

1

Generalizacja kartograficzna i sztuczna inteligencja

Najważniejsze pojęcia, na których skupiono się w niniejszym opracowaniu, to generalizacja kartograficzna oraz sztuczna inteligencja.

Generalizacja kartograficzna stanowi jeden z najistotniejszych, a zarazem nieodzownych elementów procesu redakcji map. Generalizacja definiowana jest jako celowy i logiczny proces zmniejszania szczegółowości treści mapy przy jednoczesnym uwzględnieniu przeznaczenia mapy, poziomu szczegółowości, możliwości i ograniczeń percepcyjnych odbiorcy oraz środowiska użytkownika (Kraak i inni, 2020). Generalizacja kartograficzna jest złożonym procesem decyzyjnym, w którym powinna zostać wzięta pod uwagę geograficzna specyfika generalizowanych obiektów. Istotą generalizacji kartograficznej jest modelowanie i abstrahowanie informacji geograficznej przy jednoczesnym zachowaniu, a nawet uwypukleniu specyficznych cech generalizowanych obiektów oraz relacji między nimi, odpowiednio do przeznaczenia i skali mapy (Mackaness i inni, 2017).

Termin „sztuczna inteligencja” został zaproponowany w 1956 r. przez Johna McCarthy’ego, matematyka i informatyka. Sztuczna inteligencja jest utożsamiana ze zdolnością danego systemu lub modelu opracowanego w środowisku komputerowym do nauki lub adaptacji na podstawie dostarczonej wiedzy i doświadczenia (Olszewski, 2009). Turing (1950) definiuje sztuczną inteligencję jako zdolność maszyn do komunikowania się z ludźmi (za pomocą urządzeń elektronicznych) bez ujawniania faktu, że nie są to ludzie, tylko maszyny. Autor ten zaproponował test, za pomocą którego można zweryfikować stopień zaawansowania rozwoju sztucznej inteligencji. Jest to eksperyment mający na celu sprawdzenie, czy maszyna potrafi wykazać się inteligencją na poziomie ludzkim, poprzez prowadzenie rozmowy, która nie ujawnia jej mechanicznej natury. W teście sędzia prowadzi rozmowy tekstowe z człowiekiem i maszyną, nie wiedząc, który z uczestników jest maszyną – jego zadaniem jest odgadnięcie tego. Jeśli sędzia nie jest w stanie z całą pewnością rozróżnić maszyny i człowieka, uważa się, że maszyna zdała test. W związku z rozwojem poddziedziny sztucznej inteligencji, jaką są duże modele językowe oraz generatywna forma sztucznej inteligencji, badacze raportują, że jeden z modeli generatywnej sztucznej inteligencji, czyli ChatGPT, przeszedł test Turinga (Biever, 2023). Minsky (2007), jeden

z pionierów AI, zdefiniował ją jako umożliwienie maszynom robienia rzeczy, które wymagają ludzkiej inteligencji. Chociaż opisy sztucznej inteligencji są różne, rdzeń AI jest powszechnie rozumiany jako teorie badawcze, metody, technologie i aplikacje do symulacji i wspierania ludzkiej inteligencji.

W kolejnych podrozdziałach przybliżono rolę selekcji w najnowszym modelu generalizacji. Wskazano oraz opisano fundamentalne źródła wiedzy kartograficznej niezbędnej do pozyskania, formalizacji i implementacji efektywnej generalizacji danych. Przybliżono pojęcie uczenia maszynowego, które stanowi jeden z ważniejszych elementów sztucznej inteligencji, oraz opisano typy uczenia maszynowego.

1.1. Modelowanie procesu generalizacji

W celu przeprowadzenia procesu automatyzacji generalizacji konieczne jest wydzielenie i nazwanie operacji składających się na ten proces, następnie systematyzacja tych operacji, wskazanie zależności między nimi oraz wyróżnienie ich charakterystycznych elementów. Wielu autorów podejmowało się tego zadania, np. opracowując modele generalizacji, które stanowiły odpowiedź na potrzebę opisu generalizacji w taki sposób, aby mogła ona być stosowana w systemach komputerowych. Istotnym elementem składowym modelu generalizacji jest operator generalizacji.

