wejściowej przy niewielkiej liczbie rzeczywistych przykładów. W opisywanym tu zastosowaniu mamy genotypy składające się z setek tysięcy genów, z których tylko pojedyncze mają rzeczywisty wpływ na rozwiązanie (istnienie czy nie istnienie choroby). Proste analizy statystyczne, właśnie ze względu na małą liczbę przykładów, nie radzą sobie zwykle z redukcją niepotrzebnych wymiarów (tutaj genów). Model Kohonena bywa wykorzystywany w takich zastosowaniach.

Innym ciekawym zastosowaniem metod uczenia maszynowego jest użycie algorytmu wielorękich bandytów w problemie składania genotypu z krótkich readów bez informacji o całym genomie. Najbardziej znanym takim projektem był oczywiście Human Genome Project. W problemie rekonstrukcji (asemblacji) genotypu de novo z krótkich readów z maszyny sekwencjonującej nowoczesnego typu, mamy dostępne zwielokrotnione genotypy „pocięte” na odcinki o średniej długości ok. 100 nukleotydów. Dzięki temu, że każdy genotyp jest niezależnie podzielony, ready mają wspólne, chociaż bardzo krótkie, fragmenty, w których się nakładają. Zadaniem jest utworzenie szeregu kontigów, czyli złożonych z wielu readów spójnych fragmentów. Obecnie najczęściej wykorzystywanym algorytmem jest tzw. graf de Bruijna.

Wyobraźmy sobie jednak inne rozwiązanie tego problemu. Z pomocą przychodzi algorytm tzw. wielorękiego bandyty lub K-jednorękich bandytów. Użytkownik stoi przed K maszynami typu jednorękiego bandyty, gdzie należy do jednej z maszyn wrzucić monetę, pociągnąć za ramię i oczekiwać wygranej. Wybór maszyny jest dowolny, a celem jest, naturalnie, wygranie jak największej sumy pieniędzy. Każda maszyna ma oczekiwaną stopę zwrotu, a więc średnią wartość wygranej na każdą wrzuconą monetę, przy czym żadna z tych wartości nie jest znana. Gracz potrzebuje więc znaleźć poprzez doświadczenia tą maszynę, która ma stopę zwrotu najwyższą.

Gracz staje przed rozwiązaniem problemu eksploracji i eksploatacji (ang. exploration-exploitation dilemma). Z jednej strony powinien wybierać maszynę, która w toku dotychczasowych doświadczeń dawała najwyższą stopę zwrotu (to będzie eksploatacja), z drugiej strony te uzyskane wartości mają tylko pewną szansę bycia prawdziwymi. Jednak prawdopodobieństwo, że średnia z uzyskanych zwrotów jest podobna do rzeczywistej stopy zwrotu zwiększa się wraz z liczbą prób. Gracz powinien więc próbować również grać na maszynach, które do danej chwili nie miały wysokiej stopy zwrotu (to będzie eksploracja).

Ten problem jest rozwiązany teoretycznie w postaci algorytmu Upper Confidence Bounds (UCB), który wskazuje, za którą wajchę gracz powinien pociągnąć. Podobnie istnieje rozszerzenie algorytmu dla problemów związanych z grami o nazwie Upper Confidence Bounds for Trees (UCT). Jest to obecnie najlepszy algorytm dla tak złożonych gier jak Go, w którym stopień rozwidlenia jest niezwykle wysoki.

Jak zastosować UCT dla problemu asemblacji genomu? Można sobie wyobrazić, że krótkie ready są pionami w grze, które można wyciągnąć z puli dostępnych i dołączyć do już utworzonego kontigu, o ile mają z nimi krótki wspólny fragment. Oczywiście w danym momencie będzie więcej niż jedna możliwość rozszerzenia, wobec tego tworzy się drzewo podobne do drzewa w grze z dopuszczalnymi ruchami.

Algorytm UCT pozwala nam wybrać najbardziej obiecujące ready, a następnie szybko przeprowadzić doświadczenie pokazujące na ile ten wybór będzie dawał szanse na rozwinięcie kontigu w kolejnych ruchach. W zależności od tej symulacji liście w aktualnym drzewie będą dostawały wartości odpowiadające wartości stopy zwrotu w oryginalnym problemie.

To tylko zarys rozwiązania, które nie jest jeszcze zakończone, jednak wydaje się obiecujące. Pokazuje przy tym jak tak, wydawałoby się, odległa dziedzina jak algorytmy dla gier może być przydatna dla problemów bioinformatyki. Jest jeszcze wiele innych metod uczenia maszynowego i sztucznej inteligencji, które mogą być użyte w bioinformatyce. To wciąż młoda dziedzina i oczekuje na sprytnych, którzy będą w stanie dostrzec podobieństwa dla wykorzystania innych metod.

Duże dane – duże możliwości?

Michał Foryt

Student Informatyki na Wydziale Matematyki i Informatyki UJ
Wiceprzewodniczący Koła Robotyki i Sztucznej Inteligencji

Od zarania dziejów wiedza była tym, co napędzało ludzkość. Jeśli jakieś plemię wiedziało gdzie można znaleźć wodę, schronienie i pożywienie, miało znacznie większe szanse na przeżycie. Przez wieki królowie wybierali się ze swoją służbą na łowy nie tylko po to, by zwiększyć swoje umiejętności łowieckie, ale także by poznać ukształtowanie terenu i dowiedzieć się, które miejsce najlepiej będzie się nadawać na ewentualną zasadzkę w razie wojny. Gdyby nie zdobywana przez wieki wiedza o ludzkim ciele, ludzie do dzisiaj umieraliby mając po 30-40 lat. Przez lata informacje gromadzone były na glinianych tabliczkach, zwojach papirusu czy też w bogato zdobionych księgach. Dzisiaj zdecydowanie częściej korzystamy z płyt DVD, kart pamięci, dysków twardych albo potężnych serwerów, które są w stanie przechować ogromne zbiory danych.

