



Stuart
Russell

Peter
Norvig

Sztuczna inteligencja

Nowe spojrzenie

Wydanie IV. Tom 1

Tytuł oryginału: Artificial Intelligence: A Modern Approach, 4th Edition

Tłumaczenie: Andrzej Grażyński

ISBN: 978-83-283-7608-3

Authorized translation from the English language edition, entitled ARTIFICIAL INTELLIGENCE: A MODERN APPROACH, 4th Edition by STUART RUSSELL; PETER NORVIG, published by Pearson Education, Inc, publishing as Pearson, Copyright © 2021, 2010, 2003 by Pearson Education, Inc. or its affiliates, 221 River Street, Hoboken, NJ 07030.

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from Pearson Education, Inc.

Polish language edition published by Helion S.A., Copyright © 2022.

PEARSON, ALWAYS LEARNING is an exclusive trademark owned by Pearson Education, Inc. or its affiliates in the U.S. and/or other countries.

Wszelkie prawa zastrzeżone. Nieautoryzowane rozpowszechnianie całości lub fragmentu niniejszej publikacji w jakiegokolwiek postaci jest zabronione. Wykonywanie kopii metodą kserograficzną, fotograficzną, a także kopiowanie książki na nośniku filmowym, magnetycznym lub innym powoduje naruszenie praw autorskich niniejszej publikacji.

Wszystkie znaki występujące w tekście są zastrzeżonymi znakami firmowymi bądź towarowymi ich właścicieli.

Autor oraz wydawca dołożyli wszelkich starań, by zawarte w tej książce informacje były kompletne i rzetelne. Nie biorą jednak żadnej odpowiedzialności ani za ich wykorzystanie, ani za związane z tym ewentualne naruszenie praw patentowych lub autorskich. Autor oraz wydawca nie ponoszą również żadnej odpowiedzialności za ewentualne szkody wynikłe z wykorzystania informacji zawartych w książce.

Helion S.A.

ul. Kościuszki 1c, 44-100 Gliwice

tel. 32 231 22 19, 32 230 98 63

e-mail: helion@helion.pl

WWW: <https://helion.pl> (księgarnia internetowa, katalog książek)

Drogi Czytelniku!

Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres

<https://helion.pl/user/opinie/szt41>

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

- Kup książkę
- Poleć książkę
- Oceń książkę

- Księgarnia internetowa
- Lubię to! » Nasza społeczność

SPIS TREŚCI

Zanim przemówią autorzy...	11
Przedmowa	13
O autorach	16

I Sztuczna inteligencja

1	Wstęp	17
1.1.	Czym jest sztuczna inteligencja?	17
1.2.	Podstawy sztucznej inteligencji	22
1.3.	Historia sztucznej inteligencji	35
1.4.	Stan obecny	46
1.5.	Spodziewane korzyści i ryzyko	50
	Podsumowanie	53
	Bibliografia i uwagi historyczne	54
2	Inteligentni agenci	55
2.1.	Agenci i ich środowiska	55
2.2.	Właściwe zachowanie — koncepcja racjonalności	58
2.3.	Natura środowiska	61
2.4.	Struktura agenta	67
	Podsumowanie	80
	Bibliografia i uwagi historyczne	81

II Rozwiązywanie problemów

3	Rozwiązywanie problemów za pomocą wyszukiwania	83
3.1.	Agent rozwiązujący problem	83
3.2.	Przykładowe problemy	87
3.3.	Algorytmy wyszukiwania	92
3.4.	Strategie wyszukiwania niedoinformowanego	98
3.5.	Strategie wyszukiwania poinformowanego (heurystycznego)	108
3.6.	Funkcje heurystyczne	123
	Podsumowanie	132
	Bibliografia i uwagi historyczne	133

4	Wyszukiwanie w złożonych środowiskach	137
4.1.	Wyszukiwanie lokalne i problemy optymalizacyjne	137
4.2.	Wyszukiwanie lokalne w przestrzeniach ciągłych	147
4.3.	Wyszukiwanie z niedeterministycznymi akcjami	150
4.4.	Wyszukiwanie w środowiskach częściowo obserwowalnych	155
4.5.	Wyszukiwanie online i nieznane środowiska	164
	Podsumowanie	171
	Bibliografia i uwagi historyczne	172
5	Wyszukiwanie antagonistyczne i gry	176
5.1.	Teoria gier	176
5.2.	Optymalne decyzje w grach	179
5.3.	Heurystyczne wyszukiwanie alfa-beta	187
5.4.	Wyszukiwanie Monte Carlo	192
5.5.	Gry stochastyczne	196
5.6.	Gry z częściową obserwowalnością	200
5.7.	Ograniczenia algorytmów wyszukiwania w grach	204
	Podsumowanie	206
	Bibliografia i uwagi historyczne	207
6	Problemy spełniania ograniczeń	212
6.1.	Definiowanie problemów spełniania ograniczeń	212
6.2.	Propagacja ograniczeń — wnioskowanie w CPS	218
6.3.	Wyszukiwanie z nawrotami w CPS	224
6.4.	Wyszukiwanie lokalne na usługach CSP	230
6.5.	Struktura problemów CSP	232
	Podsumowanie	238
	Bibliografia i uwagi historyczne	238
 III Wiedza, wnioskowanie i planowanie		
7	Logiczni agenci	242
7.1.	Agent bazujący na wiedzy	243
7.2.	Świat Wumpusa	244
7.3.	Podstawy logiki	248
7.4.	Rachunek zdań — bardzo prosta logika	251
7.5.	Dowodzenie twierdzeń w rachunku zdań	256
7.6.	Efektywne sprawdzanie modeli w rachunku zdań	268
7.7.	Agent na gruncie rachunku zdań	273
	Podsumowanie	283
	Bibliografia i uwagi historyczne	284

8	Logika pierwszego rzędu	287
8.1.	Ponownie o reprezentacji	287
8.2.	Składnia i semantyka logiki pierwszego rzędu	292
8.3.	Wykorzystywanie logiki pierwszego rzędu	302
8.4.	Inżynieria wiedzy w logice pierwszego rzędu	309
	Podsumowanie	315
	Bibliografia i uwagi historyczne	316
9	Wnioskowanie w logice pierwszego rzędu	318
9.1.	Wnioskowanie w rachunku zdań a wnioskowanie w logice pierwszego rzędu	318
9.2.	Unifikacja a wnioskowanie w logice pierwszego rzędu	321
9.3.	Łańcuchowanie progresywne	325
9.4.	Łańcuchowanie regresywne	332
9.5.	Rezolucja	339
	Podsumowanie	351
	Bibliografia i uwagi historyczne	352
10	Reprezentacja wiedzy	356
10.1.	Inżynieria ontologii	356
10.2.	Kategorie i obiekty	359
10.3.	Zdarzenia	366
10.4.	Obiekty mentalne i logika modalna	370
10.5.	Systemy wnioskowania dla kategorii	373
10.6.	Wnioskowanie na podstawie domniemań	378
	Podsumowanie	382
	Bibliografia i uwagi historyczne	383
11	Automatyczne planowanie	388
11.1.	Klasyczne planowanie — co to jest?	388
11.2.	Algorytmy klasycznego planowania	393
11.3.	Heurystyki w planowaniu	398
11.4.	Planowanie hierarchiczne	402
11.5.	Planowanie i działanie w domenach niedeterministycznych	411
11.6.	Czas, harmonogramy i zasoby	421
11.7.	Analiza podejść planistycznych	425
	Podsumowanie	426
	Bibliografia i uwagi historyczne	427

IV Wnioskowanie w warunkach niepewności

12	Kwantyfikowanie niepewności	432
12.1.	Działając w warunkach niepewności	432
12.2.	Notacja probabilistyczna	436
12.3.	Wnioskowanie z pełnych wspólnych rozkładów	444
12.4.	Niezależność	446
12.5.	Reguła Bayesa i jej wykorzystywanie	448
12.6.	Naiwne modele bayesowskie	451
12.7.	Odwiedzamy świat Wumpusa	453
	Podsumowanie	457
	Bibliografia i uwagi historyczne	458
13	Wnioskowanie probabilistyczne	461
13.1.	Reprezentowanie wiedzy w niepewnej domenie	461
13.2.	Semantyka sieci bayesowskich	464
13.3.	Ścisłe wnioskowanie w sieciach bayesowskich	477
13.4.	Aproksymowane wnioskowanie w sieciach bayesowskich	486
13.5.	Sieci przyczynowe	501
	Podsumowanie	506
	Bibliografia i uwagi historyczne	507
14	Probabilistyczne wnioskowanie w czasie	513
14.1.	Czas a niepewność	514
14.2.	Wnioskowanie w modelach temporalnych	518
14.3.	Ukryte modele Markowa	526
14.4.	Filtrowanie Kalmana	532
14.5.	Dynamiczne sieci bayesowskie (DBN)	540
	Podsumowanie	552
	Bibliografia i uwagi historyczne	552
15	Programowanie probabilistyczne	555
15.1.	Relacyjne modele probabilistyczne	556
15.2.	Modele probabilistyczne otwartego wszechświata	562
15.3.	Śledzenie skomplikowanego świata	570
15.4.	Programy jako modele probabilistyczne	575
	Podsumowanie	580
	Bibliografia i uwagi historyczne	580
16	Podejmowanie prostych decyzji	585
16.1.	Przekonania i pragnienia w warunkach niepewności	586
16.2.	Podstawy teorii użyteczności	587
16.3.	Funkcje użyteczności	590

16.4.	Wieloatrybutowe funkcje użyteczności	598
16.5.	Sieci decyzyjne	603
16.6.	Wartość informacji	606
16.7.	Nieznane preferencje	613
	Podsumowanie	616
	Bibliografia i uwagi historyczne	617
17	Podejmowanie złożonych decyzji	621
17.1.	Sekwencyjne problemy decyzyjne	621
17.2.	Algorytmy dla problemów MDP	633
17.3.	Problem bandyty i jego warianty	641
17.4.	Częściowo obserwowalne problemy MDP (POMDP)	649
17.5.	Algorytmy rozwiązywania problemów POMDP	651
	Podsumowanie	657
	Bibliografia i uwagi historyczne	657
18	Podejmowanie decyzji w środowisku wieloagentowym	661
18.1.	Właściwości środowisk wieloagentowych	661
18.2.	Teoria gier niekooperatywnych	667
18.3.	Teoria gier kooperatywnych	689
18.4.	Kolektywne podejmowanie decyzji	696
	Podsumowanie	708
	Bibliografia i uwagi historyczne	709
 Dodatki		
A	Kompendium matematyczne	715
A.1.	Analiza złożoności i notacja „dużego O”	715
A.2.	Wektory, macierze i algebra liniowa	718
A.3.	Rozkłady prawdopodobieństwa	720
A.4.	Wybrane operacje na zbiorach	723
	Bibliografia i uwagi historyczne	725
B	Konwencje notacyjne i pseudokod	726
B.1.	Definiowanie składni za pomocą notacji BNF	726
B.2.	Algorytmy w formie pseudokodu	727
B.3.	Uzupełniające materiały online	728
	Skorowidz	729

ROZDZIAŁ 3

ROZWIĄZYWANIE PROBLEMÓW ZA POMOCĄ WYSZUKIWANIA

W tym rozdziale wyjaśniamy, jak agent może poszukiwać ciągu akcji, które mają doprowadzić go do założonego celu.

Gdy prawidłowa akcja, którą powinien podjąć agent, nie jest dla niego od razu oczywista, zmuszony zostaje on do *planowania z wyprzedzeniem*, czyli rozważenia ciągu akcji formujących ścieżkę prowadzącą do stanu, który oznacza spełnienie celu. Agent działającego w taki sposób nazywamy **agentem rozwiązującym problemy** (ang. *problem-solving agent*), a kryjące się za nim procesy obliczeniowe określamy mianem **wyszukiwania** (ang. *search*).

Opisywany w tym rozdziale agent rozwiązujący problemy zrealizowany jest w reprezentacji **atomowej** (patrz sekcja 2.4.7) — stany postrzeganego przezeń środowiska są niepodzielne, nie objawiają jakiegokolwiek struktury dla algorytmów rozwiązywania problemów. Agentów wykorzystujących reprezentację **czynnиковą** lub **strukturalną** — zwanych **agentami planowania** — opisujemy w rozdziałach 7. i 11.

Omówimy kilka algorytmów wyszukiwania, ograniczając się do najprostszego typu środowiska — epizodycznego, jednoagentowego, w pełni obserwowalnego, deterministycznego, statycznego, dyskretnego i rozpoznanego. Odróżniać będziemy algorytmy **poinformowane** (ang. *informed*) od **niedoinformowanych** (ang. *uninformed*) — agent używający tych pierwszych ma możliwość oszacowania, jak daleko jest mu jeszcze do osiągnięcia celu, dla agenta używającego tych drugich takie szacowanie jest niedostępne. W rozdziale 4. uogólnimy naszą dyskusję na inne typy środowisk jednoagentowych, w rozdziale 5. zajmiemy się środowiskami wieloagentowymi.

W tym rozdziale wykorzystujemy intensywnie notację „dużego O ”, odzwierciedlającą asymptotyczną złożoność algorytmów — na przykład $O(n^2)$. Czytelników chcących uporządkować swą wiedzę na jej temat odsyłamy do Dodatku A.

