



10. Sztuczna inteligencja i uczenie maszynowe w łańcuchu dostaw

Czym jest sztuczna inteligencja (AI)? Czy naprawdę jest to „żywe” stworzenie zdolne do myślenia i podejmowania własnych decyzji w oparciu o swój umysł, przeszłe doświadczenia, etykę, przekonania itp. W jaki sposób jest ona powiązana z uczeniem maszynowym (ML)? Czy AI i ML to to samo?

Jakie narzędzia są wykorzystywane w AI i ML? Jaką rolę odgrywają AI i ML w kontekście biznesowym i jak można je wykorzystać w codziennych operacjach biznesowych i optymalizacjach? Czy istnieją konkretne architektury i przykłady zastosowania AI i ML w łańcuchach dostaw?

Na te i podobne pytania postaramy się udzielić odpowiedzi w następnym rozdziale, zamykając go rzeczywistym przykładem zastosowania algorytmów AI i ML w magazynie dystrybucyjnym.

10.1 Czym jest sztuczna inteligencja (AI)?

Dziedzina sztucznej inteligencji rozpoczęła się w latach pięćdziesiątych XX wieku, kiedy to informatycy zapytali: „Czy komputery mogą myśleć jak ludzie”? Naukowcy w tamtym czasie byli entuzjastycznie nastawieni do możliwości uczenia komputerów wykonywania złożonych zadań i w tym celu opracowali zestaw różnych algorytmów. Definicję tej dziedziny można określić jako wysiłek na rzecz automatyzacji zadań intelektualnych normalnie wykonywanych przez ludzi (Chollet, 2021). Najbardziej znaczące algorytmy w dziedzinie sztucznej inteligencji pochodzą obecnie z uczenia maszynowego i głębokiego uczenia, które są podzbiorami sztucznej inteligencji. Oprócz uczenia maszynowego i głębokiego uczenia, sztuczna inteligencja obejmuje wiele algorytmów nieuczenia się. Co więcej, możemy argumentować, że tego rodzaju algorytmy były bardziej dominujące we wczesnej fazie rozwoju sztucznej inteligencji. W związku z tym jest to część sztucznej inteligencji znana jako symboliczna sztuczna inteligencja, która opiera się na założeniu, że ludzki poziom wydajności i inteligencji można osiągnąć poprzez programowanie komputerów za pomocą dużego zestawu wyraźnych reguł rozwiązywania obserwowanego problemu. Podejście to zapewniało doskonałe wyniki w przypadku problemów logicznych, które były dobrze zdefiniowane, takich jak komputer



grający w szachy, ale okazało się skomplikowane w przypadku bardziej złożonych problemów. Rzeczywisty świat okazał się znacznie bardziej skomplikowany niż wszystkie jawne reguły programowania, które można było wprowadzić do komputera. Podejście to koncentruje się wokół idei dla danej sytuacji zrób to lub tamto (reguły if-then). Jest to łatwe do zrozumienia podejście, ale z drugiej strony bardzo czasochłonne i czasami bardzo trudne do określenia wszystkich możliwych scenariuszy, które należy wprowadzić do programu. Jako zabawny przykład, ale dobrze ilustrujący dany temat, przedstawiony na rysunku 10.1, który pokazuje kilka scenariuszy określania prognozy na podstawie stanu kamienia.



STAN	PROGNOZA
Kamień jest mokry	Deszcz
Kamień jest suchy	Brak deszczu
Cień na podłożu	Słonecznie
Biało na górze	Pada śnieg
Nie widać kamienia	Mgła
Kołyszący się kamień	Wietrznie
Podskakujący kamień	Trzęsienie ziemi
Kamień zniknął	Tornado

Rysunek 10.1 Reguły programowania Jeżeli – To

Źródło: Gibbs, M. 2019

Dziedzina symbolicznej sztucznej inteligencji zyskała największą popularność w latach 80. wraz z pojawieniem się systemów eksperckich (ES). ES reprezentują podzbiór systemów wspomagania decyzji (DSS) (Turban, 1998), koncentrując się na dostarczaniu skomputeryzowanych zdolności decyzyjnych podobnych do tych, które posiada ludzki specjalista w danej dziedzinie. Systemy te są stworzone do rozwiązywania skomplikowanych problemów poprzez zastosowanie szeregu reguł lub algorytmów, które symulują ludzkie procesy rozumowania. Olson i Courtney opisują systemy eksperckie (ES) jako programy komputerowe, które symulują ludzkie procesy myślowe w celu podejmowania decyzji w określonej dziedzinie, włączając pewien stopień sztucznej inteligencji, aby dopasować