Big Data – początki

Chociaż termin Big Data pojawił się po raz pierwszy dopiero w 1999 roku, to jednak o gromadzeniu informacji na tego typu skalę zaczęto mówić już w latach 40. XX wieku. Wtedy to, na łamach „The Scholar and the Future of the Research Library” Fremont Rider szacował, że z powodu ogromnego napływu danych, amerykańskie biblioteki będą podwajały swoje zbiory średnio co 16 lat. Według jego szacunków, w 2040 roku biblioteka Uniwersytetu w Yale będzie posiadała w swoich księgozbiorach 200 milionów egzemplarzy książek. Choć liczba ta do dzisiaj robi wrażenie, to jednak warto przypomnieć, że na małym czytniku możemy pomieścić 3000 książek. Tymczasem standardem w świecie Dużych Danych są klastry o pojemnościach 1 Petabajta (czyli 1 000 000 GB), natomiast firma IBM szacuje, że 90% wszystkich danych powstało w ciągu ostatnich dwóch lat.

Pytanie brzmi – jakie dane zajmują tyle miejsca i do czego są one wykorzystywane?

VVVV

Na Big Data składają się 4 wymiary, tak zwane 4V. Są to **volume**, **variety**, **velocity**, **value**.

Volume, czyli ilość danych, która jest tak wielka, że nie sposób stosować zwyczajnych narzędzi do ich analizy, gromadzenia i przetwarzania. Przy czym nie określamy tutaj jakiejś konkretnej wielkości, ale skupiamy się na technologicznych możliwościach ich przetwarzania.

Variety oznacza różnorodność danych. Ponieważ zazwyczaj rejestrujemy wszystkie elementy dotyczące interesującego nas zagadnienia, np. transakcji czy funkcjonowania serwisu społecznościowego, toteż zbierane przez nas dane będą bardzo szybko się zmieniać, a także nie będą posiadać struktury pozwalającej przeanalizować ich w tradycyjny sposób.

Velocity to szybkość z jaką napływają nowe dane i jak szybko są one analizowane, a dzieje się to nieustannie. Dlatego też, aby móc wyciągać trafne wnioski z napływających danych, należy analizować je na bieżąco.

Przez **Value** rozumiemy wartość danych. Mając w bazie np. miliard postów z Facebooka do przeanalizowania musimy wyodrębnić informacje, które są dla nas mało istotne (takie jak nowe „selfie” w lustrze nastolatki w Kanadzie), od danych mających kluczowe znaczenie, np. film przyszłego zamachowca-samobójcy, który właśnie żegna się przed misją w centrum miasta.

Analizowanie dużych zbiorów danych może służyć do najróżniejszych celów. Sprawdzenia który produkt najlepiej sprzedaje się wśród klientów o określonej zasobności portfela, wieku itp. poprzez wychwytywanie zagrożeń związanych z terroryzmem, na zastosowaniach w medycynie kończąc.

Dobrze poinformowani

Jeżeli chodzi o zastosowania technik przetwarzających duże ilości danych, to pierwszym przykładem który przytoczę będzie walka z epidemią

eboli w Afryce. Według BBC firma Orange Telecom we współpracy z Flowminder opracowała mapę wędrówek dużych grup ludzi w krajach, które były najbardziej zagrożone wybuchem epidemii. Firma w czasie rzeczywistym zbierała dane o sygnale z 150 tysięcy telefonów. Na ich podstawie opracowano miejsca w których najlepiej będzie założyć obóz medyczny, a także gdzie prawdopodobieństwo zarażenia ebolą jest największe.

Komu z nas nie zdarzyło się zapomnieć o zażyciu lekarstwa o odpowiedniej porze? Firma farmaceutyczna Express Scripts Holding przeanalizowała zachowania swoich pacjentów. Okazało się, że ci którzy deklarowali, że regularne przyjmowanie medykamentów jest dla nich najważniejsze paradoksalnie najczęściej zapominali ich przyjmować. Badanie to zainspirowało koncern do wynalezienia nowego „leku” - elektronicznej kapsułki, która przypomina o zbliżającej się porze zażycia leku.

Mamy dane – i co dalej?

Gdy już mamy zebraną wystarczającą ilość danych, musimy je w jakiś sposób przeanalizować. Tutaj z pomocą przychodzi nam data mining, czyli proces analityczny przeznaczony do badania dużych zasobów danych w poszukiwaniu wzorców oraz systematycznych współzależności między zmiennymi. Następnie wyniki są oceniane przez zastosowanie wykrytych wzorców do nowych podzbiorów danych. Zazwyczaj przy analizie informacji, które uzyskaliśmy zależy nam by przewidzieć przyszłość, np. zachowanie określonej grupy klientów, prawdopodobieństwo ich utraty. Jest to tak zwany predykcyjny data mining. Jest on niezwykle popularny, ponieważ daje bezpośrednie korzyści biznesowe. Składa się zazwyczaj z trzech etapów.

Pierwszym z nich jest eksploracja. Zaczyna się ona od przygotowania danych, czyli ich czyszczenia, wyboru przypadków i wybranie tych cech, na których analizie zależy nam najbardziej, oraz wyznaczenia ogólnej natury i stopnia złożoności modelu.

Drugim etapem jest budowa i ocena modelu. Tutaj kryterium oceny jest jakość predykcji, czyli poprawność wyznaczania wartości modelowanej zmiennej, i stabilność wyników dla różnych prób.

Etap trzeci to wdrożenie i stosowanie modeli. Jest to końcowy etap, na którym stosujemy dla nowych danych model uzyskany i uznany za najlepszy w drugim etapie. Robimy to, aby uzyskać przewidywane wartości lub sklasyfikowanie danych.

Wielka przyszłość

Wszystko wskazuje na to, że era big data dopiero nadchodzi. W czasach, gdy większość przedsiębiorstw stawia na komputeryzację i zbieranie danych, już niedługo może powstać nowa gałąź biznesu, polegająca jedynie na zbieraniu i analizowaniu danych. Dzięki danym pozyskiwanym z różnego rodzaju czujników pochodzących z inteligentnych urządzeń, takich jak RTV, AGD czy samochody, będzie można tworzyć produkty niemalże szyte na miarę dla klienta.