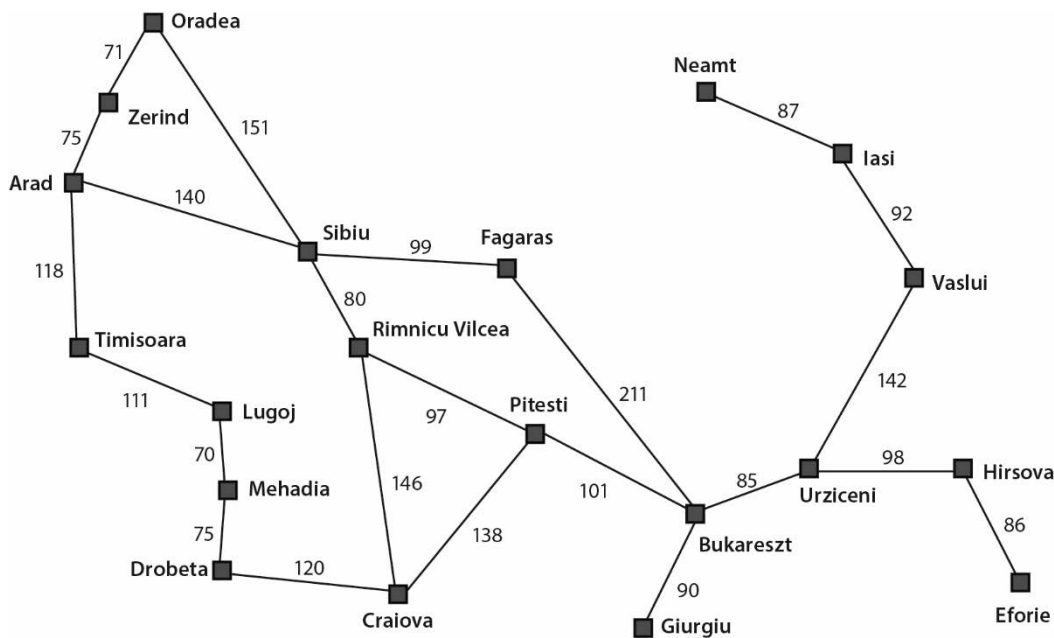
3.1. Agent rozwiązujący problem

Wyobraźmy sobie agenta, który spędza wakacje w Rumunii: zamierza zwiedzać zabytki, doskonalić znajomość języka, cieszyć się nocnym życiem (bez szkody dla samopoczucia następnego dnia) i tak dalej — wybór możliwości jest dość bogaty. Agent znajduje się obecnie w mieście Arad i ma bilet na lot z Bukaresztu następnego dnia. Według drogowskazów z Arad wychodzą trzy drogi: jedna w kierunku Sibiu, druga do Timishoary, trzecia do Zerind. Żadna z nich nie prowadzi bezpośrednio do Bukaresztu, więc podróżny, który nie zna geografii Rumunii, nie ma pomysłu na to, którą wybrać¹.

¹ Czytelników z Rumunii przepraszamy w tym miejscu za to, że omawiany przykład ma dla nich jakby mniejszy walor dydaktyczny.

Jeśli agent nie dysponuje jakimiś dodatkowymi informacjami — czyli jego środowisko jest dla niego nieznane — nie pozostaje mu nic lepszego, jak losowy wybór jednej z dróg; tę przykrą sytuację rozpatrujemy w rozdziale 4. Na użytek tego rozdziału zakładamy jednak, że nasi agenci zawsze dysponują informacjami o otaczającym świecie, chociażby takimi, jak mapa przedstawiona na rysunku 3.1. Nasz agent może tym samym rozwiązać swój problem, w czterech następujących etapach.

- **Zdefiniowanie celu.** Agent stawia sobie za główny cel dotarcie do Bukaresztu. Główny cel determinuje zachowanie agenta poprzez ograniczenie celów cząstkowych i tym samym akcji stojących do dyspozycji agenta.
- **Sformułowanie problemu.** Agent sporządza opis stanów i działań niezbędnych do osiągnięcia celu — czyli tworzy model części świata powiązanej z celem. Akcje w tym modelu ograniczają się do podróży z danego miasta do miast sąsiednich, zatem jedynym elementem stanu otaczającego świata, zmieniającym się wskutek wykonania akcji, jest bieżące położenie agenta.
- **Wyszukiwanie.** Zanim agent wykona jakąkolwiek akcję w rzeczywistym świecie, symuluje różne sekwencje akcji wynikające z jego modelu, wybierając ostatecznie tę sekwencję, która doprowadzi go do celu — sekwencję tę nazywamy **rozwiązaniem**. Oczywiście może się tak zdarzyć, że wśród sekwencji analizowanych przez agenta nie będzie żadnej, która mogłaby doprowadzić go do celu, w naszym przypadku jednak może to być (na przykład) sekwencja prowadząca wzdłuż trasy Arad – Sibiu – Fagars – Bukareszt.
- **Wykonanie.** Agent wykonuje teraz ciąg akcji, z których każda obejmuje jeden etap wspomnianej trasy.



RYSunek 3.1. Uproszczona mapa części Rumunii, z zaznaczonymi odległościami (w milach) między miastami

Zwróćmy uwagę na ważny fakt: agent, po wytyczeniu trasy, może ją przemierzać „z zamkniętymi oczami”, czyli ignorować swe spostrzeżenia w tym czasie. Formalnie oznacza to, że *agent w czasie wykonywania swych akcji może ignorować zmiany zachodzące w środowisku, oczywiście jeśli zbudowany przez niego model jest poprawny; jest to cecha charakterystyczna dla środowiska rozpoznanego, deterministycznego i w pełni obserwowalnego, jeśli rozwiązaniem problemu jest ustalona sekwencja akcji*. Specjaliści od teorii sterowania nazywają taki stan systemem **otwartej pętli** (ang. *open loop*) — w czasie wykonywania sekwencji akcji pętla sprzężenia zwrotnego między agentem a środowiskiem zostaje przerwana („otwarta”). Jeśli istnieje prawdopodobieństwo, że model może być błędnie określony, bądź środowisko nie jest deterministyczne, agent byłby bezpieczniejszy, jeśli nie przerywałby obserwacji środowiska (co teoretycy sterowania nazywają **zamkniętą pętlą** — ang. *closed loop*), o czym piszemy w podrozdziale 4.4.

W środowiskach częściowo obserwowalnych lub niedeterministycznych rozwiązanie miałyby prawdopodobnie postać rozgałęzioną, zakładającą dynamiczny wybór kolejnego miasta na trasie w zależności od wyniku obserwacji środowiska. Przykładowo, agent znajdujący się w Arad może planować jazdę do Sibiu, ale bierze pod uwagę fakt, że może natrafić na zamknięty odcinek drogi (znak „droga zamknięta” — rum. *drum inchis*) bądź też (słabo orientując się na mapie) omyłkowo dojechać do Zerind, dlatego posiada plan awaryjny uwzględniający takie sytuacje.

3.1.1. Problemy wyszukiwania i ich rozwiązania

Problem wyszukiwania można formalnie zdefiniować jako kombinację następujących elementów:

- Zbiór możliwych **stanów**, w których może znajdować się środowisko; zbiór ten nazywamy **przestrzenią stanów**.
- **Stan początkowy**, czyli stan, w jakim znajduje się środowisko w momencie, gdy zaczynamy analizować działanie agenta; dla naszego agenta-podróżnika jest to *Arad*.
- **Stan docelowy** lub zbiór stanów docelowych. Dla naszego agenta-podróżnika jest to pojedynczy stan (*Bukareszt*), ogólnie jednak stany docelowe mogą tworzyć zbiór — być może niewielki, być może nawet nieskończony. Dla agenta-odkurzacza z rozdziału 2. stanem docelowym może być brak zabrudzenia na którymkolwiek kwadracie, niezależnie od innych aspektów stanu środowiska. Będziemy brać pod uwagę wszystkie trzy wymienione rodzaje stanu docelowego (pojedynczy stan, skończony zbiór stanów, nieskończenie wiele stanów) reprezentując je w postaci metody `IS_GOAL`. Dla prostoty będziemy w dalszej części rozdziału pisać po prostu „stan docelowy”, lecz to określenie należy każdorazowo rozumieć jako „dowolny spośród możliwych stanów docelowych”.
- **Akcje** dostępne dla agenta. Dla danego stanu s metoda `ACTIONS(s)` zwraca skończony² zbiór akcji, możliwych do wykonania w stanie s . O każdej z tych akcji mówimy, że jest **stosowna** do stanu s . Na przykład

`ACTIONS(Arad) = {ToSibiu, ToTimisoara, ToZerind}`

- **Model przejścia** określa skutki wykonywania akcji — wynikiem wywołania metody `RESULT(s, a)` jest stan będący rezultatem wykonania akcji a w stanie s , na przykład

`RESULT(Arad, ToZerind) = Zerind`

- **Funkcja kosztu akcji**, oznaczana przez `ACTION-CONST(s, a, s')` lub $c(s, a, s')$ oznacza wyrażony liczbowo koszt osiągnięcia stanu s' jako skutku wykonania akcji a w stanie s . Agent rozwiązujący problemy powinien wykorzystywać tę funkcję jako odzwierciedlającą jego własną miarę wydajności, w przypadku naszego agenta-podróżnika kosztem przejścia ze stanu s do stanu s' może być odległość między miastami odpowiadającymi tym stanom, według mapy z rysunku 3.1, albo czas przejazdu między tymi miastami.

² Do rozwiązywania problemów, w których stanom odpowiadają nieskończone zbiory akcji, konieczne jest zastosowanie technik wykraczających poza zakres tego rozdziału.

Sekwencja działań tworzy **ścieżkę**, a **rozwiązanie** problemu to ścieżka prowadząca od stanu początkowego do stanu docelowego. Zakładamy, że koszty akcji są *addytywne*; czyli że całkowity koszt ścieżki jest sumą kosztów akcji tworzących tę ścieżkę. **Rozwiązaniem optymalnym** wśród wszystkich możliwych rozwiązań jest to, któremu odpowiada ścieżka o najmniejszym koszcie. W tym rozdziale, dla uniknięcia zbytnich komplikacji, zakładamy, że wszystkie koszty akcji są dodatnie³.

Przestrzeń stanów można przedstawić w postaci **grafu**, którego wierzchołki reprezentują stany, a skierowane krawędzie między nimi są odpowiednikami akcji.

Mapa przedstawiona na rysunku 3.1 jest grafem, którego krawędzie nie są skierowane. Krawędź nieskierowaną można traktować jako parę krawędzi skierowanych przeciwnie względem siebie, zatem każdą drogę na wspomnianej mapie można traktować jako parę akcji związanych z przemieszczeniami między miastami określonej pary.

3.1.2. Formułowanie problemów

Nasze sformułowanie problemu dotarcia z Arad do Bukaresztu to w istocie **model** — abstrakcyjny opis matematyczny, jedynie przybliżający rzeczywistość. W porównaniu z licznymi aspektami podróży po kraju — towarzyszącymi podróży, dostępnością ulubionych stacji radiowych, krajobrazami za oknem, bliskością patroli policyjnych, odległością między parkingami, stanem dróg, warunkami pogodowymi — model ogranicza się wyłącznie do ścieżki prowadzącej od stanu *Arad* do stanu *Bukareszt*, reszta nie ma znaczenia dla rozwiązania problemu.

Takie ignorowanie nieistotnych szczegółów nosi nazwę **abstrakcji**. Dobrze zbudowany model charakteryzuje się odpowiednim poziomem abstrakcji: gdybyśmy akcje agenta-podróżnika opisywali na poziomie szczegółowości typu „przesuń prawą stopę o centymetr do przodu” czy „obróć kierownicę o jeden stopień w kierunku przeciwnym do ruchu wskazówek zegara”, agent nie tylko nie zdążyłby na samolot w Bukareszcie, ale mnóstwo czasu zajęłoby mu w ogóle wyjechanie z parkingu.

Żeby być bardziej precyzyjnymi w kwestii **poziomu abstrakcji**: zwróć uwagę, że wybrane przez nas stany i akcje uosabiają duże zbiory szczegółowych stanów świata rzeczywistego i szczegółowych sekwencji działań. Rozważmy nieco inną trasę naszego agenta — Arad – Sibiu – Rimnicu Vilcea – Pitesti – Bukareszt — i nieco bardziej szczegółowo określoną sekwencję akcji: na odcinku od Sibiu do Rimnicu Vilcea agent słuca radia, potem jednak siła sygnału jego ulubionej stacji szybko zanika, więc agent wyłącza radio na resztę podróży.

Abstrakcja jest *poprawna*, jeśli możemy przy jej użyciu uzyskać dowolne abstrakcyjne rozwiązanie odpowiadające bardziej szczegółowemu rozwiązaniu w bardziej szczegółowym świecie, pod jednym wszakże warunkiem — każdej abstrakcyjnej akcji (na przykład *ToSibiu*) odpowiada sekwencja bardziej szczegółowych akcji, powodująca identyczną zmianę między tymi samymi stanami (na przykład *WyjeżdżamZParkinguWArad*, *UzupełniamPaliwo*, *JadęWKierunkuAutostrady*, ..., *ZjeżdżamZAutostrady*, *WjeżdżamNaParkingWSibiu*)⁴. Abstrakcja jest *użyteczna*, jeśli każda z akcji

³ Jeśli dopuścimy ujemną wartość kosztu akcji, to możliwe będzie powstawanie cykli o zerowym lub ujemnym koszcie. Optymalnym rozwiązaniem problemu zawierającego cykl o ujemnym koszcie netto jest przebycie tego cyklu nieskończoną liczbę razy. Algoritmy Bellmana-Forda i Floyda-Warshalla (nie omawiamy ich tutaj) radzą sobie z ujemnymi kosztami (wagami krawędzi), o ile w grafie nie występują (właśnie) cykle o ujemnych kosztach netto. Jeśli chodzi o akcje o *zerowym* koszcie, to łatwo je uwzględnić pod warunkiem, że liczba takich akcji następujących bezpośrednio po sobie jest a priori ograniczona. Możemy na przykład dysponować robotem, którego każdy ruch postępowy wiąże się z określonym kosztem, ale koszt obrotu o kąt prosty jest zerowy: algorytm omawiane w tym rozdziale radzą sobie z takim przypadkiem, o ile dozwolone są co najwyżej trzy takie obroty „z rzędu”, czyli nieprzedzielone ruchem postępowym. Kolejnym trudnym przypadkiem są problemy zawierające ścieżki o skończonym koszcie, jednak złożone z *nieskończonej* liczby akcji o małych kosztach — zjawisko to jest istotą paradoksu Zenona, polegającego na tym, że analizujemy jednostajny ruch żółwia, dzieląc odmierzany czas na odcinki, z których każdy równy jest połowie poprzedniego. Aby żółw mógł dotrzeć do celu, potrzebna jest nieskończona liczba takich odcinków, czyli nieskończony czas, ergo żółw nigdy do celu nie dotrze (od dawna wiadomo jednak, gdzie tkwi błąd w tym rozumowaniu). By uniknąć takich subtelnych problemów wymagamy, by koszt każdej akcji był nie mniejszy niż pewien ustalony limit $\epsilon > 0$.