wnioski, do których doszedłby ludzki ekspert (Olson & Courtney, 1992). Komponent ES jest szczególnie przydatny do wspierania decydentów w obszarach wymagających specjalistycznej wiedzy (Turban i in., 2005). Zasadniczo ES przechwytuje wiedzę specjalistyczną od ludzkiego eksperta (lub innego źródła) i przekazuje ją do komputera. Technologia ta może pomóc decydentom lub w pełni ich zastąpić, co czyni ją jedną z najczęściej stosowanych i odnoszących sukcesy komercyjne form sztucznej inteligencji (Turban i in., 2005). Jednym z kluczowych powodów rozwoju ES jest rozpowszechnianie wiedzy eksperckiej wśród szerszego grona odbiorców (Jackson, 1999). W kolejnych podrozdziałach zademonstrujemy zastosowanie ES opartego na algorytmach AI i ML w centralnym magazynie jako części ogólnego systemu DSS dla menedżerów.

Obecnie dziedzina sztucznej inteligencji składa się z różnych podejść i algorytmów, ale najczęściej używane z nich opisano na rysunku 10.2. Odeszła ona od systemów ES, a ponowne pojawienie się tej dziedziny przypisuje się głównie algorytmom głębokiego uczenia, które odniosły znaczący sukces w ciągu ostatnich 12 lat w problemach rozpoznawania obrazów, rozpoznawania mowy, segmentacji obrazów, rozpoznawania twarzy itp. Głębokie uczenie wykorzystuje wiele warstw abstrakcji do identyfikacji złożonych wzorców w danych wielowymiarowych. Podejście to pozwoliło osiągnąć znaczący postęp w takich dziedzinach jak rozpoznawanie mowy i obrazów, odkrywanie leków i przetwarzanie języka naturalnego. Głębokie uczenie ma zdolność do automatycznego odkrywania odpowiednich cech i zmniejsza potrzebę interwencji człowieka w projektowanie cech, dzięki czemu jest bardzo wydajne w wykorzystywaniu dużych zbiorów danych i mocy obliczeniowej (LeCun, Bengio i Hinton, 2015).



Rysunek 10.2 Główne dziedziny i poddziedziny sztucznej inteligencji

Źródło: Athanasopoulou i in., 2022

Dużym krokiem w kierunku dzisiejszej sytuacji były badania nad klasyfikacją cyfr przeprowadzone przez Hinton, Osindero i Teha (2006), którym udało się osiągnąć ponad 98% dokładności w klasyfikacji bazy danych Modified National Institute of Standards and Technology (MNIST). Jednym ze sposobów myślenia o tym, jak sztuczna inteligencja przeszła od symbolicznej sztucznej inteligencji do uczenia maszynowego i jaka jest istota uczenia maszynowego, jest wyobrażenie sobie algorytmów uczenia maszynowego jako amorficznej masy, która kształtuje się zgodnie z pożądanymi wynikami. System reguł, który zmienia dane wejściowe w wyjściowe, zmienia się z problemu na problem i dostosowuje się do istniejącej sytuacji. Jego celem jest znalezienie reguł, które zautomatyzują zadanie poprzez wyszukiwanie wzorców statystycznych w danych. Takie podejście do rozwiązywania różnych problemów znacznie skraca czas konfiguracji systemu (w porównaniu do reguł *if (jeżeli) – then (wtedy)*) i sprawia, że jest to bardziej uniwersalne podejście do rozwiązywania różnych problemów.



Bengio, Lecun i Hinton (2021) podkreślili, że przyszłość sztucznej inteligencji leży w głębokim uczeniu się i rewolucyjnym wpływie miękkiej uwagi i architektur transformatorowych w sztucznej inteligencji. Innowacje te pozwalają sieciom neuronowym dynamicznie koncentrować się na ważnych danych wejściowych i przechowywać informacje w zróżnicowanych pamięciach, znacznie poprawiając przetwarzanie sekwencyjne.

10.2 Czym jest ekosystem sztucznej inteligencji i uczenia maszynowego?

Ekosystem algorytmów AI i ML składa się z trzech kluczowych filarów:

- danych wejściowych,
- danych wyjściowych,
- funkcji kosztów.