Bibliografia:

[1]<http://msp.money.pl/intel/arttykul/restauracje-formula-1-oraz-medycyna-zobacz,88,0,1942872.html>

[2]http://www.statsoft.pl/textbook/stathome_stat.html?http%3A%2F%2Fwww.statsoft.pl%2Ftextbook%2Fstanman.html

[3]http://wyborcza.pl/1,91446,17427995,Gartner_big_data_przyszloscia_biznesu.html

[4]<http://natemat.pl/119085,to-oni-beda-najbardziej-poszukiwani-na-rynku-pracy-analityk-data-mining-oto-zawod-przyszlosci>

[5] https://pl.wikipedia.org/wiki/Eksploracja_danych

[6]http://www.statsoft.pl/textbook/stathome_stat.html?http%3A%2F%2Fwww.statsoft.pl%2Ftextbook%2Fstdatmin.html

[7] http://data-mining.wyklady.org/wyklad/297_eksploracja-danych-data-mining.html

Zagrożenia i nadzieje związane ze stworzeniem i rozwojem sztucznej inteligencji

Aleksander Baranowski

Student Informatyki na Wydziale Matematyki i Informatyki UJ, II rok
Prezes Koła Robotyki i Sztucznej Inteligencji

Mam nadzieję, że nie jesteśmy tylko biologicznym bootloaderem dla cyfrowej superinteligencji. Niestety jest to coraz bardziej prawdopodobne.

Elon Musk twórca PayPal, SpaceX i Tesla Motors.

Jako dziecko miałem niebywałą przyjemność czytać zbiór opowieści o pilocie Pirxie autorstwa jednego z najznamienitszych polskich pisarzy XX wieku. Najstraszniejszą, dla mojego dziecięcego rozumu, postacią w całym cyklu był android, który chciał zmusić dzielnego pilota statku kosmicznego do błędu. W efekcie jego działań androidy miałyby w przyszłości pełnić coraz ważniejsze funkcje, by w końcu zdominować gatunek ludzki. Jak pokazuje ten prosty obraz z mojego dzieciństwa sztuczna inteligencja (ang. AI) posiada zdecydowanie negatywne kolokacje, wystarczy tutaj wspomnieć choćby kilka filmowych klasyków takich jak *Terminator*, *Matrix* czy *2001: Odyseja kosmiczna*. W pracy tej postaram się, przynajmniej pobieżnie, opisać najczęściej spotykane wady i zalety stworzenia oraz rozwoju AI, a także omówić pokrótce możliwości jej zaistnienia.

Anders Sandberg i Nick Bostrom w swojej pracy opisują drogę jaką musimy przebyć do możliwości zasymulowania całego ludzkiego mózgu. Z ich obliczeń wynika, iż minimalna moc obliczeniowa potrzebna do tego celu waha się pomiędzy 10^{18} , a 10^{25} FLOPSów. Do 2018 roku IBM planuje stworzyć eksaflopowy komputer, czyli posiadający dokładnie 10^{18} FLOPSów. Potwierdza to tylko iż z hardwarewego punktu widzenia (o ile wymogi licencyjne na to

pozwalają, np. wolna licencja) stoimy u progu pełnej symulacji ludzkiego mózgu. Posiadanie sprzętu, który ma wystarczającą moc obliczeniową, by zasymulować ludzki mózg nie jest oczywiście jednoznaczne ze stworzeniem AI, aczkolwiek jest to kamień milowy w tej dziedzinie.

Inny znawca tematu Eliezer Yudkowsky, amerykański badacz samouk, zwraca uwagę, iż umysł, który byłby sztuczną inteligencją będzie, dla ludzi, tak samo zrozumiały jak hipotetyczna rasa kosmitów. Jego sposób rozumowania będzie prawdopodobnie zupełnie inny, niż każdy jaki do tej pory spotkaliśmy. Fakt ten wynika z naszego sposobu postrzegania świata, mamy w zwyczaju nadawanie twórcom których nie rozumiemy, cech ludzkich. W efekcie wiele osób wyobraża sobie AI jako byt, który będzie zainteresowany zniszczeniem ludzkości (cokolwiek miało by to znaczyć). Istnieje jednak możliwość zupełnie odwrotna, umysł taki mógłby być najbardziej altruistyczną istotą jaką znamy. Taki stan jest nazywany w literaturze przedmiotu jako przyjazna sztuczna inteligencja (ang. Friendly artificial intelligence). Termin ten został w 2008 roku zaproponowany przez wyżej wymienionego naukowca.

Szybkość z jaką mógłby się uczyć umysł AI jest prawdopodobnie ograniczona tylko i wyłącznie do zbioru danych jaki będzie mógł dostać. W czasach kiedy tzw. Big Data robi zawrotną karierę, a większość danych jest trzymana w tzw. chmurach. Ilość informacji, którą możemy dostarczyć takiemu umysłowi jest olbrzymia. Idealnym przykładem jest Watson, projekt IBM-u, który wygrał jeden z teleturniejów w Stanach Zjednoczonych - Jeopardy. Watson w chwili obecnej uczy się jak wykrywać raka. Istnieje olbrzymie prawdopodobieństwo, iż zostanie on najlepszym onkologiem w historii ludzkości. Pomimo braku przyjaznej sztucznej inteligencji, możemy zobaczyć jedno z zastosowań tego typu tworu. Oprócz diagnostyki, AI będzie miała szerokie zastosowanie w dziedzinie medycyny. Jest to pierwsza nadzieja jaką ludzkość wiąże ze sztuczną inteligencją.