⁴ Patrz podrozdział 11.4.

prowadzących do abstrakcyjnego rozwiązania jest prostsza niż odpowiadająca jej sekwencja akcji w świecie rzeczywistym; w naszym przykładzie akcja *ToSibiu*, oznaczająca „przejeżdż z Arad do Sibiu” bez sprecyzowania jakichkolwiek dalszych szczegółów, jest jednoznacznie rozumiana dla kierowcy o przeciętnych umiejętnościach. Opracowanie dobrej abstrakcji sprowadza się zatem do wyeliminowania jak największej liczby nieistotnych szczegółów, oczywiście przy zachowaniu jej poprawności i zapewnieniu, że każda abstrakcyjnie sformułowana akcja da się prosto przełożyć na konkretne działania w rzeczywistości. Gdyby nie umiejętność budowania użytecznych abstrakcji przez projektantów, każdy inteligentny agent rychło utonąłby w rzeczywistym środowisku.

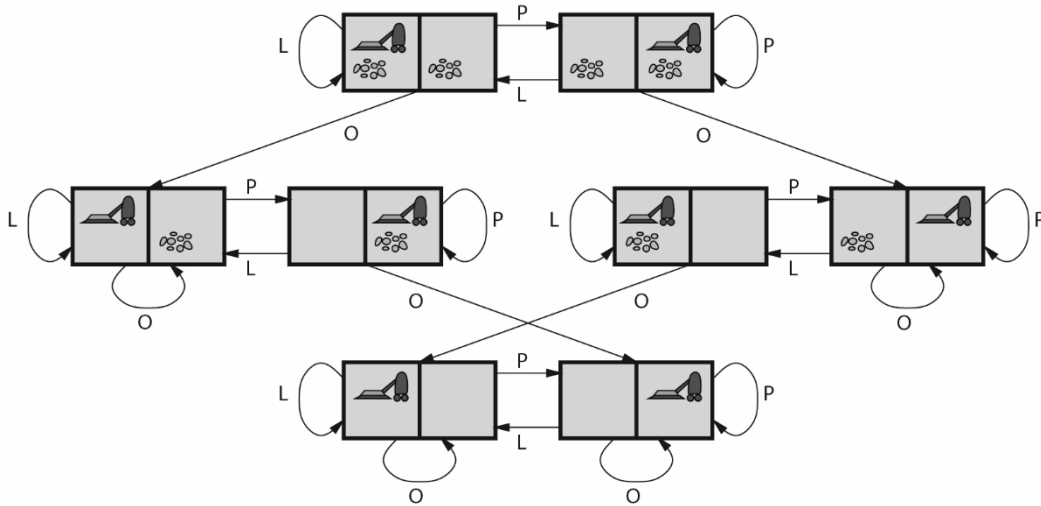
3.2. Przykładowe problemy

Opisane podejście do rozwiązywania problemów daje się zastosować do szerokiej gamy różnych środowisk zadaniowych, zaprezentujemy je w odniesieniu do środowisk najbardziej znanych. Będziemy przy tym odróżniać problemy *rzeczywiste* od problemów *standaryzowanych*. **Problemy standaryzowane** służą do ilustracji i nauki określonych metod poszukiwania rozwiązań; charakteryzują się zwięzłym, dokładnym opisem i z tego względu mogą służyć jako punkty odniesienia dla porównywania wydajności różnych algorytmów. **Problemy rzeczywiste** — takie jak nawigowanie robotów — to takie, których rozwiązania są bezpośrednio wykorzystywane przez ludzi; a sformułowanie każdego z nich jest specyficzne, dalekie od normalizacji. Przyczyną tej specyfiki jest zróżnicowanie szczegółów środowiska zadaniowego, na przykład różne roboty używają różnego zestawu czujników, odbierając w efekcie różne ciągi percepcyjne.

3.2.1. Problemy standaryzowane

Środowiskiem **problemu gridowego** jest dwuwymiarowa prostokątna tablica (macierz) komórek, często określana przejętym z języka angielskiego terminem *grid*. Agent może poruszać się po tej tablicy w poziomie, w pionie i (w przypadku niektórych instancji problemu) na skos, o ile nie istnieje przeszkoda uniemożliwiająca przemieszczenie się do sąsiedniej komórki. Komórki mogą zawierać obiekty, które agent może podnosić, pchać, ciągnąć i w jeszcze inny sposób na nie oddziaływać; dwie sąsiednie komórki mogą być oddzielone ścianą, która uniemożliwia bezpośrednie przemieszczenie obiektu między nimi. Typową odmianą problemu gridowego jest problem agenta-odkurzacza z rozdziału 2., który to problem można sformułować następująco — zgodnie z przedstawioną wcześniej definicją:

- **Stany.** Jedynymi obiektami, które mogą znajdować się w komórkach, są: agent i kurz. Każda z komórek może znajdować się w dwóch stanach: *Czysta* albo *Brudna*, w oryginalnej wersji problemu macierz zawiera dwie komórki, co daje $2 \times 2 = 4$ stany zabrudzenia. Niezależnie od tego odkurzacz może znajdować się w jednej z dwóch komórek, co daje ogólną liczbę stanów środowiska $4 \times 2 = 8$ (jak pokazano na rysunku 3.2). Ogólnie środowisko n -komórkowe może znajdować się w jednym z $n \times 2^n$ stanów.
- **Stan początkowy.** Dowolny stan może pełnić rolę stanu początkowego.
- **Akcje.** W dwukomórkowym świecie odkurzacza zdefiniowaliśmy trzy akcje: *Odkurz*, *W_lewo* i *W_prawo*, w dwuwymiarowym świecie wielokomórkowym potrzebowalibyśmy dodatkowych akcji ruchowych, przede wszystkim *W_góre* i *W_dół*, co dałoby nam cztery **absolutne** rodzaje ruchów. Alternatywą mogłyby być cztery ruchy **egocentryczne**, czyli relatywne do bieżącej pozycji odkurzacza: *W_przód*, *W_tył*, *Skręć_w_lewo*, *Skręć_w_prawo*.



RYSUNEK 3.2. Graf przestrzeni stanów dla dwukomórkowego środowiska odkurzacza, uwidoczniający 8 stanów i trzy akcje: L — *W_lewo*, P — *W_prawo*, O — *Odkurz*

- **Model przejścia.** Akcja *Odkurz* usuwa wszelkie zabrudzenia z komórki, w której aktualnie znajduje się odkurzac. Akcja *W_przód* przesuwa odkurzac do sąsiedniej komórki w kierunku, w którym jest zwrócony swym przodem, wyjątek stanowi sytuacja, gdy przed odkurzacem znajduje się ściana i wspomniany ruch nie jest możliwy, wówczas akcja nie wywołuje żadnego ruchu. Analogicznie, akcja *W_tył* powoduje ruch w przeciwnym kierunku, o ile takowy jest możliwy. Akcja *W_lewo* powoduje obrót odkurzacza o kąt prosty w kierunku antyżegarowym, analogicznie akcja *W_prawo* powoduje obrót odkurzacza o kąt prosty w kierunku zegarowym; w obu przypadkach odkurzac nie opuszcza bieżącej komórki.
- **Stany docelowe.** Stanami docelowymi są wszystkie stany, w których żadna z komórek nie zawiera brudu.
- **Koszt akcji.** Koszt każdej akcji równy jest 1.

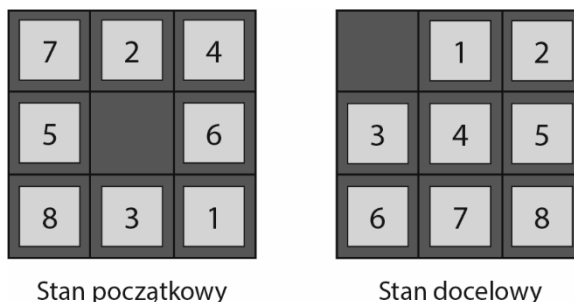
Środowiskiem innego problemu gridowego jest odmiana japońskiej układanki **Sokoban**. Każda z komórek gridu albo jest pusta, albo zawiera jedną skrzynię. „Rozwiązanie” układanki polega na przemieszczeniu każdej skrzyni do wyznaczonej dla niej a priori komórki docelowej. Gdy agent wchodzi do komórki zajętej przez skrzynię, a któraś z sąsiednich komórek jest pusta i nie jest oddzielona ścianą, agent wraz ze skrzynią przemieszcza się do tej komórki. Agentowi nie wolno przepychać skrzyni do innej komórki. W środowisku z n komórkami i b skrzyniami, bez odgradzających ścian, liczba możliwych stanów wynosi

$$\frac{n \times n!}{b! \times (n - b)!}$$

co dla gridu 8×8 komórek i 12 skrzyń daje 21×10^{13} .

Układanka doczekała się kilku znanych wariantów. W wariantcie **przesuwanych kafelków** (ang. *sliding-tile puzzle*) w gridzie istnieje co najmniej jedna wolna komórka, na którą można przesuwać sąsiadujące kafelki. W wariantcie o nazwie *Rush Hour* („godziny szczytu”) w komórkach gridu 6×6 znajdują się pojazdy — 12 samochodów osobowych i 4 ciężarówki; samochód osobowy ma wymiary 1×2 komórki, ciężarówka 1×3 komórki; zadaniem gracza jest uwolnienie

wskazanego pojazdu z „korka” ulicznego. Najbardziej chyba znanymi odmianami układanki są tzw. **puzzle-8**⁵ — 8 kafelków przesuwalnych wewnątrz gridu 3×3 komórki (rysunek 3.3) oraz **puzzle-15** — 15 kafelków w gridzie 4×4 komórki. Celem gry jest „uporządkowanie kafelków”, czyli sprawienie, by każdy z nich znalazł się na swoim miejscu (jak w prawej części rysunku 3.3).



RYSUNEK 3.3. Przykładowa instancja gry puzzle-8

Standaryzowane sformułowanie problemu dla odmiany puzzle-8 prezentuje się następująco:

- **Stany.** Stan to położenia poszczególnych kafelków.
- **Stan początkowy.** Dowolny stan może pełnić rolę stanu początkowego. Zauważmy, że własność parzystości dzieli przestrzeń stanów na pół — dokładnie połowa wszystkich stanów początkowych umożliwia osiągnięcie danego stanu docelowego. (patrz Ćwiczenie online 3. PART).
- **Akcje.** W świecie rzeczywistym akcją stanowi oczywiście przesunięcie kafelka, lecz ze względu na prostotę opisu wygodniej będzie rozpatrywać manipulowanie układanką jako przesuwanie się pustej komórki *W_lewo*, *W_prawo*, *W_góre* i *W_dół*. Oczywiście gdy pusta komórka znajduje się w którymś ze skrajnych położań, nie wszystkie z tych akcji mają zastosowanie.
- **Model przejścia.** Mapuje parę (stan bieżący, akcja) w nowy stan bieżący, na przykład zastosowanie akcji *W_lewo* do stanu widocznego w lewej części rysunku 3.3 spowoduje zamianę kafelka nr 5 z pustą komórką.
- **Stan docelowy.** Chociaż możemy wybrać dowolny stan w charakterze stanu docelowego, to zwykle przyjmuje się w tej roli uporządkowanie kafelków według numerów, jak w prawej części rysunku 3.3.
- **Koszt akcji.** Koszt każdej akcji równy jest 1.

Zauważmy, że każde sformułowanie problemu zawiera abstrakcję. Definicja każdej z akcji układanki puzzle-8 obejmuje wyłącznie układ kafelków przed i po wykonaniu akcji — pomijamy możliwość zablokowania się kafelka w trakcie przesuwania (i potrząsania tabliczką w celu likwidacji blokady), nie dopuszczamy także „wydłubowania” kafelków z tabliczki (na przykład za pomocą noża) i wstawiania ich na nowe miejsce. Uwzględnienie takich i podobnych fizycznych manipulacji niepotrzebnie zaciemniłoby obraz problemu.

⁵ W amerykańskiej terminologii układanka z k kafelkami określana jest terminem „ k -puzzle”, w polskim wydaniu tej książki konsekwentnie będę używał sformułowań „puzzle-8”, „puzzle-15” i „puzzle-24” dla układanki z (odpowiednio) 8, 15 i 24 kafelkami poruszającymi się na kwadratowej planszy o boku (odpowiednio) 3, 4 i 5 — *przyp. tłum.*

Ostatni z rozpatrywanych przez nas problemów zestandaryzowanych pochodzi od Donalda Knutha (1964) i ilustruje pojawienie się nieskończonej przestrzeni stanów. Knuth wyraził mianowicie przypuszczenie, że wychodząc od liczby 4 i stosując odpowiednio funkcje pierwiastka kwadratowego, obcięcia w dół do liczby całkowitej (ang. *floor*) i silni, możemy otrzymać jako wynik dowolną liczbę naturalną — na przykład 5:

$$\text{floor} \left(\sqrt{\sqrt{\sqrt{\sqrt{\sqrt{(4!)!}}}}} \right) = 5$$

Definicja tego problemu jest nadzwyczaj prosta:

- **Stany.** Dodatnie liczby rzeczywiste.
- **Stan początkowy.** 4.
- **Akcje.** Użycie funkcji $\sqrt{}$, *floor* lub silni (ta ostatnia tylko do argumentów będących liczbami naturalnymi).
- **Model przejścia.** Wynika z definicji użytych funkcji matematycznych.
- **Stan docelowy.** Żądana liczba naturalna.
- **Koszt akcji.** Koszt każdej akcji równy jest 1.

Przestrzeń stanów dla tego problemu jest nieskończona — silnia liczby naturalnej większej niż 2 jest od tej liczby większa. Stany problemu mogą mieć postać bardzo dużych liczb: w przedstawionym przykładzie wychodzimy od 4 i w pierwszym kroku obliczamy $4! = 24$, jednak już w drugim kroku mamy (bagatela)

$$24! = 620\,448\,401\,733\,239\,439\,360\,000$$

Nieskończone przestrzenie stanów pojawiają się często w zadaniach obejmujących generowanie wyrażeń matematycznych, dowodów twierdzeń, programów, obwodów logicznych i innych obiektów definiowanych rekurencyjnie.