Dane wejściowe reprezentują nagrania danych określonej cechy lub cech (w zależności od zaobserwowanego problemu). Dokładność danych wejściowych jest kluczowa dla budowania dokładnych algorytmów. Często nie ma to miejsca w rzeczywistych zastosowaniach i zwykle poświęca się dużo czasu i wysiłku na zbieranie danych, czyszczenie, porządkowanie, ujednolicanie, sprawdzanie fałszywych danych wejściowych itp. Oprócz dokładnych danych, innym ważnym aspektem charakterystyki danych wejściowych jest ich reprezentacja i kodowanie. Różne sposoby kodowania danych mogą ujawnić różne cechy danych i znacząco „pomóc” modelom ML w ujawnianiu ukrytych wzorców w danych. Tutaj widzimy piętę achillesową algorytmów ML. Często zbyt wiele uwagi poświęca się tworzeniu inteligentnych metod wydobywania informacji z danych (tj. samych algorytmów), podczas gdy niewystarczająca uwaga jest poświęcana danym wejściowym i ich związkom z danymi wyjściowymi. Ogólnie przyjmuje się za pewnik, że dane wejściowe mają związek przyczynowy z danymi wyjściowymi, co czasami wcale nie ma miejsca. Dlatego kolejnym dużym krokiem w rozwoju ML powinno być znalezienie lepszych sposobów gromadzenia, reprezentowania i kodowania danych.

Dane wyjściowe reprezentują pomiary konkretnego problemu, który próbujemy rozwiązać. W problemie klasyfikacji będzie to etykieta klasy. W regresji będzie to liczba rzeczywista, którą próbujemy przewidzieć. W kontekście konkretnego problemu, w przypadku problemów z rozpoznawaniem mowy, danymi wyjściowymi mogą być wygenerowane przez człowieka



transkrypcje plików dźwiękowych. W przypadku rozpoznawania obrazów, danymi wyjściowymi mogą być etykiety klas obrazów itp.

Funkcja kosztu reprezentuje sposób pomiaru wydajności AI i ML. Zasadniczo chcielibyśmy, aby odpowiedzi algorytmów pasowały do danych wyjściowych dla danych wejściowych. Funkcja kosztu jest również sygnałem zwrotnym do zestawu parametrów, które kierują pracą algorytmu, tj. pozwala na optymalizację ogólnej wydajności algorytmu poprzez proces uczenia się (znalezienie optymalnego zestawu parametrów). Proces uczenia się zazwyczaj obejmuje uczenie nadzorowane, w którym model jest szkolony na oznaczonych danych w celu zminimalizowania błędów przewidywania za pomocą technik takich jak stochastyczne opadanie gradientu i propagacja wsteczna. Umożliwia to modelowi skuteczne dostosowanie jego wewnętrznych parametrów, co prowadzi do poprawy wydajności w zadaniach takich jak wykrywanie i klasyfikacja obiektów (LeCun, Bengio, & Hinton, 2015).

10.3 Jakie narzędzia są wykorzystywane w uczeniu maszynowym?

Ogólnie rzecz biorąc, algorytmy uczenia maszynowego można podzielić na dwie główne kategorie: uczenie nadzorowane i nienadzorowane.

Uczenie nadzorowane obejmuje szkolenie algorytmów na etykietowanym zbiorze danych, w którym każdy punkt danych wejściowych jest sparowany z prawidłowym wynikiem. Ten jasny „obraz” tego, jaka powinna być prawidłowa odpowiedź dla danego wejścia, pozwala algorytmowi nauczyć się funkcji mapowania z wejść na wyjścia. W związku z tym znane są zarówno dane wejściowe, jak i wyjściowe (Athanasopoulou i in., 2022). Typowe zastosowania uczenia nadzorowanego obejmują zadania klasyfikacji (np. określanie, czy wiadomość e-mail jest spamem, czy nie) i zadania regresji (np. przewidywanie cen domów na podstawie różnych cech). Niektóre z najpopularniejszych algorytmów, które zostały sprawdzone w licznych zastosowaniach, to uogólnione modele addytywne, Random Forest, boosting, drzewa klasyfikacyjne i regresyjne, maszyny wektorów nośnych, rozszerzona regresja liniowa, regresja logistyczna, algorytm k-najbliższych sąsiadów, liniowa analiza dyskryminacyjna, lasso, sieci neuronowe, adaptacyjny system wnioskowania neurorozmytego itp. (Rostami-Tabar & Mircetic, 2023). Uczenie nadzorowane jest potężne, ponieważ wykorzystuje dane z adnotacjami człowieka, aby osiągnąć wysoką dokładność przewidywań. Jednak jego skuteczność zależy w dużej mierze od jakości i ilości etykietowanych danych.