Innym dobrym zastosowaniem sztucznej inteligencji może być wykorzystanie jej do skomplikowanych lub stresujących prac. Wyobraźmy

sobie maszynę, która sortuje odpady, ale robi to dużo lepiej niż obecnie. Rozwiązałoby to problem elektrośmieci, które są jednymi z najbardziej niebezpiecznych i często trudnych do recyklingu pozostałości. Niektóre kolejki już w chwili obecnej są obsługiwane przez komputery, w przyszłości AI mogłaby sterować wszystkimi autami na danym obszarze, co zdecydowanie upłynniłoby ruch i przyczyniłoby się do zmniejszenia zużycia energii. Jazda po terenach górzystych i trudno dostępnych często wymaga stalowych nerwów oraz wiąże się z olbrzymim ryzykiem błędu ludzkiego. Posyłając pojazd sterowany przez sztuczną inteligencję zmniejszylibyśmy ryzyko do minimum. Chciałbym tutaj wspomnieć o aspekcie ekonomicznym takiego zastosowania AI. Koszt wychowania człowieka do momentu kiedy może pracować to ponad 320 tysięcy dolarów (dane dla dziecka uczęszczającego do uczelni państwowych, w USA, koszt obejmuje utrzymanie do 18 roku życia) . Dużo taniej od przygotowania człowieka do danej pracy, wyjdzie zainwestowanie w sprzęt mogący obsłużyć umysł AI. Wynika to z tego, iż w świecie IT kopiowanie praktycznie nic nie kosztuje, posiadając już jeden taki umysł moglibyśmy go skopiować na inny sprzęt praktycznie za darmo. Co więcej mówimy tutaj tylko o jednym człowieku, w przypadku kiedy sztuczna inteligencja miałaby odpowiadać na przykład za flotę ciężarówek sytuacja wygląda jeszcze korzystniej dla takiego rozwiązania.

AI posiada także szereg zagrożeń, łącznie z takimi które ciężko nam na chwilę obecną przewidzieć. Należą do nich eksterminacja ludzkości (celowo lub w ramach awarii). Możliwość automodyfikacji (kierowanej ewolucji) w celu obejścia zabezpieczeń. Odczłowieczenie ludzi, rozleniwienie ich do tego stopnia, iż sami staną się niewolnikami własnej ignorancji. Transfer wiedzy z zasobów ludzkich do swoich własnych. Zagrożenie to jest o tyle istotne, iż wraz z czasem ludzkość mogłaby stracić bezpowrotnie wiedzę na temat pewnych odkryć i osiągnięć. W przypadku posiadania odpowiedzialnej funkcji, awaria AI mogłaby być bardzo brzemenna w skutkach.

Na sam koniec chciałbym zwrócić uwagę na bezbronność gatunku ludzkiego wobec potencjalnej sztucznej inteligencji. Czas jaki będzie

potrzebować ludzkość do zwiększania choćby dwukrotnie swojej inteligencji i możliwości poznawczych jest olbrzymi, szczególnie jeśli pomyślimy o tym o ile szybciej rozwija się moc obliczeniowa komputerów. Na dodatek zmiany jakie będą zachodzić na korzyść sztucznej inteligencji są zrównoleglone, szybsza jednostka obliczeniowa i pamięć, większe zbiory do nauki, zoptymalizowane oprogramowanie, szersze sieci połączeniowe i szyny wymiany danych stawiają ewolucję ludzkiego intelektu na przegranej pozycji. Daje to prostą konkluzję, że stworzymy byt, którego ewolucja będzie dla nas zupełnie nieprzewidywalna i pomimo wszelkich starań nie będziemy w stanie jej zrozumieć i nad nią zapanować.

Bibliografia:

- [1] Russell S. i Norvig P., Artificial Intelligence: A Modern Approach, Pearson, 2009.
- [2] Sotala K., „Advantages of Artificial Intelligences, Uploads, and Digital Minds”, 2012.
- [3] Sandberg A. i Bostrom N., Whole Brain Emulation: A Roadmap, 2008.
- [4] Yudkowsky E., „Artificial Intelligence as a Positive and Negative Factor in Global Risk”, 2008.
- [5] <http://www.ibm.com/smarterplanet/us/en/ibmwatson/watson-oncology.html>
- [6] https://en.wikipedia.org/wiki/Watson_%28computer%29

Uczenie przez wzmocnienie – od plastyczności synaptycznej do uczenia maszynowego

Agata Szlaga

Studentka Neurobiologii, IV rok
Zakład Neurofizjologii i Chronobiologii

Uczenie przez wzmocnienie polega na zapamiętywaniu pewnej reakcji lub faktu przez skojarzenie z drugim bodźcem, który jest nagradzający. Może być domeną samorozwijającej się sztucznej inteligencji, ale przede wszystkim jest biologicznym faktem. W mózgu takie skojarzenie może nastąpić na poziomie dwóch neuronów, które tworząc pomiędzy sobą synapsę konsolidują dwie informacje, które wcześniej były z sobą niepowiązane, tworząc ślad pamięciowy. Tymi dwiema informacjami mogą być umiejętności ruchowe, które zostają połączone w płynną sekwencję ruchów albo procesy czysto umysłowe. Reagowanie na nagrodę jest dla żyjących organizmów ważne z punktu widzenia przystosowania do środowiska. Nagroda działa na dość pierwotne ewolucyjnie struktury śródmózgowia, które wysyłają projekcje do wielu innych struktur. Takie rozgałęzione unerwienie sprawia, że mózg jest ukierunkowany na osiągnięcie celu, który został wyznaczony przez komórki dopaminergiczne śródmózgowia. Komórki te sygnalizują wystąpienie nagrody, a ich aktywność wpływa na aktywność innych rejonów mózgu. Dlatego uczenie przez wzmocnienie jest tak skuteczne – koresponduje z pierwotnymi zapędami do wykonania czynności ważnych dla przeżycia organizmu (jak zdobycie pokarmu), czy przekazania genów (realizacja potrzeb seksualnych). Zauważmy, że takie czynności uważane są za nagradzające. Z punktu widzenia uczenia maszynowego właśnie ta wartościująca informacja trenująca jest kluczowa, gdyż określa stan następny (od występującego przed wzmocnieniem) jako wyższy.