3.2.2. Problemy ze świata rzeczywistego

Widzieliśmy już, jak *problem ze znajdowaniem trasy* jest definiowany w kategoriach określonych lokalizacji i przemieszczania się po łączących je drogach. Podobne algorytmy są używane w wielu różnych aplikacjach. Niektóre, takie jak systemy (webowe lub wbudowane w pokładowe komputery samochodów) prostego wyznaczania trasy na podstawie wyłącznie sieci dróg, stanowią proste rozszerzenia naszego rumuńskiego przykładu; ograniczenie ich użyteczności jest konsekwencją ignorowania zróżnicowanych kosztów przejazdu przez poszczególne drogi — zróżnicowanych z powodu różnego natężenia ruchu czy też zamknięcia pewnych odcinków z powodu robót drogowych. Aplikacje bardziej skomplikowane, służące (między innymi) do trasowania strumieni wideo w sieciach komputerowych, planowania operacji wojskowych czy układania rozkładów lotów, charakteryzują się znacznie większą złożonością.

Rozważmy problem projektowania serwisu webowego ułatwiającego planowanie podróży lotniczych — poniższa specyfikacja, i tak mocno uproszczona, jest o wiele bardziej złożona niż ta z rumuńskiego przykładu.

- **Stany.** Oczywistymi składnikami stanu są: bieżąca lokalizacja pasażera i lokalny czas. Ponadto, jako że koszt danego odcinka podróży może zależeć od kosztów odcinków poprzednich, ich taryf przewozowych i statusów (krajowy albo międzynarodowy), stan musi zawierać informacje o dotychczasowej historii podróży.
- **Stan początkowy.** Port lotniczy, z którego pasażer rozpoczyna podróż.
- **Akcje.** Wybierz dowolny lot z bieżącej lokalizacji, w dowolnej klasie miejsc, rozpoczynający się po bieżącej godzinie, z uwzględnieniem czasu niezbędnego na przesiadkę i formalności na lotnisku.

- **Model przejścia.** Parametrami nowego stanu, po przebyciu odcinka podróży, są między innymi: port końcowy tego odcinka i obowiązujący w tym porcie lokalny czas.
- **Stan docelowy.** Port końcowy ostatniego odcinka. Niekiedy cel może być bardziej skomplikowany, na przykład może uwzględniać wymaganie podróży bez międzylądowań.
- **Koszt akcji.** Kombinacja różnych czynników, między innymi kosztu pieniężnego, całkowitego czasu podróży i czasu oczekiwania przy przesiadkach, klasy miejsc, formalności celnych i imigracyjnych, pory dnia, typu samolotu i punktów lojalnościowych dla osób często podróżujących.

W przypadku komercyjnych systemów wspomagających planowanie podróży specyfikacja obejmuje wiele innych elementów, wynikających z (na przykład) bizantyjsko zagmatwanych systemów taryf. Poza tym doświadczony pasażer znakomicie zdaje sobie sprawę z faktu, że samoloty niekoniecznie kursują zgodnie z planem; dobry system planowania podróży powinien zatem przewidywać plan awaryjny na wypadek, gdy opóźnienie lub odwołanie lotu zniweczy zaplanowany rozkład podróży.

Problemy turystyczne tym się różnią od opisywanego problemu planowania podróży, że ich celem jest nie pojedyncza, końcowa lokalizacja, lecz *zbiór lokalizacji*, które mają zostać w ramach wycieczki odwiedzone. Jednym z takich problemów jest **problem komiwojażera** (ang. TSP — *traveling salesperson problem*): zadaniem komiwojażera (akwizytora, sprzedawcy, przedstawiciela handlowego) jest odwiedzenie wszystkich miast z ustalonego zbioru, oczywiście zgodnie z istniejącym systemem połączeń (dróg) między miastami. Każde połączenie między miastami cechuje się określonym kosztem, a celem komiwojażera jest wybór takiej kolejności odwiedzania miast (czyli takiej trasy), by łączny koszt wszystkich przebytych dróg był mniejszy od ustalonej wartości C (w optymalizacyjnej wersji problemu celem komiwojażera jest wybór trasy o *najmniejszym możliwym* koszcie). Informatycy włożyli mnóstwo wysiłków w badania i rozwój algorytmów rozwiązujących problem komiwojażera, ponieważ algorytmy te bardzo łatwo dały się przystosowywać do problemu zarządzania flotą pojazdów — ich zastosowanie do optymalizacji wykorzystania autobusów szkolnych w Bostonie przyniosło oszczędności rzędu 5 milionów USD, a ponadto zmniejszenie zanieczyszczenia środowiska, zmniejszenie natężenia ruchu oraz oszczędność czasu uczniów i kierowców (Bertsimas i in., 2019). Pokrewne algorytmy znalazły zastosowanie do planowania ruchów automatycznych wiertel do robienia otworów w płytkach drukowanych i optymalizacji wykorzystania wózków widłowych w magazynach.

Układy wielkiej skali integracji (ang. VLSI — *Very Large Scale of Integration*) to miliony tranzystorów i połączeń między nimi w pojedynczym chipie, które należy rozmieścić w taki sposób, by zminimalizować wymaganą powierzchnię, opóźnienia sygnałów i pojemności rozproszone, a jednocześnie zmaksymalizować wydajność produkcji. Problem **topografii** chipu pojawia się po zakończeniu fazy projektowania logicznego i rozwiązywany jest zwykle w dwóch fazach: projektowania **układu komórek** i **trasowania kanałów**. W pierwszej z tych faz pierwotne komponenty pogrupowane zostają w komórki, z których każda spełnia dobrze określoną funkcję, charakteryzuje się określonym kształtem i rozmiarem oraz wymaga określonej liczby połączeń z każdą z pozostałych komórek. Oczywiście komórki muszą zajmować rozłączne obszary (nie mogą zachodzić na siebie), a po ich rozmieszczeniu musi pozostać wystarczająco dużo miejsca na łączące je przewody. Projektowanie określonego *przebiegu* wspomnianych przewodów to właśnie trasowanie kanałów. Dla projektantów są to zadania ekstremalnie skomplikowane, ale warte rozwiązywania, i — o czym świadczą coraz to nowsze generacje procesorów — rozwiązywalne.

Nawigowanie robotów to uogólnienie opisanego wcześniej problemu znajdowania trasy: zamiast podążać ustalonymi ścieżkami (jak w naszym rumuńskim przykładzie), robot może wędrować po okolicy, wytyczając sobie własne ścieżki. Z perspektywy robota krążącego po płaskiej powierzchni problem ten jest w zasadzie dwuwymiarowy, ale już wyposażenie robota w (sterowalne) ramiona i nogi przydaje problemowi dodatkowych wymiarów, po jednym na każdy staw i przegub kończyny. Komplikuje to ów problem na tyle, że już samo zapewnienie *skończoności* przestrzeni poszukiwań, geometrycznie ciągłej przecież, wymaga zaawansowanych technik (patrz rozdział 26.). A to dopiero początek złożoności: funkcjonujący w praktyce robot musi dodatkowo radzić sobie z błędami w odczytach czujników i sterowaniu silnikami, z częściową obserwowalnością otoczenia i innymi agentami, którzy mogą wpływać na stan otoczenia.

Automatyczne sekwencjonowanie montażu złożonych obiektów (takich jak silniki elektryczne) przez roboty jest standardową praktyką przemysłową stosowaną od niemal półwiecza — minimalizowanie ręcznej pracy ludzkiej na liniach montażowych przynosi zwykle znaczne oszczędności czasu i kosztów. Algorytm sekwencjonowania najpierw znajduje jakąś wykonalną sekwencję montażu, po czym dąży do jej optymalizacji; jeżeli wybierze sekwencję niewykonalną, na którymś jej etapie okaże się, że dodanie kolejnej części jest niemożliwe bez cofnięcia (czyli zmarnotrawienia) części już wykonanej pracy. Badanie wykonalności sekwencji działań stanowi trudny problem wyszukiwania geometrycznego, ściśle powiązany z problemem nawigowania robotów, zatem generowanie kolejnych dopuszczalnych akcji jest najdroższą częścią procesu sekwencjonowania. Przeszukiwanie kompletnej przestrzeni możliwych akcji jest zadaniem niewykonalnym obliczeniowo, więc każdy praktyczny algorytm sekwencjonowania musi ograniczać się do niewielkiej części tej przestrzeni.

Pokrewnym sekwencjonowaniu czynności montażowych jest problem **projektowania białek**. Jego celem jest znajdowanie sekwencji aminokwasów, które związać (fałdować) się będą w trójwymiarowe struktury białek, posiadających własności leczenia niektórych chorób.

3.3. Algorytmy wyszukiwania

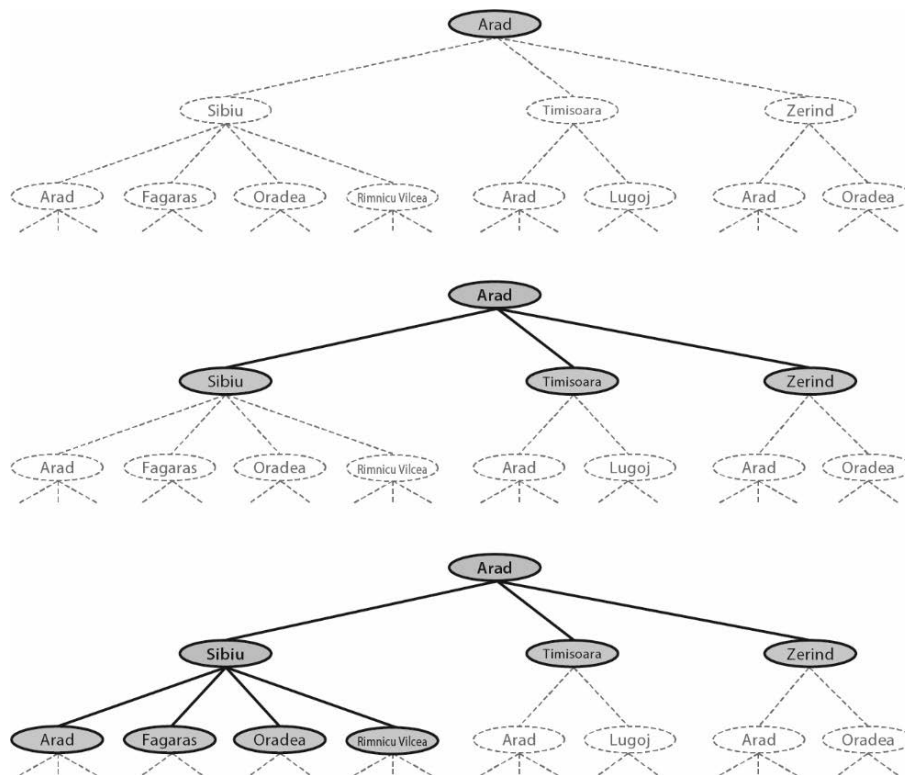
Wejściem do **algorytmu wyszukiwania** (ang. *search algorithm*) jest problem do rozwiązania, a wynikiem — rozwiązanie problemu albo informacja, że takie rozwiązanie nie istnieje. W tym rozdziale opisujemy algorytmy, które nakładają **drzewo wyszukiwawcze** na graf przestrzeni stanów; drzewo to zawiera różne ścieżki rozpoczynające się od stanu początkowego, a algorytm próbuje odnaleźć ścieżkę, która kończy się w stanie docelowym. Każdy **węzeł** (ang. *node*) tego drzewa reprezentuje jeden stan z przestrzeni stanów, a każda **gałąź** (ang. *branch*) reprezentuje akcję powodującą przejście między stanami. Węzeł stanowiący wspólny początek wszystkich ścieżek, odpowiadający stanowi początkowemu, nazywamy **korzeniem** (ang. *root*).

Należy wyraźnie odróżniać drzewo wyszukiwawcze od grafu przestrzeni stanów. W grafie tym każdy stan z przestrzeni (być może nieskończonej) odwzorowywany jest w jeden wierzchołek, a każda krawędź reprezentuje wykonanie określonej akcji w określonym stanie. W drzewie wyszukiwawczym kolejne węzły na danej ścieżce reprezentują pewien ciąg stanów; w drzewie tym może istnieć wiele ścieżek, a w konsekwencji wiele węzłów odpowiadających danemu stanowi. Drzewo z definicji nie zawiera cykli, zatem każdemu węzłowi odpowiada dokładnie jedna ścieżka prowadząca wstecz do korzenia.

Rysunek 3.4 przedstawia część drzewa wyszukiwawczego, reprezentującą kilka pierwszych kroków w poszukiwaniu trasy z Arad do Bukaresztu. Korzeń odpowiada stanowi głównemu, czyli miejscowości *Arad*; uwzględniając akcje (ACTIONS) dostępne w tym stanie możemy węzeł **rozszerzyć**, czyli za pomocą funkcji RESULT, zobaczyć dokąd prowadzą te akcje i dla każdej z nich **wygenerować** nowy węzeł, zwany **węzłem potomnym** (ang. *child node*) lub **następnikiem** (ang. *successor*) korzenia, który dla węzła potomnego jest **węzłem macierzystym** (ang. *parent node*) lub **poprzednikiem** (ang. *predecessor*) — operacja ta nosi nazwę **rozwinęcia** (ang. *expansion*) węzła macierzystego.

Teraz musimy wybrać, który z tych trzech węzłów potomnych rozważymy jako następny — to właśnie stanowi istotę wyszukiwania: podążanie za jedną wybraną opcją i pozostawianie innych na później. Wybierzmy zatem jako pierwszy do rozwinęcia węzeł *Sibiu*. Rezultatem tego rozwinęcia (patrz dolna część rysunku 3.4) jest wygenerowanie 6 nierozwiniętych węzłów (etykietowanych pogrubioną czcionką). Zbiór ten nazywamy **granicą** (ang. *frontier*). Stan, dla którego został wygenerowany węzeł, nazywamy stanem **osiągniętym** (ang. *reached*) niezależnie od tego, czy węzeł ten został rozwinięty, czy nie⁶; a węzeł reprezentujący ten stan — **węzłem osiągniętym**. Kolejne stadia rozwojowe drzewa wyszukiwawczego nałożone na graf przestrzeni stanów widoczne są na rysunku 3.5.

⁶ Niektórzy autorzy nazywają granicę drzewa **otwartą listą** (ang. *open list*), co jest zarówno mniej sugestywne pod względem geograficznym, jak i mniej adekwatne z perspektywy obliczeń — kolejka (ang. *queue*) jest w tej sytuacji bardziej efektywna niż lista. Konsekwentnie zbiór węzłów wcześniej rozwiniętych nazywany jest **listą zamkniętą** (ang. *closed list*), co w naszym ujęciu równoważne jest zbiorowi węzłów osiągniętych pomniejszonemu o granicę.