Jeden ze współczesnych modeli generalizacji zakłada hierarchiczny podział procesu na trzy główne etapy: wstępne przetworzenie danych (*pre-processing*), który stanowi etap nadrzędny, oraz dwa etapy podrzędne – grupę operatorów generalizacji ilościowej (*visual quantity*) i grupę operatorów generalizacji jakościowej (*visual quality*) (Stanislawski i inni, 2014). Operator generalizacji określany jest jako ogólny opis modyfikacji obiektu (Regnauld i McMaster, 2007; Roth i inni, 2011). Modyfikacja ta może mieć charakter przestrzenny lub atrybutowy, może zatem dotyczyć geometrii lub cech obiektu. Stanislawski i inni (2014) w ramach poszczególnych etapów wskazują grupy związanych z nimi operatorów generalizacji. Na etapie wstępnego przetworzenia danych są to wzbogacanie danych (*enrichment*) oraz reklasyfikacja danych (*reclassification*). Wzbogacanie danych służy dodaniu istotnych, dodatkowych informacji do bazy danych źródłowych, natomiast reklasyfikacja pozwala na podział klas obiektów (warstw tematycznych zawartych w źródłowej bazie danych) na bardziej szczegółowe podklasy, np. podział dróg na podklasy w zależności od ich kategorii zarządzania. W ramach grupy operatorów związanych z generalizacją ilościową autorzy wyróżniają m.in. selekcję (eliminację, wybór), następnie agregację, operator zmiany geometrii obiektów (*collapse*) oraz upraszczanie. Natomiast w ramach grupy operatorów generalizacji jakościowej wskazują takie operatory, jak przemieszczanie, wygładzenie, przewiększanie.

Operator generalizacji jest implementowany za pomocą algorytmu generalizacji stanowiącego np. określoną formułę matematyczną lub statystyczną. Istnieje wiele klasyfikacji operatorów generalizacji (Regnauld i McMaster, 2007; Roth i inni, 2011). Operatorem, który występuje w każdej z nich, a zarazem kluczowym i zawsze jednym z pierwszych w kolejności zastosowania, jest operator selekcji. Selekcja zakłada

usunięcie jednego lub wielu obiektów, zatem jej zadaniem jest redukcja treści mapy lub zawartości bazy danych odpowiednio do docelowej skali lub zakładanego poziomu szczegółowości (Stanisławski i inni, 2014; Karsznia i Weibel, 2018; Karsznia i Sielicka, 2020).

1.2. Źródła i eksploracja wiedzy kartograficznej

Powstające obecnie bazy danych przestrzennych zawierają wiele informacji dotyczących geometrii obiektów, niestety często brak im atrybutów opisujących własności semantyczne obiektów, a także kontekstowe cechy geometryczne dotyczące relacji z innymi obiektami. Jednocześnie topograficzne bazy danych, stanowiące referencję dla różnych opracowań kartograficznych, powinny być na tyle rozbudowane i bogate w swojej strukturze, aby na ich podstawie można było opracować poprawne mapy i pochodne bazy danych. Efektywne zaimplementowanie procesu generalizacji – co obecnie znaczy: optymalne i automatyczne – wymaga uwzględnienia nie tylko własności geometrycznych obiektów, ale również, co jest oczywiste, ich zróżnicowanych cech semantycznych oraz relacji przestrzennych między obiektami (Mackaness, 2006). Proces polegający na dodawaniu istotnych informacji o obiektach określany jest w literaturze mianem wzbogacania danych (*data enrichment*) (Neun, 2007; Mackaness i Gould, 2014; Mackaness i inni, 2014; Stoter i inni, 2014). Natomiast proces analizowania grup obiektów oraz relacji między nimi nazywany jest rozpoznawaniem struktury (*pattern recognition*) (Brassel i Weibel, 1988; Steiniger i Weibel, 2007; Karsznia i Weibel, 2018). Z tego względu wzbogacanie danych jest częstą praktyką zarówno w badaniach naukowych, jak i w pracach redakcyjnych.