Klasyczne warunkowanie Pawłowa, w którym pies po łączeniu podawania pokarmu (bodźca bezwarunkowego wywołującego bezwarunkową, biologicznie uwarunkowaną i nie wymagającą uczenia się reakcję ślinienia) z sygnałem dzwonka lub zaświeceniem światła (bodziec warunkowy) nauczył się ślinić na sam dźwięk, a także warunkowanie instrumentalne polegające na nagradzaniu pożądanых zachowań – prowadzą do zwiększenia częstości występowania zachowania, które było wzmacniane. Pokarm w warunkowaniu Pawłowa jest bodźcem bezwarunkowym, ale poza tym z biologicznego punktu widzenia jest nagrodą, która wspomaga proces uczenia się łącznie ślinienia z bodźcem warunkowym i tworzeniem reakcji warunkowej (ślinienia bez podania jedzenia). Naukowcy zastanawiając się nad możliwym mechanizmem leżącym u samych podstaw uczenia się korelacji 'zachowanie' ↔ 'nagroda' odnieśli się do reguły Donalda Hebba, pamiętając o dynamicznej strukturze mózgu (tzw. plastyczności), gdyż każdy proces uczenia się musi być związany z funkcjonalną lub strukturalną zmianą w mózgu (zmianą plastyczną, plastycznością). Zmiany funkcjonalne dotyczą zmiany efektywności działania połączenia synaptycznego – efektywność może być większa lub mniejsza, a jej miernikiem jest amplituda prądów postsynaptycznych w neuronie. Zmiana strukturalna, to taka w której powstają zupełnie nowe połączenia lub przebudowywane są połączenia istniejące, np. wyrastają nowe kolce dendrytyczne, na powierzchni których tworzone są nowe synapsy albo istniejące kolce dendrytyczne zmieniają powierzchnię kontaktową z zakończeniami neuronów presynaptycznych i tworzona jest nowa synapsa na istniejącym już kolcu.

Reguła Hebba jest teorią postulującą jak działa plastyczność synaptyczna: „kiedy dwa neurony wykazują aktywność w niemal tym samym momencie, połączenie synaptyczne między nimi wzmacnia się”. Jak potem się okazało, Hebb nie do końca miał rację – taka synchroniczna aktywność musi się powtarzać, a siła połączenia synaptycznego może się nie tylko zwiększać, ale też słabnąć. Zauważyć można analogię między zmianami zgodnymi

z regułą Hebba, a częstszym występowaniem wzmocnionych zachowań: kiedy dwa bodźce współwystępują z sobą, zostają powiązane – dzwonek i ślinienie się, nagroda i wzmocniona reakcja. Zanim bodźce zostały sparowane (w sensie połączone w parę, zaprezentowane krótko po sobie) mózg miał przygotowaną reakcję, ale tylko na każdy z tych bodźców z osobna. Żeby poprawnie funkcjonować w środowisku potrzebne jest wykształcenie odpowiedniej reakcji: skoro dzwonek poprzedza podanie jedzenia (w fazie uczenia), to powinna być wydzielana ślina, która jest potrzebna w procesie trawienia jedzenia. Na takiej zasadzie po sparowaniu bodziec warunkowy wywołuje reakcję bezwarunkową. Natomiast kiedy wykonanie danego zadania wiąże się z otrzymaniem nagrody, to żeby zwiększyć szanse na otrzymanie nagrody, zachowanie jest powtarzane. Plastyczność, to taka zmiana połączeń w mózgu, która umożliwia pełnienie przez niego swoich funkcji: uczenie, adaptację do warunków środowiska, osiągnięcie celów.

Uaktualnioną regułę Hebba określono jako STDP – *spike-timing dependent plasticity*, czyli plastyczność zależna od wzorca aktywności. Kiedy neuron presynaptyczny (z którego tradycyjnie uwalniane są neuroprzekaźniki) wykaże aktywność elektryczną przed neuronem postsynaptycznym (odbierającym informację niesioną przez neuroprzekaźniki za pomocą obecnych na błonie receptorów), połączenie między nimi wzmocni się. Jeżeli błona komórkowa neuronu postsynaptycznego będzie zdepolaryzowana (neuron będzie wykazywał aktywność), przed neuronem presynaptycznym połączenie osłabi się. Zmiana siły połączenia synaptycznego może być mierzona w częstości uwalniania pakietów neuroprzekaźnika z komórki presynaptycznej lub w wysokości połowych potencjałów postsynaptycznych. Wzmocnienie siły synapsy objawia się zwiększeniem wartości obu tych parametrów. Oczywiście potrzebny jest pewien okres uczenia się, czyli sparowania aktywności neuronów presynaptycznego z aktywnością neuronu postsynaptycznego (w tej kolejności lub odwrotnie). STDP jest prostą zasadą o charakterze obliczeniowym dzięki czemu jest

stosunkowo łatwy w zastosowaniu przy badaniu modulujących właściwości substancji na plastyczność synaptyczną - paradygmat badawczy indukowania zmian plastycznych za pomocą STDP może pozostać ten sam.

Żeby wyjaśnić wpływ nagrody na to jak jeden bodziec jest kojarzony z drugim (asocjacyjny charakter STDP – aktywność jednego neuronu jest kojarzona z aktywnością drugiego przez zmianę siły połączenia synaptycznego między nimi) rozwinęto teorię modulowanego nagrodą STDP. Polega ono na tym, że obecność nagrody sygnalizowana jest przez neurony dopaminergiczne, uwalniające neuroprzekaźnik dopaminę w okolice synaps, co moduluje zmiany plastyczne, które na nich zachodzą. Nadal kluczowe jest czasowe współwystępowanie i kolejność aktywności neuronów po sobie, ale dzięki obecności dopaminy procesy te zachodzą inaczej niż bez niej. Można więc powiedzieć, że nagroda moduluje uczenie się, zwiększając możliwości powiązania kontekstu nagrody z samą nagrodą.