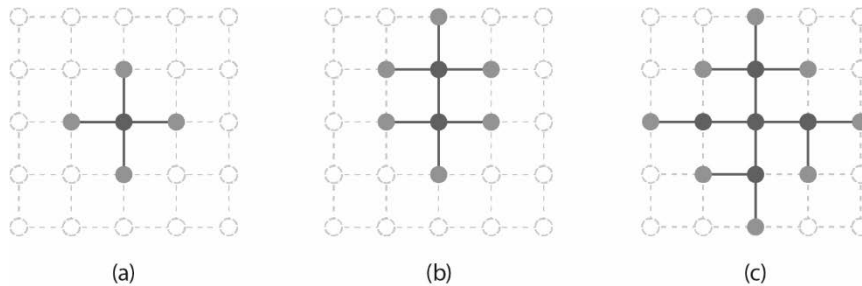


RYSUNEK 3.4. Trzy częściowe drzewa wyszukiwawcze odpowiadające poszukiwaniu drogi z Arad do Bukaresztu. Węzły rozwinięte mają kolor fioletowy i pogrubioną czcionkę w etykiecie, węzły graniczne, wygenerowane, ale jeszcze nie rozwinięte, mają kolor zielony i etykietowane są przy użyciu zwykłej czcionki. Stany odpowiadające obu tym rodzajom węzłów nazywane są stanami osiągniętymi. Węzły, które mogą zostać wygenerowane jako kolejne, otoczone są cienką linią przerywaną. Zauważmy, że w dolnym drzewie występuje cykliczne następstwo stanów Arad – Sibiu – Arad, wykluczające optymalność ścieżki, zatem ścieżka ta nie powinna być dalej rozwijana



RYSUNEK 3.5. Sekwencja wybranych etapów rozwojowych drzewa wyszukiwawczego generowanego przez graf z rysunku 3.1. Na każdym etapie rozwijamy wszystkie węzły tworzące granicę, rozszerzając w ten sposób ścieżki, ograniczając się do akcji, które mają sens i nie prowadzą do któregoś ze stanów już osiągniętych. Zauważmy, że w trzecim etapie rozwinięcie węzła Oradea doprowadziłoby do wygenerowania dwóch węzłów potomnych, reprezentujących stany już osiągnięte (Zerind i Sibiu), dlatego rezygnujemy z tego rozwinięcia

Zauważmy, że granica **oddziela** w grafie przestrzeni stanów dwa regiony: wewnętrzny, zawierający węzły już rozwinięte, i zewnętrzny, zawierający węzły reprezentujące stany, które nie zostały jeszcze osiągnięte. Właściwość tę zilustraliśmy na rysunku 3.6.



RYСУNEK 3.6. Zjawisko separacji przy przeszukiwaniu grafu, zilustrowane za pomocą prostokątnej siatki (gridu). Granica (kolor zielony) oddziela region wewnętrzny (kolor fioletowy) od zewnętrznego (brak koloru). Granica stanowi zbiór węzłów reprezentujących stany osiągnięte, ale jeszcze nierozwiniętych. Region wewnętrzny to zbiór węzłów już rozwiniętych, węzły w regionie zewnętrznym reprezentują natomiast stany jeszcze nieosiągnięte. W stadium (a) rozwinięty jest jedynie korzeń drzewa, w stadium (b) rozwinięty jest najwyżej położony węzeł granicy, stadium (c) to efekt rozwinięcia pozostałych węzłów potomnych korzenia, odwiedzanych w kolejności zgodnej z ruchem wskazówek zegara

3.3.1. Wyszukiwanie „najpierw najlepszy”

Skoro rozbudowa drzewa wyszukiwawczego odbywa się poprzez rozwijanie węzłów tworzących granicę, to naturalną kwestią staje się *kolejność* wyboru poszczególnych węzłów do rozwinięcia. Ogólna strategia rozwiązywania tej kwestii nazywana jest potocznie „**najpierw najlepszy**” (ang. *best first*) i wyznacza tę kolejność na podstawie pewnej **funkcji ewaluacyjnej** (ang. *evaluation function*) — im *mniej* wartość tej funkcji dla konkretnego węzła, tym wcześniej węzeł ten zostanie rozwinięty. W programie na listingu 3.1 funkcja ta jest parametrem o nazwie *f*. W każdej iteracji wybieramy, spośród nierozwiniętych jeszcze węzłów tworzących granicę, ten węzeł *n*, dla którego wartość $f(n)$ jest najmniejsza; jeżeli węzeł ten reprezentuje stan docelowy, zwracamy go jako wynik, w przeciwnym razie wywołujemy dla tego węzła funkcję EXPAND generującą węzły potomne. Jeżeli dany węzeł potomny reprezentuje stan jeszcze nieosiągnięty, zostaje dodany do zbioru węzłów tworzących nową granicę. Węzeł ten zostaje jednak dodany do nowej granicy mimo reprezentowania stanu osiągniętego, jeśli ścieżka łącząca go z korzeniem ma mniejszy koszt niż wszystkie inne ścieżki prowadzące od stanu początkowego (korzenia) do tegoż stanu. Wynikiem działania algorytmu jest węzeł reprezentujący stan docelowy, bądź sygnał o niepowodzeniu (czyli nieznaalezieniu ścieżki prowadzącej od stanu początkowego do stanu docelowego). Zarówno szczegóły działania algorytmu, jak i zwracany przezeń wynik, zależne są oczywiście od wyboru funkcji *f* — zajmiemy się tym w dalszej części rozdziału.

LISTING 3.1. Algorytm wyszukiwania według strategii „najpierw najlepszy” i funkcja EXPAND dokonująca rozwijania węzła. Wykorzystywane struktury danych opisujemy w sekcji 3.3.2, znaczenie instrukcji **yield** wyjaśnione jest w Dodatku B

```
function BEST-FIRST-SEARCH(problem, f) returns węzeł docelowy albo sygnalizacja niepowodzenia
    węzeł  $\leftarrow$  NODE(STATE=problem.INITIAL)
    granica  $\leftarrow$  kolejka priorytetowa uporządkowana przez funkcję f, początkowo węzeł jest jej jedynym elementem
```

```

osiągnięte ← tabela przeglądowa kojarząca stany z reprezentującymi je węzłami.
             Początkowo zawiera jeden element — stan problem.INITIAL
             i reprezentujący go węzeł.

while not IS-EMPTY(granica) do
    węzeł ← POP(granica)
    if problem.IS-GOAL(węzeł.STATE) then return węzeł
    for each potomny in EXPAND(problem, węzeł) do
        s ← potomny.STATE
        if (s nie wchodzi w skład osiągnięte) or potomny.PATH-COST < osiągnięte[s].PATH-COST
        then
            osiągnięte[s] ← potomny
            dodaj potomny do granica
    return niepowodzenie

function EXPAND(problem, węzeł) yields węzły
    s ← węzeł.STATE

    for each akcja in problem.ACTIONS(s) do
        s' ← problem.RESULT(s, akcja)
        koszt ← węzeł.PATH-COST + problem.ACTION-COST(s, akcja, s')
        yield NODE(STATE=s', PARENT=węzeł, ACTION=akcja, PATH-COST=koszt)

```

3.3.2. Struktury danych związane z wyszukiwaniem

Algorytmy wyszukiwania wykorzystują struktury danych odzwierciedlających bieżącą postać drzewa wyszukiwanego. Węzeł tego drzewa reprezentowany jest przez strukturę `NODE` złożoną z czterech następujących komponentów, oznaczających kolejno:

- `STATE` — stan reprezentowany przez węzeł;
- `PARENT` — wskaźnik na węzeł macierzysty, czyli ten, który wygenerował niniejszy węzeł;
- `ACTION` — akcję, której zastosowanie do węzła macierzystego spowodowało wygenerowanie niniejszego węzła;
- `PATH-COST` — całkowity koszt ścieżki od stanu początkowego do niniejszego węzła. We wzorach matematycznych odpowiada mu oznaczenie $g(\text{węzeł})$.

Poruszanie się wstecz po ścieżce przy użyciu wskaźnika `PARENT` umożliwia dotarcie do korzenia drzewa i przy okazji prześledzenie historii stanów i akcji poprzedzających bieżący węzeł. Gdy ów bieżący węzeł reprezentuje stan docelowy, taka wycieczka udostępnia kompletną ścieżkę stanowiącą rozwiązanie problemu.

Potrzebujemy także struktury danych reprezentującej granicę drzewa. Doskonale nadają się do tego **kolejki**, zważywszy na następujące operacje wykonywane w ramach algorytmów wyszukiwania:

- `IS-EMPTY(granica)` — zwraca *prawda* wtedy i tylko wtedy, gdy *granica* nie zawiera żadnego węzła;
- `POP(granica)` — usuwa czołowy węzeł *granicy* i zwraca go jako wynik;
- `TOP(granica)` — zwraca jako wynik czołowy węzeł *granicy*, ale w przeciwieństwie do operacji `POP` nie usuwa go;
- `ADD(węzeł, granica)` — dodaje węzeł do *granicy*.

W algorytmach wyszukiwania wykorzystywane są trzy rodzaje kolejek:

- **Kolejka priorytetowa** — jej węzłem czołowym jest ten, któremu pewna funkcja ewaluacyjna f przypisuje najmniejszą wartość wśród wszystkich węzłów w kolejce. Taka kolejka wykorzystywana jest przez algorytm „najpierw najlepszy”.
- **Kolejka FIFO** (ang. *First-In-First-Out* — „pierwszy wchodzi, pierwszy wychodzi”) — jej węzłem czołowym jest węzeł dodany najdawniej.
- **Kolejka LIFO** (ang. *Last-In-First-Out* — „ostatni wchodzi, pierwszy wychodzi”) — zwana także **stosem** (ang. *stack*); jej węzłem czołowym jest węzeł ostatnio dodany. Jak niebawem pokażemy, ten rodzaj kolejki wykorzystywany jest przez wyszukiwanie w głąb.

Zbiór węzłów reprezentujących stany osiągnięte może być reprezentowany w formie tabeli przeglądowej, której kłuczami są węzły, a wartościami odpowiadające tym węzłom stany.

3.3.3. Redundantne ścieżki

Drzewo wyszukiwawcze pokazane w dolnej części rysunku 3.4 obejmuje ścieżkę z Arad do Sibiu i z powrotem do Arad. *Arad* jest więc **stanem powtórzonym** na tej ścieżce, w tym przypadku wskutek wystąpienia **cyklu** (zwanego także **zapętloną ścieżką**). Mimo iż cała przestrzeń liczy tylko 20 stanów, kompletne drzewo wyszukiwawcze jest *nieskończone*, ponieważ nie istnieje ograniczenie co do tego, jak często wolno nam przechodzić przez pętlę.

Cykl jest szczególnym przypadkiem **redundantnej ścieżki**. Z Arad do Sibiu można dostać się dwiema drogami: bezpośrednio (140 mil) lub przez Zerind i Oradea (297 mil). Ta druga ścieżka nie ma jednak znaczenia, jako „gorsza”, czyli dłuższa; przejście między tą samą parą stanów — od *Arad* do *Sibiu* — może bowiem dokonać się po krótszej ścieżce. Ścieżka *Arad – Zerind – Oradea – Sibiu* jest ścieżką redundantną — nie należy brać jej pod uwagę przy poszukiwaniu optymalnych ścieżek.

Wyobraźmy sobie agenta poruszającego się po gridzie o rozmiarach 10×10 komórek, z możliwością jego swobodnego (nie ma przeszkód między komórkami) przemieszczania się do każdej z komórek sąsiednich względem bieżącej. Agent znajdujący się na dowolnym polu może dotrzeć do dowolnego innego w co najwyżej 9 ruchach. Ale liczba ścieżek o długości 9 to trochę mniej niż 89 („trochę mniej” ze względu na krawędzie gridu), czyli około 100 milionów. Oznacza to, że do każdej ze 100 komórek gridu prowadzi ok. miliona redundantnych ścieżek o długości 9, ergo — jeśli wyeliminujemy redundantne ścieżki z poszukiwań, możemy skrócić czas wyszukiwania milion razy! Zwykle się mawia, że *kto zapomina o swej przeszłości, ściągga na siebie klątwę jej powtarzania* i ta uwaga idealnie wpisuje się w algorytmy wyszukiwania. Brzmi groźnie, ale tak naprawdę tylko pozornie, bo złowieszcze fatum można oddalić co najmniej na trzy sposoby.

Po pierwsze, możemy przytoczone ostrzeżenie potraktować dosłownie i zapamiętywać wszystkie osiągnięte stany (jak ma to miejsce w wyszukiwaniu „najpierw najlepszy”). Pozwala to na wykrywanie wszystkich ścieżek wiodących do poszczególnych stanów i zapamiętywanie dla każdego stanu tylko tej najlepszej. Jest to preferowana metoda dla przestrzeni stanów, w których istnieje wiele redundantnych ścieżek, ale jej ograniczeniem są wymagania pamięciowe związane z zapamiętywaniem osiągniętych stanów — ich tablica może nie mieścić się w pamięci.

Po drugie, możemy wybrać drugą skrajność — oddalić irracjonalną obawę przed „powtarzaniem przeszłości” i potraktować je jako rzecz normalną. Istnieją pewne sformułowania problemów, w których istnienie dwóch ścieżek prowadzących do tego samego stanu jest bardzo rzadkie lub wręcz wykluczone. Rezygnując z wykrywania redundantnych ścieżek, nie musimy zapamiętywać historii osiągniętych stanów i zmniejszamy wymagania pamięciowe dla algorytmu. Takim problemem jest na przykład problem montowania jakiegoś mechanizmu czy podzespołu z części (bądź też kompletowanie zamówienia do wysyłki), gdzie każda akcja polega na dodaniu jednej części, ale kolejność dodawania poszczególnych części podlega pewnym ograniczeniom — jeżeli należy dodać pewne dwie części *A* i *B*, to konieczne musi się to odbyć w kolejności „najpierw *A*, potem *B*”, nie odwrotnie, i takich ograniczeń nałożonych na pary części może być bardzo wiele. Algorytm wyszukiwania nazywamy **wyszukiwaniem grafowym**, jeśli prowadzi on kontrolę

redundantnych ścieżek, i **wyszukiwaniem drzewiastym**⁷ w przeciwnym razie. Algorytm BEST-FIRST-SEARCH z listingu 3.1 to wyszukiwanie grafowe; stanie się on wyszukiwaniem drzewiastym, gdy usuniemy z niego wszystkie odwołania do tablicy *osiqgnięte* — algorytm będzie wymagał mniej pamięci, ale jednocześnie będzie działał wolniej.