Z kolei uczenie bez nadzoru dotyczy zbiorów danych, w których brakuje etykietowanych odpowiedzi. W związku z tym algorytmy uczenia maszynowego bez nadzoru wykorzystują nieoznakowane zbiory danych, które zawierają tylko dane wejściowe (Athanasopoulou i in., 2022). W tym przypadku algorytm otrzymuje tylko dane wejściowe, a jego celem jest znalezienie podstawowych wzorców, struktur lub relacji w danych. Typowe techniki uczenia bez nadzoru obejmują grupowanie (np. grupowanie klientów według zachowań zakupowych) i redukcję wymiarowości (np. zmniejszenie liczby zmiennych w zbiorze danych przy jednoczesnym zachowaniu ważnych informacji). Uczenie bez nadzoru jest cenne w przypadku analizy danych eksploracyjnych i odkrywania ukrytych struktur w danych. Jest często używane, gdy etykietowane dane są rzadkie lub niedostępne.

Inną ważną kategorią, choć różniącą się od uczenia nadzorowanego i nienadzorowanego, jest uczenie ze wzmocnieniem. W tym przypadku algorytm uczy się poprzez interakcję ze środowiskiem i otrzymywanie informacji zwrotnych w postaci nagród lub kar. To podejście oparte na metodzie prób i błędów pomaga algorytmowi nauczyć się optymalnych działań w celu maksymalizacji skumulowanych nagród. Uczenie ze wzmocnieniem jest szeroko stosowane w takich dziedzinach jak robotyka, gry i systemy autonomiczne.

Jednym z najbardziej popularnych i skutecznych algorytmów ML jest gałąź sieci neuronowych. Sieci neuronowe istnieją od lat 50. ubiegłego wieku, ale zyskały popularność w latach 80. w ciągu ostatnich 12 lat. Są one zbudowane na wzór biologicznych neuronów i sposobu, w jaki dzielą się one informacjami w mózgu, ale poza tym nie ma między nimi znaczących powiązań. Obecnie najczęściej stosowaną formą sieci neuronowych są sieci głębokiego uczenia, które reprezentują kilka warstw ukrytych między cechami wejściowymi i wyjściowymi, które wykonują kilka nieliniowych transformacji cech wejściowych. Ponieważ okazało się to bardzo skuteczne, uczenie głębokie jest obecnie jedną z najważniejszych poddziedzin ML (Chollet, 2021).

10.4 Studium przypadku

Ustalenia Wenzel, Smit i Sardesai (2019) dotyczące ML w zarządzaniu łańcuchem dostaw wskazują na rosnącą integrację aplikacji ML w różnych zadaniach SC. W związku z tym obserwowane studium przypadku przedstawia zastosowanie sztucznej inteligencji i uczenia maszynowego w centralnym magazynie fabryki żywności (Mirčetić i in., 2016; Mircetic i in., 2014). W kompleksie fabrycznym znajduje się 30 wózków widłowych. Wózki widłowe



są zaangażowane w różne operacje wewnątrz kompleksu, które są kluczowe dla operacji logistycznych w produkcji, magazynowaniu i wysyłce produktów. Centralny magazyn ma pojemność 11 100 miejsc paletowych i roczną produkcję od 300 000 do 350 000 palet. Obecnie fabryka zaopatruje około 20 000 supermarketów za pośrednictwem dostaw bezpośrednich.

Problem zaangażowania wózków widłowych jest związany z faktem, że nadmierne lub niedostateczne zaangażowanie wózków widłowych w różne procesy fabryczne prowadzi do znacznych strat finansowych i rynkowych. Obecnie proces podejmowania decyzji, gdzie i co będzie robił każdy wózek widłowy, opiera się na decyzjach ekspertów (menedżerów). Decyzje ekspertów opierają się na ich doświadczeniu, bez pomocy jakiegokolwiek systemu wspomagania decyzji (DSS). Liczne dowody empiryczne sugerują, że ludzka intuicyjna ocena i podejmowanie decyzji często nie są optymalne, szczególnie w warunkach złożoności i stresu (Druzdzel i Flynn, 2002). Podkreśla to znaczenie włączenia systemów wspomagania decyzji (DSS) celem wspomagania ekspertów w procesie podejmowania decyzji.

W tej aplikacji wybraliśmy kilka algorytmów ML, aby pomóc w optymalizacji pracy magazynu załadunkowego. Algorytmy ML zostały zebrane w unikalną strukturę decyzyjną, która służy jako DSS dla menedżerów i ekspertów w danej firmie. Co więcej, cały system decyzyjny DSS można postrzegać jako platformę sztucznej inteligencji, ponieważ stale przelicza on sugestie z kilku modeli ML (ile wózków widłowych należy użyć i które z nich) i automatycznie wybiera najlepsze z nich, biorąc pod uwagę dostarczone dane wejściowe operatorów.