Modulowane nagrodą STDP było testowane na szeregu sztucznych sieci neuronowych i robotów, co z punktu widzenia rozwoju sztucznej inteligencji i konieczności uczenia się przez nią otaczającego ją środowiska jest bardziej istotne niż dokładny przebieg tych procesów w tkance biologicznej.

Robota zdolnego do nawiązania interakcji z człowiekiem i uczenia się na podstawie zachowań człowieka zbudowali Ting-Shuo Chou, Liam D. Bucci i Jeffrey L. Krichmar z Uniwersytetu Kalifornia. Robot CARL-SJR został pokryty czujnikami dotyku, gdyż w przeciwieństwie do tradycyjnych robotów, bardziej miał się kierować czuciową interakcją z człowiekiem, niż przesłankami wzrokowymi. Zaimplementowana generująca potencjały czynnościowe sieć neuronowa (spiking neural network – SNN) uczyła się przez wzmocnienie. Sieć składała się z neuronów pobudzających i hamujących, oba rodzaje modulowane przez dopaminę. Robotowi na korpusie wbudowano diody LED, którymi świecił na różne kolory (w kolejności losowej). Człowiek wchodzący

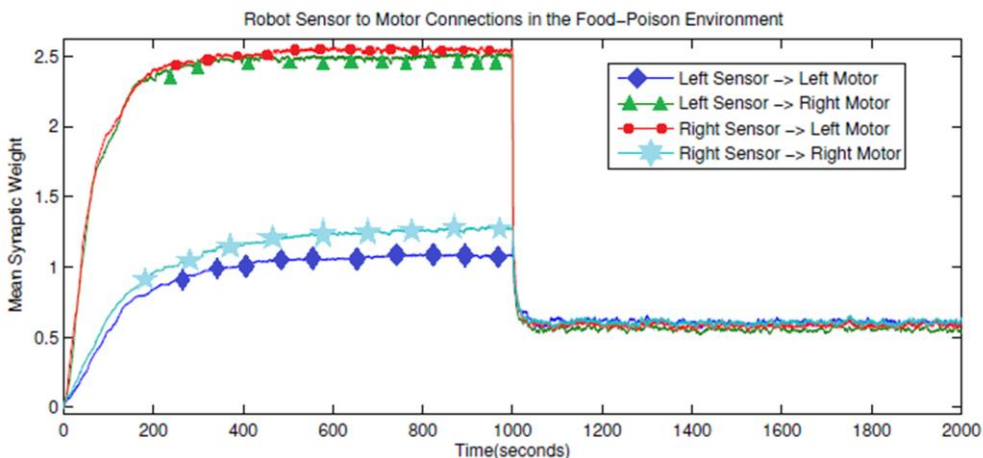
w interakcje z robotem mógł nagrodzić go, jeżeli kolor światła mu się podobał. Robot odbierał głaskanie jako nagrodę, jeżeli kierunek głaskania był odpowiedni (zaprogramowana preferencja kierunku głaskania) oraz jeżeli dotyk nastąpił odpowiednio szybko, czyli w czasie 2 sekund od zaświecenia diod (zgodnie z zasadami STDP i oknem czasowym na zmiany plastyczne). Robot uczył się więc przez modulowane przez nagrodę (a konkretnie dopaminę) STDP.

Odbierany przez CARL-SJR dotyk był bodźcem bezwarunkowym, gdyż reakcja na niego nie wymagała uczenia się, zaś kolor światła emitowanego przez niego był bodźcem warunkowym. Uczenie polegało na skorelowaniu informacji o tym jaki kolor światła wywołuje głaskanie i wydzielanie dopaminy. Po 40 połączeniach bodźca warunkowego (świecenie diod robota) i bezwarunkowego (dotyku przez człowieka) robot częściej świecił kolorami, które były nagradzane, a „wydzielanie” dopaminy przesunęło się u robota z bodźca bezwarunkowego na warunkowy. Takie przesunięcie jest charakterystyczne dla etapu, w którym uczenie przyniosło efekty i dla biologicznych neuronów wydzielających dopaminę nazywa się błędem przewidywania. Oznacza ono, że obiekt uczący się spodziewa się uzyskania nagrody, ponieważ wykonał czynność, która w etapie uczenia się przynosiła tę nagrodę. Dla lepszego zrozumienia przypominam eksperyment Pawłowa: pies po pewnym czasie warunkowania zaczął ślinić się na sam bodziec warunkowy (dzwonek/ światło lampy), ponieważ nauczył się, że oznacza on jedzenie. Zainteresowanych równaniami opisującymi właściwości neuronów budujących SNN, dokładną budową sieci i wszelkiego rodzaju wykresami odsyłam do oryginalnej pracy podanej w literaturze.

Drugim przykładem maszynowego uczenia przez wzmocnienie będzie robot uczący się funkcjonować w środowisku, z którego może pobierać korzyści (zbierać pojemniki z „jedzeniem” będącym nagrodą), ale czekają go również niebezpieczeństwa (pojemniki z „trucizną”). Umiejętność adaptacji

jest kluczowa dla przeżycia, pozwoliłaby sztucznej inteligencji nauczyć się środowiska i optymalnie w nim funkcjonować. Ogólna reguła uczenia się jest taka sama: dopamina wydzielona w odpowiedzi na otrzymanie „jedzenia” ułatwia wzmocnienie siły połączenia pomiędzy neuronami odpowiedzialnymi za pozyskanie tej nagrody. W warunkach, gdy w otoczeniu robota (pole 100cm × 100cm) znajdowało się tylko jedzenie (przypadkowo rozrzuconych 20 porcji) nauka zwracania się ku jedzeniu i zbierania go trwała 200 sekund, a poziom wykonania zadania (tempo zbioru) ustabilizował się na poziomie 1,5 zebranych porcji/sekundę. Zainteresowanie jedzeniem i możliwość zebrania go była wywołana wzmocnieniem połączeń synaptycznych pomiędzy prawym sensorem, a prawym i lewym motorem, czyli nauką skręcania. Połączenia lewego sensora z silnikami również się wzmocniły, ale w mniejszym stopniu.