Do osiągnięcia optymalizacji procesu automatycznej generalizacji kartograficznej niezbędna jest eksploracja wiedzy kartograficznej oraz wykorzystanie jej w programach wspomagających ten proces. Wiedza ta, aby mogła zostać z sukcesem zastosowana w podejściu automatycznym, w środowisku komputerowym, powinna zostać w odpowiedni sposób uściślona (sformalizowana w postaci reguł lub warunków). Opisana i sformalizowana wiedza kartograficzna wraz z jej implementacją w postaci odpowiednich narzędzi, czyli m.in. operatorów, algorytmów oraz parametrów generalizacji, określana jest mianem kartograficznej bazy wiedzy lub, krócej, bazy wiedzy (Karsznia, 2010; 2015).

Można wyróżnić trzy typy wiedzy kartograficznej i odpowiednio ich źródła. Autorzy wyszczególniają wiedzę geometryczną, strukturalną (semantyczną) oraz proceduralną (Armstrong, 1991; Weibel, 1995).

Wiedza geometryczna pozwala na opisanie geometrii generalizowanych obiektów, np. współrzędnych, kształtu. Jest łatwa do pozyskania oraz formalizacji. Wiedza strukturalna (semantyczna) dostarcza informacji o ogólnej strukturze i znaczeniu obiektów, np. ekonomicznym, kulturowym (Kulik i inni, 2005). Jej pozyskanie i formalizacja są znacznie większym wyzwaniem niż w przypadku wiedzy geometrycznej, zwykle zawartej bezpośrednio w bazach danych. Pozyskanie wiedzy semantycznej wymaga często wspomnianego wcześniej wzbogacania danych źródłowych

(*data enrichment*). Wiedza proceduralna dotyczy informacji o tym, jakie narzędzia i w jakiej kolejności powinny zostać zastosowane w generalizacji (Stoter, 2009; 2014). Jest to bardzo ważny typ wiedzy, którego pozyskanie jest największym wyzwaniem.

Źródłami wiedzy semantycznej oraz proceduralnej są instrukcje redakcji map, doświadczenie i wiedza kartografa oraz poprawnie opracowane przez kartografów mapy. Eksploracja wiedzy z wykorzystaniem produktu końcowego, czyli map, określana jest w literaturze mianem inżynierii odwrotnej (*reverse engineering*) (Leitner i Buttenfield, 1995). Natomiast proces pozyskania wiedzy kartograficznej nazywany jest eksploracją wiedzy głębokiej (*deep knowledge*) (Müller i Mouwes, 1990).

W procesie eksploracji oraz formalizacji wiedzy kartograficznej – w szczególności pozyskiwania zasad generalizacji, wskazywania najważniejszych atrybutów (zmiennych) charakteryzujących generalizowane obiekty (wiedza strukturalna), eksploracji i uczytelniania procesu decyzyjnego, a także ustalania kolejności stosowania narzędzi generalizacji (wiedza proceduralna) – kluczowe i bardzo przydatne jest wykorzystanie sztucznej inteligencji oraz uczenia maszynowego.

1.3. Uczenie maszynowe

Uczenie maszynowe stanowi jedną z poddziedzin sztucznej inteligencji. Związane jest z konstrukcją programów komputerowych oraz algorytmów i modeli, które mają zdolność uczenia się na podstawie dostarczonych danych oraz doskonalenia z wykorzystaniem zdobywanego doświadczenia (Mitchell, 1997). Na podstawie dostarczonych do programu danych oraz przykładów algorytmy uczenia maszynowego mogą dokonać eksploracji wiedzy (nowych informacji) oraz wnioskować na podstawie zależności, relacji oraz prawidłowości wykrytych w danych. Uczenie maszynowe jest związane z programowaniem komputerów w celu optymalizacji ich wydajności przy wykorzystaniu przykładowych danych lub doświadczenia (Karsznia, 2023).