Opis problemu

Proces załadunku ma kluczowe znaczenie dla logistyki magazynowej, wpływając na poziom obsługi rynku. Podczas wysyłek, ekspert magazynowy określa liczbę i wybór wózków widłowych do załadunku, kierując się trzema czynnikami: (1) ukończenie załadunku w określonych ramach czasowych, (2) zminimalizowanie zakłóceń w innych zadaniach wózków widłowych oraz (3) dostosowanie wykorzystania wózków widłowych do możliwości konserwacyjnych, które mogą obsługiwać dwa przeglądy jednocześnie. Każdy wózek widłowy przechodzi od czterech do pięciu przeglądów konserwacyjnych rocznie.

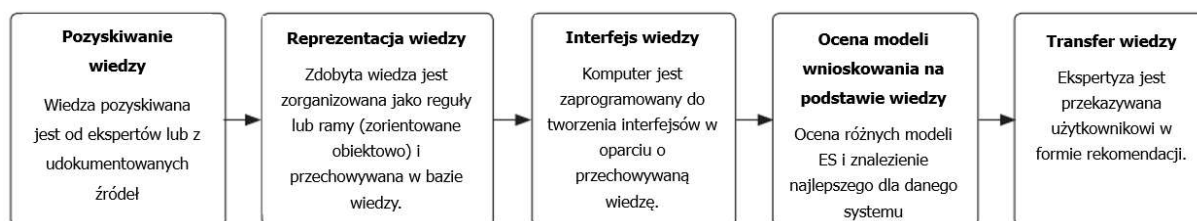
Wózki widłowe mają kluczowe znaczenie dla operacji załadunku, które muszą wspierać strategię marketingową firmy, zapewniając jednocześnie płynne działanie innych czynności. Nieprawidłowe przydzielanie wózków widłowych może prowadzić do niepełnego wykorzystania zasobów lub zaszkodzić reputacji firmy i poziomowi świadczonych usług. Opóźnienia w załadunku wiążą się z karami. Menedżer musi koordynować wykorzystanie wózków



widłowych we wszystkich działaniach, aby uniknąć jednoczesnych przeglądów i zarządzać różnymi potrzebami konserwacyjnymi. Chociaż menedżerowie zazwyczaj podejmują trafne decyzje, praca w środowisku charakteryzującym się wysokim poziomem stresu może prowadzić do błędów. Dlatego potrzebne jest wykorzystanie systemu DSS, aby zwiększyć pewność i niezawodność podejmowania decyzji.

AI i ML jako DSS dla centralnego magazynu

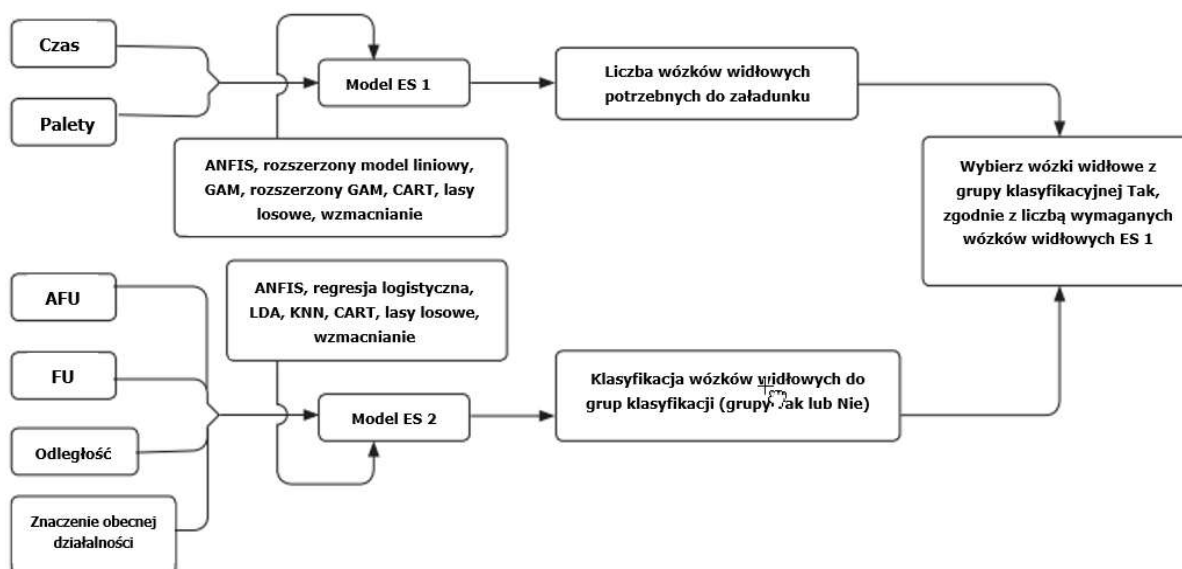
Pierwszym krokiem do wygenerowania systemów AI i ML jest pozyskanie stabilnego i poprawnego źródła wiedzy (bazy danych) oraz rzucenie światła na role biznesowe, które musi wspierać. Dlatego też rysunek 10.3 przedstawia metodologię budowy systemu AI i ML DSS.