Po trzecie wreszcie, wykazując pewną dozę ostrożności (licho nie śpi!) możemy pójść na pewien kompromis i ograniczyć się do sprawdzania cykli, rezygnując ze sprawdzania innych ścieżek redundantnych. Cykle możemy łatwo wykrywać bez zaangażowania dodatkowej pamięci, wykorzystując wskaźniki na węzły macierzyste (PARENT) w strukturach NODE reprezentujących węzły i sprawdzając, czy dany stan nie pojawił się wcześniej w łańcuchu utworzonym przez te wskaźniki. Niektóre implementacje tego pomysłu przeprowadzają kompletną analizę tego łańcucha, wykrywając w ten sposób wszystkie cykle, inne ograniczają się do kilku kroków wstecz (na przykład do rodzica, dziadka, prapradziadka itp.), wykrywając tym samym jedynie cykle, które nie przekraczają założonej długości; zazwyczaj funkcjonują wówczas inne mechanizmy radzące sobie z dłuższymi cyklami.

3.3.4. Wydajność rozwiązywania problemów

Zanim zajmiemy się szczegółami różnych algorytmów wyszukiwania, przyjrzyjmy się podstawowym kryteriom ich oceny — jakościowej i ilościowej — stanowiącym jednocześnie kryteria wyboru konkretnego algorytmu, jako najlepszego w kontekście danego problemu i wymogów dotyczących jego rozwiązania. Najważniejsze spośród tych kryteriów to:

- **Zupełność** (ang. *completeness*) — jeśli istnieje rozwiązanie problemu, algorytm daje gwarancję jego znalezienia; w przypadku braku rozwiązania algorytm daje gwarancję sygnalizacji tego faktu.
- **Optymalność kosztowa** (ang. *cost optimality*) — spośród wszystkich rozwiązań algorytm zwraca to charakteryzujące się najmniejszym kosztem.
- **Złożoność czasowa** (ang. *time-complexity*)⁸ — to czas potrzebny na znalezienie rozwiązania, w zależności od rozmiaru problemu: może być wyrażony w jednostkach czasu rzeczywistego lub (bardziej abstrakcyjnie) liczbą analizowanych stanów bądź akcji.
- **Złożoność pamięciowa** (ang. *space-complexity*) — to wielkość pamięci niezbędnej do wykonania algorytmu, w zależności od rozmiaru problemu.

Aby zrozumieć znaczenie zupełności algorytmu, rozważmy problem wyszukiwania z jednym celem. Ten cel może znajdować się w dowolnym miejscu w przestrzeni stanów; dlatego algorytm zupełny musi być w stanie systematycznie badać każdy stan osiągalny ze stanu początkowego. W skończonych przestrzeniach stanów można to osiągnąć dość prosto: badając konsekwentnie wszystkie ścieżki i odrzucając te, które są cyklami (na przykład *Arad – Sibiu – Arad*) ostatecznie dotrzemy do każdego (osiągalnego) stanu.

W nieskończonych przestrzeniach stanów musimy być jednak bardziej ostrożni. Przypomnijmy sobie hipotezę Knutha (ostatni ze standaryzowanych problemów w sekcji 3.2.1). Jedna ze ścieżek, wychodzących ze stanu początkowego „4”, mogłaby prowadzić w nieskończoność przez kolejne złożenia funkcji „silnia” (4, 4!, (4!)!, ((4!)!) itd.). Analogicznie mają się sprawy w przypadku agenta poruszającego się po *nieskończonym* gridzie (bez przeszkód między komórkami) — jedna z (nieskończonych) ścieżek reprezentuje ruch po linii prostej. W obu przypadkach algorytm nie powraca do stanów już osiągniętych, ale jego eksplorowanie przestrzeni stanów nigdy się nie skończy. Nie jest to więc algorytm zupełny.

⁷ Sformułowanie „wyszukiwanie drzewiaste” bierze się stąd, że choć przestrzeń stanów reprezentowana jest przez graf, to w procesie wyszukiwania traktujemy go tak, *jakby był drzewem* — czyli jakby dla każdego stanu istniała dokładnie jedna ścieżka łącząca go ze stanem początkowym.

⁸ Synonimiczne określenie to „dopuszczalność” (ang. *admissibility*) i (po prostu) „optymalność” — to ostatek jest jednak mylące, bo może odnosić się do optymalności według innego kryterium.

Aby algorytm wyszukiwania był kompletny, musi być **systematyczny** w tym sensie, że eksplorując nieskończoną przestrzeń stanów, osiągnie ostatecznie dowolny stan połączony ze stanem początkowym. Na przykład dla nieskończonego gridu systematyczna eksploracja komórek, rozpoczynająca się od pewnej komórki stanowiącej stan początkowy. Mogłaby przebiegać wzdłuż spiralnej linii przebiegającej w kierunku zegarowym najpierw komórki sąsiednie wspomnianej komórki początkowej, potem komórki odległe o 2 pozycje od tejże, potem o 3 itd. Niestety, jeżeli *żadna* komórka gridu nie będzie rozwiązaniem problemu, algorytm będzie kreślił nieskończoną spiralę i nigdy nie zwróci informacji o braku rozwiązania.

Złożoność czasową i złożoność pamięciową rozpatruje się w odniesieniu do rozmiaru problemu, czyli pewnego aspektu jego trudności. W informatyce teoretycznej za rozmiar problemu wyszukiwania przyjmuje się zwykle wielkość grafu przestrzeni stanów, czyli sumę wierzchołków i krawędzi tego grafu ($|V|+|E|$) — czyli sumaryczną liczbę stanów i par (stan, akcja możliwa w tym stanie). Jest to odpowiednie w sytuacji, gdy wspomniany graf ma postać jawnie określonej struktury, takiej jak mapa Rumunii z rysunku 3.1. W wielu problemach sztucznej inteligencji graf przestrzeni stanów określony jest *niejawnie*, poprzez stan początkowy, akcje i model przejścia. W takim przypadku złożoność problemu można mierzyć w kategoriach **głębokości** (ang. *depth*) d , czyli liczby akcji w optymalnym rozwiązaniu, maksymalnej liczby akcji na dowolnej ścieżce m lub **czynnika rozgałęzienia** (ang. *branching factor*) b , czyli liczby następników węzła wymagającego analizy.

3.4. Strategie wyszukiwania niedoinformowanego

Algorytm wyszukiwania, w którym niemożliwe jest określenie, jak bardzo odległy jest dany stan od celu (celów), nazywamy algorytmem **niedoinformowanym** (ang. *uninformed*). Nasz agent znajdujący się w Arad, zamierzający dotrzeć do Bukaresztu, nieznający geografii Rumunii, nie wie, które z sąsiednich miast — Zerind czy Sibiu — znajduje się bliżej celu jego podróży i do którego z nich się udać. Z kolei agent poinformowany (patrz podrozdział 3.5) wiedziałby, że znacznie bliżej Bukaresztu znajduje się Sibiu, a więc prawdopodobnie prowadząca przezeń trasa będzie krótsza niż ta wiodąca przez Zerind.

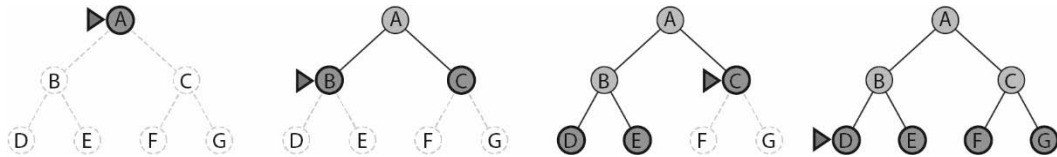
3.4.1. Wyszukiwanie wszерz

Gdy wszystkie akcje mają taki sam koszt, odpowiednią strategią wyszukiwania jest **wyszukiwanie wszерz** (ang. *breadth-first search*), zgodnie z którym bezpośrednio po rozwinięciu danego węzła rozwijane są wygenerowane właśnie węzły potomne — zasadę tę stosuje się rekurencyjnie, poczynając od korzenia. Algorytm stosujący tę strategię jest systematyczny i jest zupełny nawet w nieskończonych przestrzeniach stanów. Algorytm wyszukiwania wszерz możemy uzyskać z algorytmu BEST-FIRST-SEARCH z listingu 3.1 poprzez zdefiniowanie funkcji ewaluacyjnej $f(n)$ jako głębokości węzła n , czyli akcji potrzebnych do jego osiągnięcia.

Jednak dzięki kilku sztuczkom możemy uzyskać algorytm znacznie bardziej efektywny — bardziej efektywną od kolejki priorytetowej okazuje się bowiem w tej sytuacji kolejka FIFO, zapewniająca właściwą kolejność przetwarzania węzłów. Nowe węzły (które oczywiście mają większą głębokość niż ich węzeł macierzysty) wstawiane są na koniec kolejki, dzięki czemu elementy tej kolejki uporządkowane są rosnąco względem głębokości — węzły są więc przetwarzane (rozwijane) poczynawszy od „najpłytszych”. Co więcej, struktura *osiągnięte*, w algorytmie BEST-FIRST-SEARCH będąca tabelą przeglądową mapującą stany na węzły, tym razem jest jedynie zbiorem stanów; jest to wystarczające, ponieważ gdy osiągniemy dany stan, nigdy później nie znajdziemy już krótszej ścieżki prowadzącej do niego⁹. Kapitalną tego konsekwencją jest możliwość **wczesnej weryfikacji celu** — jeżeli stan reprezentowany przez nowo wygenerowany

⁹ Czyli ścieżki o mniejszym koszcie, jako że wszystkie akcje mają taki sam koszt — *przyp. tłum.*

węzeł okaże się stanem docelowym, będzie takim stanem osiągniętym po możliwie najkrótszej (najtańszej) ścieżce, czyli *będzie rozwiązaniem problemu*. W algorytmie BEST-FIRST-SEARCH taka weryfikacja możliwa jest dopiero po po-braniu węzła z kolejki, jest więc przykładem **późnej weryfikacji celu**. Rysunek 3.7 przedstawia przebieg wyszukiwania wszerz w drzewie binarnym, zaś na listingu 3.2 widoczny jest kod algorytmu BREADTH-FIRST-SEARCH — zwróć uwagę na podobieństwa i różnice w stosunku do BEST-FIRST-SEARCH.



RYСУNEK 3.7. Wyszukiwanie wszerz w prostym drzewie binarnym. Na każdym etapie węzeł przeznaczony do rozwinięcia zaznaczony jest trójkątnym wskaźnikiem

LISTING 3.2. Algorytmy wyszukiwania wszerz i wyszukiwania przy jednolitym koszcie

```
function BREADTH-FIRST-SEARCH(problem) returns węzeł docelowy albo sygnalizacja niepowodzenia
    węzeł ← NODE(problem.INITIAL)
    if problem.IS-GOAL(węzeł.STATE) then return węzeł
    granica ← kolejka FIFO, początkowo węzeł jest jej jedynym elementem
    osiągnięte ← {problem.INITIAL}
    while not IS-EMPTY(granica) do
        węzeł ← POP(granica)
        for each potomny in EXPAND(problem, węzeł) do
            s ← potomny.STATE
            if problem.IS-GOAL(s) then return potomny
            if (s nie wchodzi w skład osiągnięte) then
                dodaj s do osiągnięte
                dodaj potomny do granica
    return niepowodzenie

function UNIFORM-COST-SEARCH(problem) returns węzeł docelowy albo sygnalizacja niepowodzenia
    return BEST-FIRST-SEARCH(problem, PATH-COST)
```

Wyszukiwanie wszerz zawsze znajduje rozwiązanie z minimalną liczbą akcji, ponieważ gdy generowany jest węzeł na głębokości d , wygenerowane są już wszystkie węzły na głębokości $d-1$; gdyby któryś z tych ostatnich był rozwiązaniem, algorytm zauważyłby ten fakt i zwrócił to rozwiązanie. Jeśli wszystkie akcje na ścieżce będącej rozwiązaniem mają jednakowy koszt, to ścieżka ta (jako przeciw najkrótsza) jest jednocześnie ścieżką o najmniejszym koszcie, zatem w takiej sytuacji algorytm wyszukiwania wszerz jest algorytmem optymalnym kosztowo. Niezależnie od jednolitego czy zróżnicowanego kosztu akcji, algorytm wyszukiwania wszerz jest algorytmem zupełnym. Aby oszacować jego złożoność (czasową i pamięciową), wyobraźmy sobie wyszukiwanie drzewa, którego każdy węzeł (który nie jest liściem) ma dokładnie b węzłów potomnych. Korzeń drzewa generuje b węzłów, z których każdy generuje b kolejnych węzłów itd.; gdy wygenerowane zostaną wszystkie węzły na głębokości d , ogólna liczba wygenerowanych węzłów wynosi

$$1 + b + b^2 + b^3 + \dots + b^d = \frac{b^{d+1} - 1}{b - 1} = O(b^d)$$

Wszystkie węzły egzystują jednocześnie w pamięci, zatem zarówno czasowa, jak i pamięciowa złożoność algorytmu jest rzędu $O(b^d)$. Konkretnie wartości kryjące się za *wykładniczą* złożonością mogą być porażające. Wyobraźmy sobie typowy problem ze świata rzeczywistego, w którym czynnik rozgałęzienia wynosi $b = 10$; problem ten rozwiązywany jest za pomocą komputera, którego procesor zdolny jest generować milion węzłów w ciągu sekundy, a każdy węzeł wymaga 1 kilobajta pamięci. Dojście algorytmu do głębokości $d = 10$ zajmie procesorowi niecałe 3 godziny, lecz wymagania pamięciowe sięgną wówczas 10 terabajtów¹⁰ — zatem *nie czas wyszukiwania, ale wymagania pamięciowe są prawdziwym problemem w przypadku algorytmu wyszukiwania wszerek*. Złożoności czasowej nie należy bynajmniej lekceważyć — przy głębokości $d = 14$ wyszukiwanie (nawet przy zaspokojeniu wymagań pamięciowych rzędu 100 000 terabajtów) zajęłoby 3,5 roku. Jak widać, *problemy wyszukiwania o złożoności wykładniczej nie dadzą się rozwiązywać za pomocą algorytmów niedoinformowanych, z wyjątkiem być może bardzo prostych przypadków*.