Następnie chciano sprawdzić jak szybko robot jest w stanie nauczyć się awersji pokarmowej. Zetknięcie z trucizną powoduje zmniejszenie wydzielania dopaminy poniżej poziomu bazowego, gdyż jest doświadczeniem niekorzystnym. Po tym jak robot nauczył się być zainteresowany jedzeniem (ponownie pole 100cm × 100cm, 20 sztuk jedzenia, 1000 sekund), jedzenie zastąpiono truciznami i próbę kontynuowano przez dalsze 1000 sekund. Wszystkie połączenia synaptyczne gwałtownie obniżyły swoją wartość, do wartości dwukrotnie poniżej poziomu siły połączeń lewego sensora z silnikami z części w której w środowisku występowało tylko jedzenie. W środowisku, gdzie jedzenie zostało zastąpione trucizną, robot zbierał przedmioty na poziomie losowym, nie unikając ich jednak, gdyż nie był nagradzany za unikanie i nie nauczył się go.



(a) Synaptic Weights in the Food-Poison Environment

Ilustracja1: fragment Figure 4.14 z oryginalnejpracyRichardaEvansa "Reinforcement learning in a neurally controlled robot using dopamine modulated STDP", zaadaptowano"

Gdyby u robota bazowym poziomem dopaminy nie było zero, a wartości dodatnie, byłby on nagradzany za unikanie (nie napotykaając trucizny nie obniża poziomu dopaminy, która jest wtedy na poziomie bazowym, który gdy jest dodatni i może pozytywnie modulować plastyczność połączeń odpowiedzialnych za unikanie) i szybko mógłby się nauczyć unikania, tak jak nauczył się jechać w kierunku jedzenia. W warunkach biologicznych dopamina jest stale uwalniana z neuronów dopaminergicznych, pozwalając na wykonywanie ruchów, a nagroda powoduje zmianę wzorca aktywności neuronów dopaminergicznych na owocującą zwiększonym od bazowego wydzielaniem neuroprzekaźnika.

Podsumowując, modulowana dopaminą plastyczność zależna od wzorca aktywności to biologicznie prawdziwy sposób na uczenie maszynowe, umożliwiający naukę przez doświadczenie, a co za tym idzie adaptację do otoczenia i interakcję z nim. Może to prowadzić do uczenia coraz bardziej skomplikowanych zachowań i sekwencji ruchów, a dzięki względnej prostocie samego STDP nie będzie stanowiło problemu programistycznego.

Bibliografia:

- [1] Chou T-S i in., „Learning touch preferences with a tactile robot using dopamine modulated STDP in a model of insular cortex”, *Frontiers in Neurorobotics*, 22 lipca 2015, wolumen 9, artykuł 6.
- [2] Evans R., „Reinforcement learning in a neurally controlled robot using dopamine modulated STDP”, Imperial College London Department of Computing, badawczapracamagisterska, wrzesień 2012.

Randomizowane sieci neuronowe

Agnieszka Słowik

Studentka III roku Informatyki

Wiceprzewodnicząca Naukowego Koła Robotyki i Sztucznej Inteligencji

Wspólną cechą struktury wszystkich sieci neuronowych jest obecność neuronów połączonych ze sobą synapsami. Obserwacja reakcji neuronów biologicznych na sygnały z zewnątrz doprowadziła do wniosku, że neuron nie tylko przewodzi informację w postaci sygnału elektrycznego, ale także potrafi określić znaczenie (wagę) każdego sygnału dla organizmu. Obserwacje te zostały wykorzystane przy tworzeniu modeli sztucznych neuronów i sztucznych sieci neuronowych. Uczenie sieci neuronowych oznacza m.in. dobieranie takich wartości wag wejść do neuronów sieci, aby ostateczny wynik obliczeń był poprawny. Pojęcie randomizowanych sieci neuronowych obejmuje modele zakładające, że nauce podlega jedynie warstwa wyjściowa, a wagi warstwy ukrytej są losowane z rozkładu prawdopodobieństwa, dzięki czemu problem sprowadza się do uczenia liniowego modelu. Celem artykułu jest krótkie wprowadzenie do rodzajów sieci spełniających to założenie: RadialBasisFunction Network (RBF), Extreme Learning Machines (ELM) i Extreme Entropy Machines (EEM).

RBF Network

Standardowa sieć neuronowa wykorzystuje jedną funkcję aktywacji, co pozwala podzielić przestrzeń cech na obszary odpowiadające poszczególnym klasom. Wszystkie neurony sieci biorą udział w procesie takiego podziału. Tymczasem w sieciach RBF zamiast globalnej funkcji aktywacji stosujemy różne tzw. radialne funkcje bazowe, które definiujemy dla każdego neuronu z osobna. Pozwala to neuronom na lokalne wyodrębnianie obszarów wokół wybranych punktów [1].

Funkcje radialne to podzbiór funkcji rzeczywistych, których wartość maleje lub wzrasta wraz z odległością od określonego punktu. Najczęściej wykorzystywanym przykładem jest funkcja Gaussa.

Teoretycznie radialne funkcje mogą być zastosowane dla każdego modelu (liniowego i nieliniowego) oraz dla każdej sieci (jedno- lub wielowarstwowej), jednak najczęściej sieci RBF mają jedną warstwę ukrytą.

Extreme Learning Machines

Sieci ELM to jednokierunkowe sieci neuronowe o pojedynczej warstwie ukrytej. Najprostszy algorytm trenujący opisuje wzór:

$$\hat{Y} = W_2 \sigma(W_1 x)$$

gdzie W_1 - macierz wag warstwy ukrytej, σ - wybrana funkcja aktywacji, W_2 - macierz wag warstwy wyjściowej. Wartości W_1 są losowane z rozkładu prawdopodobieństwa jak w pozostałych typach randomizowanych sieci neuronowych. Macierz W_2 estymuje się z użyciem metody najmniejszych kwadratów i pseudoodwrotności macierzy:

$$W_2 = \sigma(W_1 x)^+ Y$$

Twórca terminu Extreme Learning Machines Guang-Bin Huang spotkał się z zarzutami wymyślania nowej nazwy dla już istniejących sieci (RBF). O cechach wyróżniających sieci ELM oraz ich znaczeniu można poczytać w jego pracy „What are Extreme Learning Machines? Filling the Gap Between Frank Rosenblatt's Dream and John von Neumann's Puzzle” [2].