Rysunek 10.3 Etapy metodologii tworzenia systemu AI i ML DSS

Źródło: Rainer i Turban, 2008; Turban, Aronson oraz Liang, 2005

Pozyskanie wiedzy zostało osiągnięte na drodze wywiadów z menedżerami, obserwację ich procesów decyzyjnych oraz przegląd dokumentacji magazynowej. Aby opracować system DSS dla danego problemu, utworzono dwie bazy wiedzy. Pierwsza baza wiedzy obejmuje decyzje dotyczące liczby wózków widłowych rozmieszczonych w strefie załadunku (434 decyzje eksperckie), podczas gdy druga obejmuje decyzje dotyczące tego, które wózki widłowe były używane (368 decyzji eksperckich) w różnych scenariuszach operacyjnych. Na etapie wnioskowania o wiedzy zastosowano kilka algorytmów ML przy użyciu oprogramowania Matlab: Adaptacyjny neuro-rozmyty system wnioskowania (ANFIS), uogólnione modele addytywne (GAM), Random Forest, boosting, drzewa klasyfikacji i regresji (CART), rozszerzona regresja liniowa, regresja logistyczna, algorytm k-najbliższych sąsiadów (KNN) i liniowa analiza dyskryminacyjna (LDA). Oceniono różne modele ML i zidentyfikowano te o najlepszej wydajności. ANFIS i CART wykazały lepsze wyniki i zostały wybrane jako ostateczne systemy DSS do praktycznego zastosowania w firmie. Transfer wiedzy został ułatwiony poprzez interfejs użytkownika ostatecznych modeli DSS. Strukturę i logikę systemu DSS zilustrowano na rysunku 10.4.



Rysunek 10.4 Budowanie struktury magazynu DSS w oparciu o algorytmy AI i ML

Struktura DSS składa się z warstwy wejściowej, zbudowanej z kilku kluczowych czynników, które wpływają na zaangażowanie wózków widłowych. Warstwa ML zawiera modele ML, które ponownie obliczają sugestie dotyczące liczby i wyboru wózków widłowych do wykorzystania w danym scenariuszu wejściowym. Najlepiej działające modele są wybierane jako modele systemu eksperckiego (modele ES), ponieważ baza wiedzy, na podstawie której tworzone są modele ML, jest pozyskiwana od ekspertów. Pierwszy model koncentruje się na określeniu liczby wózków widłowych wymaganych w strefie załadunku (model ES 1). Drugi model odnosi się do problemu wyboru konkretnych wózków widłowych, które powinny zostać zaangażowane (model ES 2). Oba modele zostały opracowane przy użyciu nadzorowanych technik uczenia maszynowego. Według Turban, Aronson i Liang (2005), uczenie maszynowe wykazało doskonałe wyniki w projektowaniu inteligentnych systemów wspomagania decyzji (DSS). Modele ES wysyłają sygnały (sugestie i propozycje ML) dalej do operacji sortowania, gdzie każdy wózek widłowy, który jest sklasyfikowany w grupie sortowania „Tak”, może być zaangażowany w daną operację załadunku.