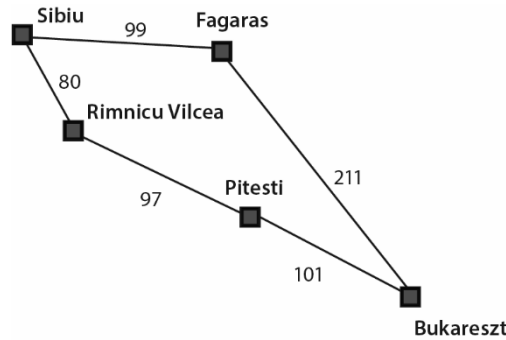
3.4.2. Algorytm Dijkstry — wyszukiwanie przy jednolitym koszcie

Gdy koszty akcji są zróżnicowane, oczywistym wyborem wydaje się użycie wyszukiwania „najpierw najlepszy” z funkcją ewaluacyjną $f(n)$ zwracającą koszt ścieżki rozpoczynającej się w korzeniu, a kończącej w węźle n . Wynalazcą tego algorytmu jest Edsger Wybe Dijkstra (1930 – 2002), holenderski pionier informatyki; oczywiście z tego względu społeczność informatyków nazywa ten algorytm „algorytmem Dijkstry” (choć to tylko jeden z wielu cennych wkładów Dijkstry w rozwój informatyki). Na gruncie sztucznej inteligencji algorytm ten nosi nazwę **wyszukiwania przy jednolitym koszcie** (ang. *uniform-cost search*). Podczas gdy algorytm wyszukiwania wszerek porusza się po drzewie wyszukiwawczym „falą” przebiegającą kolejne głębokości — $d = 1$, $d = 2$ itd. — algorytm Dijkstry rozprzestrzenia się po drzewie poprzez ścieżki o jednolitym koszcie. Funkcję UNIFORM-COST-SEARCH realizująca ten algorytm otrzymamy, wywołując funkcję BEST-FIRST-SEARCH z listingu 3.1 z funkcją ewaluacyjną, która dla węzła będącego jej argumentem zwraca wartość jego atrybutu PATH-COST — co pokazaliśmy w dolnej części listingu 3.2.

Spójrzmy na rysunek 3.8., ilustrujący problem dostania się z Sibiu do Bukaresztu. Węzeł Sibiu ma dwa następniki: Rimnicu Vilcea i Fagaras o kosztach (odpowiednio) 80 i 99. Wybieramy węzeł o mniejszym koszcie (Rimnicu Vilcea) i rozwijamy go, dochodząc do węzła Pitesti (nie mamy innego wyboru), którego koszt wynosi $80+97 = 177$. Węzeł Fagaras ma mniejszy koszt (99), więc rozwijamy go, dochodząc do Bukaresztu, czyli węzła o koszcie $99+211 = 310$. Węzeł ten zostaje dodany do *granicy*; choć reprezentuje on stan docelowy, algorytm tego nie zauważa — sprawdzanie, czy węzeł reprezentuje stan docelowy, następuje *dopiero po pobraniu tego węzła z granicy za pomocą operacji POP*, nie w momencie dodania go do *granicy*.

Algorytm kontynuuje pracę, przechodząc do węzła Pitesti i rozwijając go, co daje nową ścieżkę do Bukaresztu, o koszcie $80+97+101 = 278$. Jest to koszt mniejszy od poprzedniego (310), więc nowa ścieżka zastępuje poprzednią w tablicy *osiągnięte*, kończący ją węzeł Bukareszt zostaje dodany do *granicy*. Jako węzeł o najmniejszym ze wszystkich kosztów zostanie wkrótce pobrany z granicy (za pomocą operacji POP) i rozpoznany jako rozwiązanie. Zwróćmy uwagę na ważny fakt: gdybyśmy testowali węzeł na okoliczność reprezentowania stanu docelowego już w momencie dodawania tego węzła do *granicy*, algorytm zwróciłby ścieżkę wiodącą przez Rimnicu Vilcea i Pitesti, która nie jest najkrótsza.

¹⁰ Dzięki pamięci wirtualnej program może uzyskać tak dużą przestrzeń adresową, nawet jeśli fizyczna pamięć RAM komputera jest znacznie mniejsza, na przykład liczona w gigabajtach. Zwykle oznacza to jednak drastyczną degradację wydajności, spowodowaną intensywną wymianą danych między pamięcią RAM a plikiem wymiany — być może nawet do poziomu *jednego* węzła na sekundę (zamiast wspomnianego miliona). A wtedy szacowane 3 godziny rozciągną się na *kilka tygodni*. Nota bene to kolejne świadectwo wymagań, jakie badania nad sztuczną inteligencją stawiają przed infrastrukturą sprzętową i programową — *przyjp. tłum.*



RYСУNEK 3.8. Część przestrzeni stanów podróży po Rumunii, wybrana dla zilustrowania wyszukiwania przy jednolitym koszcie

Złożoność algorytmu wyszukiwania przy jednolitym koszcie rozpatruje się w kategoriach C^* — kosztu optymalnego rozwiązania¹¹ — oraz ε — dolnego ograniczenia kosztu każdej akcji, $\varepsilon > 0$. Złożoność — czasowa i pamięciowa — najgorszego przypadku wynosi wówczas $O(b^{1+\lceil C^*/\varepsilon \rceil})$, czyli może być znacznie większa od b^d ($\lceil x \rceil$ oznacza największą liczbę całkowitą nieprzekraczającą x). Jest tak dlatego, że algorytm wyszukiwania przy jednolitym koszcie może eksplorować duże drzewa „tanich” akcji przed zbadaniem kosztownych ścieżek zawierających być może użyteczne akcje. Gdy koszty wszystkich akcji są jednakowe, wartość $b^{1+\lceil C^*/\varepsilon \rceil}$ równa jest b^{d+1} i wyszukiwanie przy jednolitym koszcie równoważne jest wyszukiwaniu wszerz.

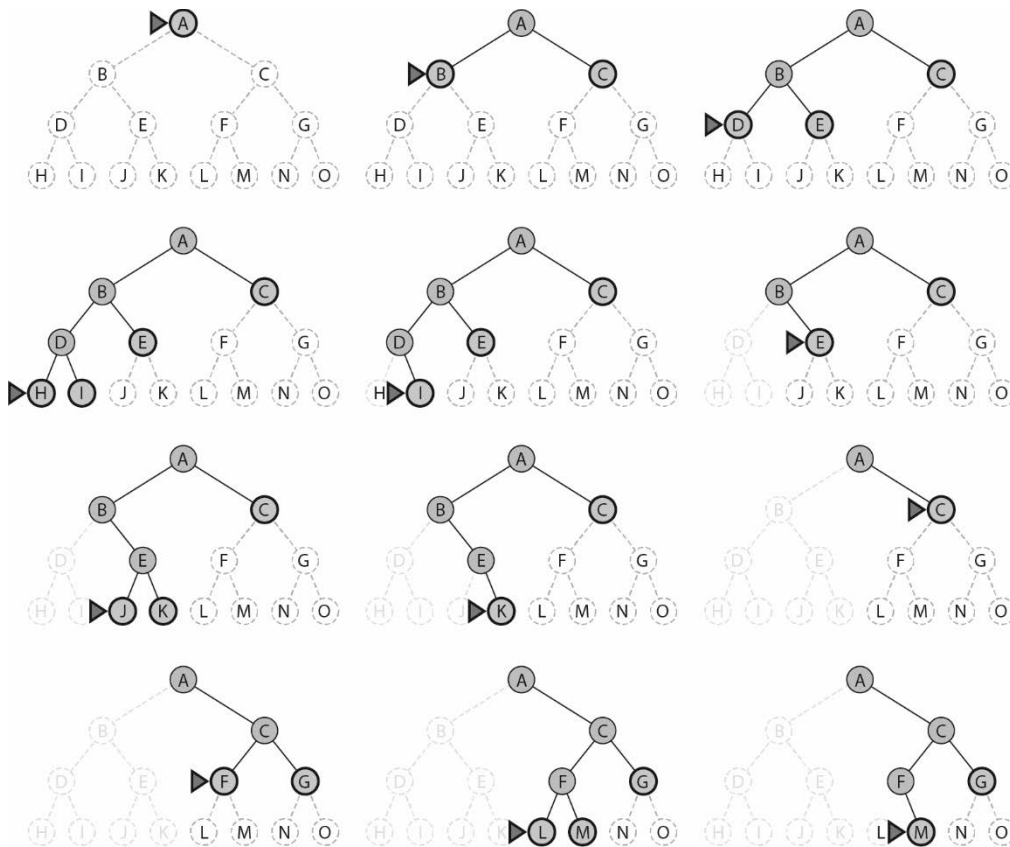
Algorytm wyszukiwania przy jednolitym koszcie jest algorytmem zupełnym i optymalnym kosztowo, ponieważ pierwsze znalezione przez niego rozwiązanie ma koszt nie większy niż koszt każdego innego węzła należącego do granicy. Algorytm ten systematycznie przegląda wszystkie ścieżki w kolejności wzrastających kosztów, nigdy nie dając się złapać w pojedynczą nieskończoną ścieżkę, pod warunkiem, że koszt każdej akcji jest większy od pewnego dodatniego ε .

3.4.3. Wyszukiwanie w głąb i problemy z pamięcią

Algorytm **wyszukiwania w głąb** (ang. *depth-first search*) zawsze rozwija w pierwszej kolejności *najgłębszy* węzeł na granicy. Można go otrzymać przez wywołanie algorytmu BEST-FIRST-SEARCH z funkcją ewaluacyjną $f(n)$ zwracającą *zanegowaną głębokość* węzła n , zwykle jednak implementuje się go jako wyszukiwanie drzewiaste, nie grafowe, nieutrzymujące tablicy osiągniętych stanów. Postęp przykładowego wyszukiwania ilustruje rysunek 3.9; po rozwinięciu węzła wyszukiwanie przechodzi natychmiast do jego węzłów potomnych, po czym „cofa się” do najbliższego najgłębszego węzła, który nadal posiada nierozwinięte węzły potomne. Algorytm wyszukiwania w głąb *nie jest optymalny kosztowo* — zwraca pierwsze znalezione rozwiązanie nawet jeśli nie jest najtańsze.

W skończonych przestrzeniach stanów o strukturze drzewiastej algorytm jest wydajny i zupełny — w przestrzeni acyklicznej, nawet jeśli dany stan będzie rozwijany wielokrotnie (w wyniku przemierzania różnych ścieżek prowadzących do niego), ale ostatecznie i tak odwiedzony zostanie każdy stan.

¹¹ W tym miejscu i w dalszym ciągu książki gwiazdka towarzysząca wielkości C oznacza optymalną wartość tej wielkości.



RYСУNEK 3.9. Dwanaście kroków (od lewej do prawej, z góry na dół) przebiegu wyszukiwania w głąb w drzewie binarnym, począwszy od stanu A, skończywszy na stanie docelowym M. Kolorem zielonym oznaczono węzły tworzące bieżącą granicę, trójkątny wskaźnik identyfikuje węzeł do aktualnego rozwijania. Węzły już rozwinięte oznaczone są kolorem fioletowym, węzły stanowiące kandydatury do następnych rozwinięć są niepokolorowane. Słabiej widoczne rozwinięte węzły niemające węzłów potomnych w granicy wykluczane są z dalszych poszukiwań

W przestrzeni stanów zawierającej cykle algorytm może utknąć w nieskończonej pętli, dlatego niektóre jego implementacje sprawdzają każdy nowy węzeł pod kątem cykli. Wreszcie — w nieskończonej przestrzeni stanów algorytm może utknąć na nieskończonej ścieżce, nawet jeśli przestrzeń ta nie zawiera cykli, nie jest więc zupełny w nieskończonych przestrzeniach.

Dlatego więc w ogóle sięgać po wyszukiwanie w głąb, skoro dostępne są wyszukiwanie wszerz i wyszukiwanie „najpierw najlepszy”? Otóż w przypadku problemów, do których można zastosować wyszukiwanie drzewiaste, wyszukiwanie w głąb cechuje się znacznie mniejszym zapotrzebowaniem na pamięć niż jego wspomniani konkurenci: nie ma potrzeby utrzymywania tablicy stanów osiągniętych, także granice są zwykle mniej liczne. Jeżeli mianowicie wyobrazimy sobie granicę wyszukiwania wszerz jako rozszerzającą się sferę, to granica wyszukiwania w głąb będzie w tej analogii wydłużającym się promieniem tej sfery.

W przypadku skończonej przestrzeni stanów o strukturze drzewa (takiego jak na rysunku 3.9) wyszukiwanie w głąb zajmuje czas proporcjonalny do liczby stanów, a jego wymagania pamięciowe są rzędu $O(bm)$, gdzie b jest czynnikiem rozgałęzienia, a m największą głębokością węzła w drzewie. Dzięki temu niektóre problemy, które przy wyszukiwaniu wszerek wymagałyby eksabajtów pamięci, można za pomocą wyszukiwania w głąb rozwiązywać kosztem zaledwie kilobajtów. Właśnie ze względu na ową oszczędność pamięci drzewiaste wyszukiwanie w głąb zostało powszechnie przyjęte w wielu obszarach sztucznej inteligencji jako przysłowiowy wół roboczy, między innymi w zadaniach dotyczących spełniania ograniczeń (patrz rozdział 6.), spełnialności w rachunku zdań (rozdział 7.) i programowania w logice (rozdział 9.).