Extreme Entropy Machines

Model zaproponowany w tym roku przez doktora Wojciecha Czarneckiego i profesora Jacka Tabora z naszego uniwersytetu [3]. Posiada cechy wspólne z modelami Extreme Learning Machines (krótki czas uczenia, randomizacja) i Support Vector Machines. EEM opiera się na teorii informacji i zakłada użycie miary entropii do optymalizacji klasyfikatora. Ta metoda pozwala na osiągnięcie efektów porównywalnych z działaniem SVM, ELM i LS-SVM (Least Squares Support Vector Machines) przy lepszej skalowalności. Co więcej, model ma zwięzłą implementację i małą ilość prostych do zoptymalizowania hiperparametrów. Twórcy porównali działanie EEM i konkurencyjnych modeli na zbiorach danych liczących od kilkuset do setek tysięcy elementów - osobno dla zbalansowanych problemów o małej ilości wymiarów oraz dla wielkich, niezbalansowanych zbiorów. Przy rozwiązywaniu mniej złożonych problemów EEM dla niektórych zbiorów danych osiągnął najlepsze wyniki ze wszystkich testowanych modeli niezależnie od użytej funkcji aktywacji. Wykorzystanie niezbalansowanych zbiorów danych pokazało, że EEM zachowuje krótki czas uczenia nawet dla bardzo dużej liczby przykładów. Ponadto model EEM jest bardzo stabilny w kontekście rozmiaru warstwy ukrytej.

Zastosowania w bioinformatyce

Randomizowane sieci neuronowe mają zastosowanie przy małych i niedużych zbiorach danych, w szczególności dla problemów wymagających szybkiego uczenia i przeuczenia.

W bioinformatyce sieci RBF zostały już wykorzystane m.in do predykcji ruchu białek transportujących [4], Algorytmy ELM do klasyfikacji sekwencji białek [5], a Extreme Entropy Machines w celu predykcji aktywności związków chemicznych [6].

Bibliografia:

- [1] Orr M. J. L., Introduction to Radial Basis Function Networks, 1996.
- [2] Huang G-B., „What are Extreme Learning Machines? Filling the Gap Between Frank Rosenblatt’s Dream and John von Neumann’s Puzzle”, 2015.
- [3] Czarnecki W.M. i Tabor J., „Extreme entropy machines: robust information theoretic classification”, 2015.
- [4] Chen S-A i in., „Prediction of transporter targets using efficient RBF networks with PSSM profiles and biochemical properties”, 2011.
- [5] Cao J. i Xiong L., „Protein Sequence Classification with Improved Extreme Learning Machine Algorithms”, 2014.
- [6] Czarnecki W.M. i Rataj K., „Compounds Activity Prediction in Large Imbalanced Datasets with Substructural Relations Fingerprint and EEM”, 2015.

Sprawozdanie z działań Koła

Paulina Kania

Prezes Koła Naukowego Bioinformatyki BIT
Studentka Informatyki na Wydziale Matematyki i Informatyki UJ, IV rok

Zgodnie z zapowiedziami z poprzedniego numeru trochę się w naszym Kole dzieje!

Ruszyły w końcu przygotowania już piątej edycji ogólnopolskiej konferencji **Liczy Komputery Życie**. Chcemy przygotować ciekawe warsztaty, sesje i dotrzeć do jak największej liczby osób z całej Polski. Takie wydarzenie jednakże samo się nie zorganizuje! Dlatego jeśli ktoś chciałby się dowiedzieć jak to wygląda „od kuchni” zachęcam do kontaktu z nami poprzez fanpage lub adres mailowy podany poniżej.

Ruszamy również pełną parą z **cyklem warsztatów** z programowania w języku **Python**. Są one przeznaczone dla wszystkich osób, które chcą się tego nauczyć.

Zaczęliśmy również pracę nad nową **stroną internetową**. Przez najbliższy okres obecna strona nie będzie utrzymywana więc najświeższe informacje będzie można zawsze znaleźć na Facebook’u.

Jak się z nami skontaktować?

- Adres kołowy: **kolobioinformatyki@gmail.com**
- Adres redakcji: **bioinfoworld.redakcja@gmail.com**
- Nasz fanpage: **www.facebook.com/koloBITUJ**
- Jesteśmy w Sali **1160** na Wydziale Matematyki i Informatyki

Słowniczek dla biologów

Bootloader – program rozruchowy który uruchamia system operacyjny.

Entropia - średnia ważona ilości informacji niesionej przez pojedynczą wiadomość, gdzie wagami są prawdopodobieństwa nadania poszczególnych wiadomości.

FLOPS - FLoating point Operations Per Second – liczba operacji zmiennoprzecinkowych na sekundę.

Funkcja aktywacji - pojęcie używane w sztucznej inteligencji do określenia funkcji, według której obliczana jest wartość wyjścia neuronów sieci neuronowej.

Hardware – fizyczna część maszyny

Software – oprogramowanie.

Pseudoodwrotność macierzy - uogólnienie macierzy odwrotnej na macierze prostokątne.

Teoria informacji - dyscyplina zajmująca się problematyką informacji oraz metodami przetwarzania informacji, powiązana naukowo z matematyką dyskretną.

Słowniczek dla informatyków

Neuron - elektrycznie pobudliwa komórka zdolna do przetwarzania i przewodzenia informacji w postaci sygnału elektrycznego, podstawowy element układu nerwowego.

Warstwa sieci neuronowej - sieć posiada warstwy zbudowane z pojedynczych neuronów, w obrębie których nie ma połączeń między neuronami (połączenia stosuje się pomiędzy poszczególnymi warstwami).