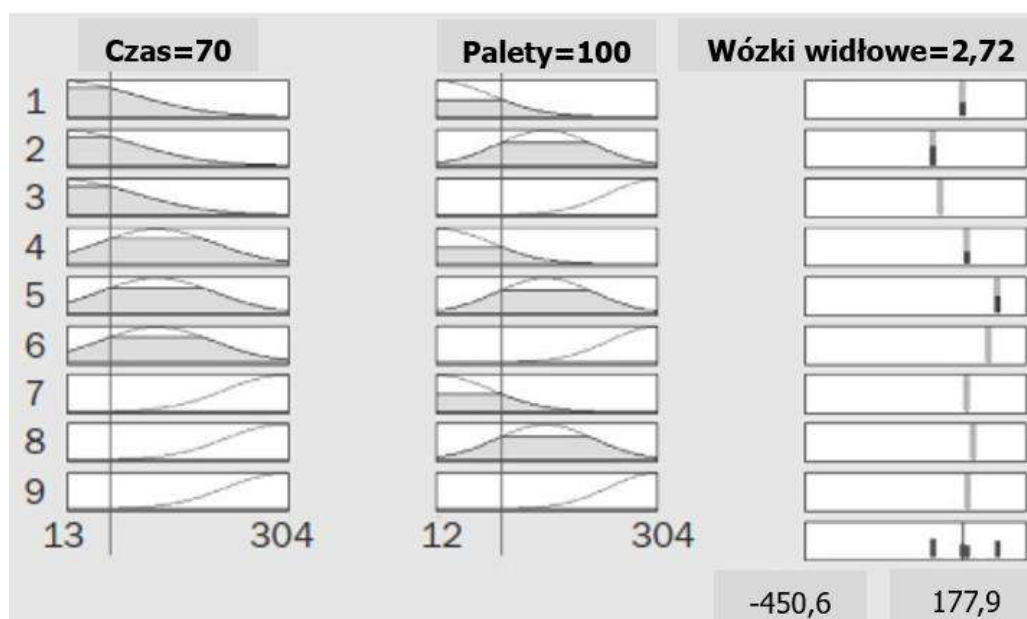
Czynniki wpływające na decyzje kierownika zostały zidentyfikowane w drodze konsultacji. W celu określenia liczby wózków widłowych do rozmieszczenia w strefie załadunku, kluczowymi czynnikami są określony czas załadunku (od 15 do 135 minut) i wielkość ładunku (od 15 do 225 palet). Wybierając wózki widłowe do zaangażowania w operacje, menedżer bierze pod uwagę znaczenie bieżącej aktywności (ocenianej od 1 do 9 zgodnie z polityką firmy), wskaźnik wykorzystania wózka widłowego, jego odległość od doku załadunkowego oraz średni



wskaźnik wykorzystania wszystkich wózków widłowych. Każdy wózek widłowy ma określoną liczbę godzin pracy, zanim konieczny będzie przegląd, a jego użytkowanie jest ograniczone po przekroczeniu tego limitu. Wykorzystanie wózka widłowego (FU) to procent godzin pracy zrealizowanych przez pojedynczy wózek widłowy, podczas gdy średnie wykorzystanie wózka widłowego (AFU) to średnia godzin pracy zrealizowanych przez wszystkie wózki widłowe. Wyższy wskaźnik AFU sugeruje, że większość wózków widłowych będzie wkrótce wymagać przeglądu.

Interfejs użytkownika DSS

W większości sytuacji wejściowych najlepszą wydajność wykazały ANFIS i CART. W związku z tym zostały one wybrane jako silniki danego DSS i jego ES. Interfejs użytkownika modelu ES 1 jest przedstawiony na rysunku 10.5 i umożliwia operatorom szybkie i łatwe podejmowanie decyzji dotyczących liczby wózków widłowych do rozmieszczenia, po prostu przesuwając pionową linię przez domenę zmiennych wejściowych, w oparciu o określony czas załadunku i ilość ładunku.