Uczenie maszynowe znajduje zastosowanie w rozwiązywaniu problemów związanych z rozpoznawaniem obrazów i rozpoznawaniem mowy, jest stosowane w robotyce i innych dziedzinach. Dzięki rozwojowi zaawansowanych technologii komputerowych możemy obecnie przechowywać oraz analizować duże zbiory danych (*big data*), do których dostęp mamy zapewniony w internecie, nawet z odległych lokalizacji. Zastosowanie uczenia maszynowego w przetwarzaniu dużych zbiorów danych określa się w literaturze mianem eksploracji danych (*data mining*). Eksploracja danych polega na przetwarzaniu i analizie zbiorów danych w taki sposób, aby na ich podstawie można było skonstruować model, który jest wiarygodny, poprawny oraz reprezentatywny (Alpaydin, 2010). Systemy sztucznej inteligencji i modele uczenia maszynowego są wyposażone w możliwość uczenia i doskonalenia oraz eksploracji nowej wiedzy. Wykorzystanie uczenia maszynowego zakłada podział zbioru danych na dwie grupy: zbiór danych treningowych oraz zbiór danych testowych. Dane trenin-gowe służą do trenowania modeli uczenia maszynowego (tzw. dane uczące) i są wyposażone w przykłady poprawnych rozwiązań (odpowiedzi), natomiast dane testowe (tzw. dane walidacyjne) służą do weryfikacji efektywności uczenia maszynowego.

1.4. Typy uczenia maszynowego

Można wyróżnić cztery główne typy uczenia maszynowego: uczenie nadzorowane (*supervised learning*), uczenie nienadzorowane (*unsupervised learning*), uczenie wzmocnione (*reinforcement learning*) oraz uczenie ewolucyjne (*evolutionary learning*) (Marshland, 2015).

W przypadku uczenia nadzorowanego, określanego również mianem predykcyjnego, zakładamy, że do modelu uczenia maszynowego dostarczone są dane treninowe zawierające przykłady poprawnych rozwiązań. Na podstawie danych treningowych wraz z poprawnymi przykładami model uczenia maszynowego zdobywa wiedzę (jest uczony), a następnie ją wykorzystuje, aby rozwiązywać zadania na podstawie nowych danych. Ten typ uczenia określany jest również mianem uczenia na podstawie przykładów (*learning from examples*) (Murphy, 2012; Marshland, 2015).

W uczeniu nienadzorowanym, nazywanym również uczeniem opisowym (*descriptive learning*), przykłady poprawnych rozwiązań nie są dostarczane do modelu uczenia maszynowego, a działanie modelu oraz podejmowanie decyzji odbywa się na podstawie analizy danych i wykrywania prawidłowości oraz zależności między danymi. Zadaniem tego modelu uczenia maszynowego jest poszukiwanie podobieństwa między danymi wejściowymi, tak aby obiekty scharakteryzowane przez dane wejściowe mające wspólne cechy były rozpatrywane łącznie. W przypadku gdy celem modelu jest eksploracja grup obiektów o podobnych właściwościach, mamy do czynienia z klasteryzacją (*clustering*), natomiast gdy celem jest określenie wzajemnego rozmieszczenia obiektów w przestrzeni, możemy mówić o szacowaniu gęstości (*density estimation*) (Bishop, 2006).

Uczenie wzmocnione, nazywane także uczeniem przez wzmacnianie (*reinforcement learning*), może być rozpatrywane jako typ uczenia pośredniego między uczeniem nadzorowanym i nienadzorowanym (Sutton i Barto, 1998). W przypadku tego typu uczenia model otrzymuje informację o rozwiązaniach i przykładach niepoprawnych, jednak nie jest wyposażony w informację dotyczącą poprawnych rozwiązań. Działanie modelu w tym przypadku polega na eksploracji danych metodą prób i błędów w celu odnalezienia i wskazania poprawnych rozwiązań. Uczenie przez wzmacnianie stanowi więc rodzaj kompromisu między eksploracją, w której w modelu sprawdzane są nowe rodzaje działań, aby ocenić ich skuteczność, a eksploatacją, kiedy w modelu wykorzystywane są działania, o których wiadomo, że są skuteczne (Bishop, 2006).

Ostatni typ uczenia to uczenie ewolucyjne (*evolutionary learning*), które nawiązuje do ewolucji w znaczeniu biologicznym, gdzie organizmy adaptują się i rozwijają, aby zwiększyć szanse przetrwania. Uczenie ewolucyjne jest związane z implementacją modeli, w których zastosowano wybrane aspekty naśladujące naturalną ewolucję w celu usprawnienia i optymalizacji działania programów komputerowych (Back, 1996; Bishop, 2006).

Najbardziej popularną i najczęściej wykorzystywaną w badaniach techniką jest uczenie nadzorowane. W niniejszych badaniach wykorzystano ten typ uczenia maszynowego.