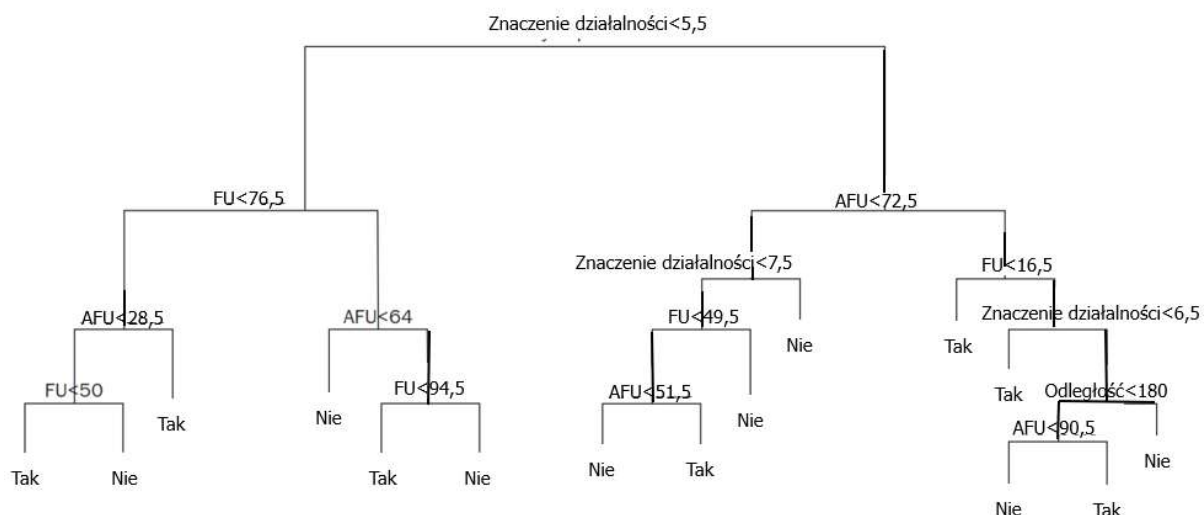


Rysunek 10.5 System wnioskowania rozmytego modelu ES 1

Model ES 2 służy jako narzędzie uzupełniające do Modelu ES 1, usprawniając podejmowanie decyzji poprzez dostarczanie informacji o tym, czy dany wózek widłowy powinien zostać rozmieszczony w strefie załadunku (rysunek 10.6). Biorąc pod uwagę pozycję wózka widłowego (odległość od strefy załadunku), jego bieżącą aktywność (znaczenie aktywności), zrealizowane godziny pracy (FU) i średnie wykorzystanie wszystkich wózków widłowych (AFU),



użytkownicy mogą łatwo określić, czy dany wózek widłowy nadaje się do załadunku, czy też należy wybrać inny. Drzewo decyzyjne CART jest proste w interpretacji, eliminując potrzebę wprowadzania wartości do oprogramowania. Zamiast tego drzewo z rysunku 10.6 można wydrukować i wyświetlić w widocznym miejscu w magazynie w celu szybkiego sprawdzenia.



Rysunek 10.6 Drzewo decyzyjne modelu ES 2 dotyczące zaangażowania wózka widłowego

Menedżerowie mogą codziennie korzystać z prezentowanego systemu DSS, pomagając w osiągnięciu wyższej responsywności łańcucha dostaw na potrzeby klientów i zapewniając wysokie prawdopodobieństwo terminowej dostawy. Zaproponowany AI i ML DSS wykazał skuteczne wyniki w pozyskiwaniu wiedzy eksperta „know-how” i przechwytywaniu jego „logiki wnioskowania”. Korzystając z tej metody, wiedza menedżerów może być wyodrębniona i zastosowana do innych operacji magazynowych. Jest to szczególnie cenne dla praktyków, ponieważ zatrudnianie ekspertów magazynowych jest często kosztowne. Ponadto DSS może również służyć jako narzędzie szkoleniowe dla początkujących menedżerów, pomagając im zdobyć doświadczenie i poprawić umiejętności podejmowania decyzji w miarę upływu czasu. Dlatego systemy AI i ML, które mogą symulować decyzje menedżera, są niezbędnymi narzędziami, oferującymi znaczne oszczędności kosztów i zwiększoną wydajność operacji magazynowych.



BIBLIOGRAFIA DO ROZDZIAŁU 10.

- Athanasopoulou, K., Daneva, G. N., Adamopoulos, P. G., & Scorilas, A. (2022). Artificial intelligence: The milestone in modern biomedical research. *BioMedInformatics*, 2(4), 727-744. <https://doi.org/10.3390/biomedinformatics2040049>
- Chollet, F. (2021). *Deep learning with Python*. Simon and Schuster.
- Druzdel, M. J., & Flynn, R. R. (2002). Decision support systems. In A. Kent (Ed.), *Encyclopedia of library and information science* (2nd ed.). Marcel Dekker, Inc.
- Gibbs, M. (2019, December 17). Table-driven programming and the weather forecasting stone. *Global Nerdy*. <https://www.globalnerdy.com/2019/12/17/table-driven-programming-and-the-weather-forecasting-stone/>
- Jackson, P. (1999). *Introduction to expert systems*. Addison-Wesley.
- Mirčetić, D., Ralević, N., Nikolić, S., Maslarić, M., & Stojanović, Đ. (2016). Expert system models for forecasting forklifts engagement in a warehouse loading operation: A case study. *Promet-Traffic & Transportation*, 28(4), 393-401.
- Mircetic, D., Lalwani, C., Lirn, T., Maslaric, M., & Nikolicic, S. (2014, July). ANFIS expert system for cargo loading as part of decision support system in warehouse. In *19th International Symposium on Logistics (ISL 2014)*.
- Rostami-Tabar, B., & Mircetic, D. (2023). Exploring the association between time series features and forecasting by temporal aggregation using machine learning. *Neurocomputing*, 548, 126376.
- Olson, D. L., & Courtney, J. F. (1992). *Decision support models and expert systems*. Macmillan.
- Rainer, R. K., & Turban, E. (2008). *Introduction to information systems: Supporting and transforming business*. John Wiley & Sons.
- Turban, E. (1998). *Decision support and expert systems* (2nd ed.). Macmillan.
- Turban, E., Aronson, J., & Liang, T.-P. (2005). *Decision support systems and intelligent systems* (7th ed.). Pearson Prentice Hall.