Jeszcze skromniejszymi wymogami pamięciowymi charakteryzuje się odmiana wyszukiwania w głąb, zwana **wyszukiwaniem z nawrotami** (ang. *backtracking search*), którą szczegółowo opisujemy w rozdziale 6. W procesie rozwijania węzła generowany jest każdorazowo tylko jeden węzeł potomny (zamiast wszystkich), w węźle częściowo rozwiniętym przechowywana jest informacja o tym, który z węzłów potomnych wygenerować jako następny. Ponadto generowanie kolejnych węzłów potomnych odbywa się poprzez *modyfikowanie* opisu stanu w węźle macierzystym, zamiast przydzielania pamięci do zapamiętywania nowych stanów. Konieczność zapamiętywania tylko jednego stanu dla każdego węzła redukuje wymagania pamięciowe z $O(bm)$ do $O(m)$. Możliwe jest także kontrolowanie stanu bieżącej ścieżki za pomocą efektywnej struktury (w postaci zbioru), dzięki czemu wykrywanie cyklu następuje w czasie $O(1)$ (czyli natychmiast) zamiast w czasie $O(m)$. Warunkiem koniecznym do zastosowania wyszukiwania z nawrotami jest możliwość *anulowania* wykonanych już akcji (to właśnie stanowi istotę nawrotów). Wyszukiwanie z nawrotami ma krytyczne znaczenie dla pomyślnego rozwiązywania problemów charakteryzujących się rozbudowanym opisem stanu, na przykład zrobotyzowanych linii montażowych.

3.4.4. Wyszukiwanie z ograniczeniem głębokości i iteracyjne zagłębianie

Aby wykluczyć ryzyko utknięcia wyszukiwania w głąb na nieskończonej ścieżce, można **ograniczyć głębokość wyszukiwania**. Przyjmujemy mianowicie pewien limit głębokości λ i wszystkie węzły na głębokości λ traktujemy tak, jakby nie miały węzłów potomnych (patrz listing 3.3). Złożoność czasowa takiego wyszukiwania wynosi $O(b^\lambda)$, a złożoność pamięciowa $O(b \lambda)$. Niestety, jeżeli niefortunnie wybierzemy λ , algorytm może nie znaleźć rozwiązania (mimo iż ono istnieje), zatem stanie się niezupełny.

Ponieważ wyszukiwanie w głąb jest wyszukiwaniem drzewiastym, nie możemy generalnie zapobiec marnowaniu czasu na analizę redundantnych ścieżek, ale możemy eliminować cykle, za cenę dodatkowego czasu obliczeń. Możemy wykrywać większość cykli, posuwając się jedynie o kilka pozycji w łańcuchu węzłów macierzystych, dłuższe cykle eliminowane będą automatycznie ze względu na ograniczenie głębokości wyszukiwania.

W sensownym wyborze ograniczenia głębokości (λ) może pomóc wstępna analiza problemu. Skoro na przykład na mapie z rysunku 3.1 znajduje się 20 miast, to rozsądnym wyborem okazuje się $\lambda = 19$. Gdybyśmy jednak bliżej przyrzekli się problemowi zauważylibyśmy, że z każdego miasta można dotrzeć do dowolnego innego w maksymalnie 9 etapach. Liczba ta, nazywana **średnicą grafu** przestrzeni stanów, daje nam lepsze ograniczenie, skutkujące bardziej efektywnym wyszukiwaniem. Dla większości problemów właściwe ograniczenie staje się znane dopiero... po rozwiązaniu problemu.

Wyszukiwanie z iteracyjnym zagłębianiem (ang. *iterative deepening search*) to strategia doboru właściwej wartości λ metodą prób i błędów — $\lambda = 0$, $\lambda = 1$ itd. aż do skutku, czyli znalezienia rozwiązania albo zwrócenia przez algorytm sygnału *niepowodzenia* (zamiast sygnału przekroczenia *limitu* głębokości). Algorytm takiego zagłębiania wiódący jest na listingu 3.3. Iteracyjne zagłębianie wyszukiwania łączy w sobie zalety wyszukiwania w głąb i wyszukiwania wszerek. Podobnie jak wyszukiwanie w głąb, cechuje się skromnymi wymaganiami pamięciowymi — $O(bd)$ gdy istnieje rozwiązanie i $O(bm)$ przy jego braku (m jest liczbą stanów w przestrzeni). Podobnie jak wyszukiwanie wszerek, jest optymalne kosztowo dla problemów, w których wszystkie akcje mają taki sam koszt; jest ponadto zupełne w skończonych przestrzeniach stanów — acyklicznych oraz w przypadku, gdy analizujemy węzły pod kątem cykli, poruszając się w górę w łańcuchu węzłów macierzystych.

LISTING 3.3. Iteracyjne zagłębianie wyszukiwania (ITERATIVE-DEEPENING-SEARCH) zaimplementowane jako powtarzane wywoływanie wyszukiwania z ograniczeniem głębokości (DEPTH-LIMITED-SEARCH) dla coraz większych limitów głębokości. Wynik zwracany przez tę ostatnią technikę może być trojaki: węzeł reprezentujący stan docelowy, sygnalizacja przekroczenia limitu głębokości albo sygnalizacja braku rozwiązania po eksploracji wszystkich węzłów. Zwróćmy uwagę na brak struktury przechowującej stany osiągnięte, co znacznie redukuje zapotrzebowanie na pamięć (w porównaniu z wyszukiwaniem „najpierw najlepszy”), ale jednocześnie niesie ryzyko wielokrotnego eksplorowania tych samych stanów na różnych ścieżkach. Ponadto, jeśli funkcja IS-CYCLE nie wykrywa wszystkich cykli, algorytm może utknąć w nieskończonej pętli

```

function ITERATIVE-DEEPENING-SEARCH(problem) returns rozwiązanie albo sygnalizacja niepowodzenia
  for głębokość = 0 to ∞ do
    wynik ← DEPTH-LIMITED-SEARCH(problem, głębokość)
    if wynik ≠ przekroczenie_limitu then return wynik

function DEPTH-LIMITED-SEARCH(problem, λ) returns węzeł wynikowy albo sygnalizacja niepowodzenia
                                     albo sygnalizacja przekroczenia limitu głębokości
  granica ← kolejka LIFO (stos) początkowo z węzłem NODE(problem.INITIAL) jako jedynym elementem
  wynik ← niepowodzenie

  while not IS-EMPTY(granica) do
    węzeł ← POP(granica)
    if problem.IS-GOAL(węzeł.STATE) then return węzeł
    if DEPTH(węzeł) > λ then
      wynik ←przekroczenie_limitu
    else if not IS-CYCLE(węzeł) then
      for each potomny in EXPAND(problem, węzeł) do
        dodaj potomny do granica
  return wynik

```

Złożoność czasowa wyszukiwania z iteracyjnym zagłębianiem wynosi $O(b^d)$, gdy istnieje rozwiązanie, i $O(b^m)$, gdy nie ma żadnego. W każdej iteracji generowany jest nowy poziom węzłów, podobnie jak wyszukiwanie wszerz, ale jest istotna różnica: wyszukiwanie wszerz przechowuje w pamięci wszystkie węzły, podczas gdy iteracyjne zagłębianie ponownie przechodzi przez wszystkie poziomy, oszczędzając w ten sposób pamięć kosztem czasu obliczeń. Na rysunku 3.10 widoczne są cztery iteracje zagłębiania w drzewie binarnym, rozwiązanie znalezione zostaje w czwartej iteracji.

Wielokrotne odwiedzanie początkowych węzłów drzewa w ramach iteracyjnego zagłębiania może wydawać się marnotrawieniem czasu; przestrzenie stanów wielu problemów mają jednak tę cechę, że znakomita większość węzłów w ich drzewach wyszukiwawczych znajduje się na głębszych poziomach i czas stracony na ponowną eksplorację płytszych węzłów nie jest aż tak znaczący. Węzły najniższego poziomu (d) generowane są tylko raz, węzły poziomu bezpośrednio wyższego — dwukrotnie itd., węzły potomne korzenia generowane są d razy. Całkowita liczba aktów generowania węzłów wynosi więc w najgorszym przypadku

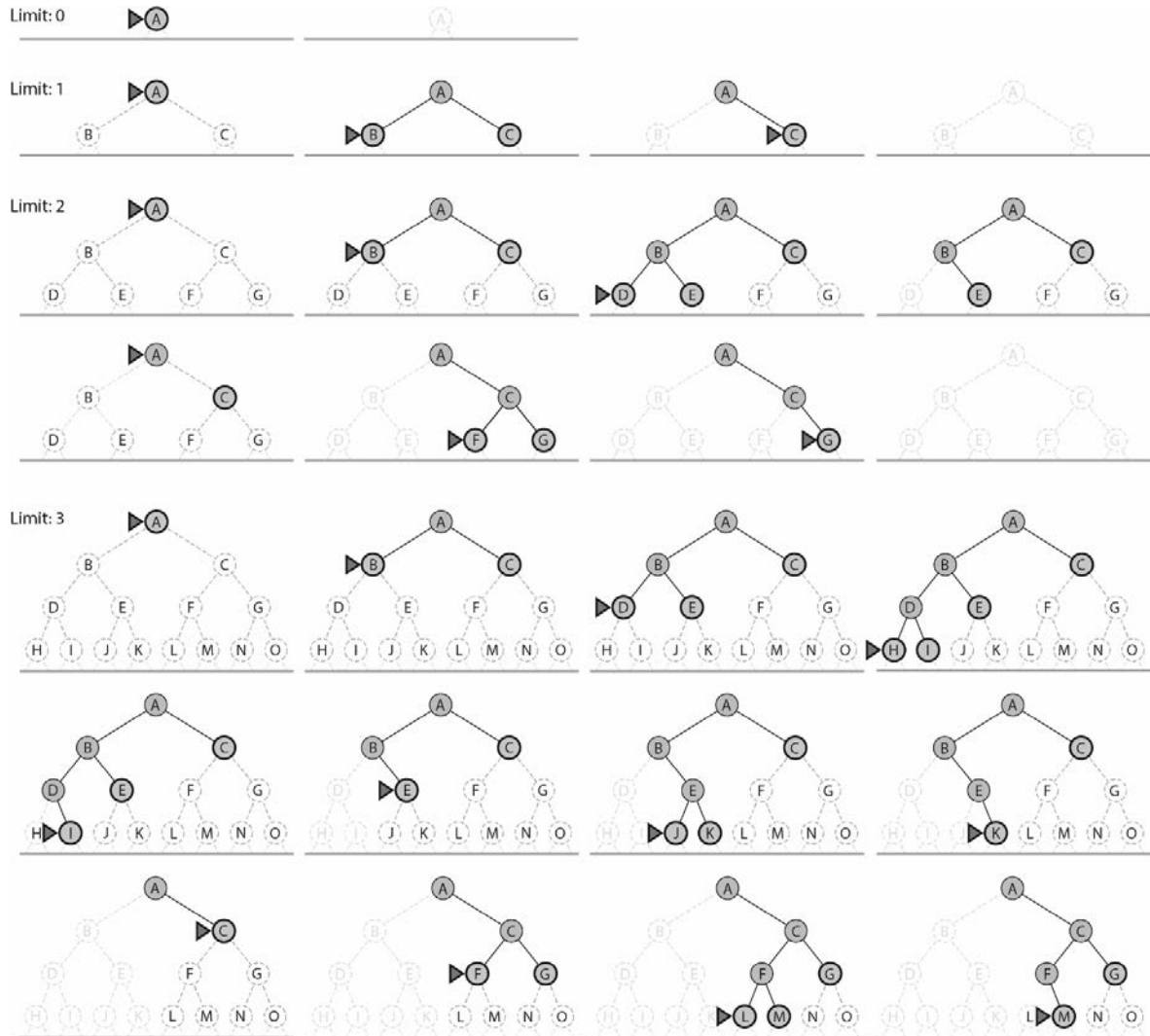
$$N(\text{IDS}) = (d)b^1 + (d-1)b^2 + (d-2)b^3 + \dots + b^d$$

co jest złożonością rzędu $O(b^d)$, asymptotycznie równą tej z wyszukiwania „najpierw najlepszy”, na przykład dla $b = 10$, $d = 5$ mamy¹²

$$N(\text{IDS}) = 50 + 400 + 3\,000 + 20\,000 + 100\,000 = 123\,450$$

$$N(\text{BFS}) = 10 + 100 + 1\,000 + 10\,000 + 100\,000 = 111\,110$$

¹² IDS to skrót od *Iterative Deepening Search* (przeszukiwanie z zagłębianiem), BFS to skrót od *Best-First Search* (wyszukiwanie „najpierw najlepszy”) — *przyp. tłum.*



RYSunEK 3.10. Cztery iteracje wyszukiwania z zagłębieniem w drzewie binarnym, z limitem głębokości zmieniającym się od 0 do 3, stan docelowy reprezentowany jest przez węzeł *M*. Zauważmy, że węzły wewnętrzne tworzą pojedynczą ścieżkę. Trójkąty wskazują węzły do rozwinięcia; zielone węzły z ciemną obwódką tworzą granicę, węzły słabiej widoczne wykluczone są jako rozwiązania na danym poziomie zagłębienia

Jeśli jednak strata czasu na wielokrotne generowanie tych samych węzłów jest prawdziwym problemem, można użyć następującej strategii hybrydowej: uruchamiamy wyszukiwanie wszczepione i kontynuujemy aż do wyczerpania dostępnej pamięci, a następnie przełączamy się na iteracyjne zagłębienie wszystkich węzłów na granicy. *Ogólnie rzecz biorąc, wyszukiwanie z zagłębieniem jest preferowaną metodą wyszukiwania niedoinformowanego w sytuacji, gdy liczebność przestrzeni stanów przekracza możliwości pamięciowe, a głębokość, na jakiej znajduje się węzeł-rozwiązanie, nie jest znana.*

3.4.5. Wyszukiwanie dwukierunkowe

W dotychczas omawianych algorytmach wyszukiwanie rozpoczyna się w stanie początkowym i kończy w jednym ze stanów docelowych. Podejście alternatywne, zwane **wyszukiwaniem dwukierunkowym** (ang. *bidirectional search*), polega na uruchomieniu równoległych wyszukiwań: jednego rozpoczynającego się od stanu początkowego, podążającego w przód, i drugiego rozpoczynającego się od węzła docelowego, podążającego wstecz — w nadziei, że te dwa wyszukiwania spotkają się gdzieś pośrodku drzewa. Motywacja takiego podejścia jest oczywista: złożoność dwukierunkowego wyszukiwania równa $b^{d/2} + b^{d/2} = O(b^{d/2})$ jest znacząco niższa niż $O(b^d)$ w wyszukiwaniu jednokierunkowym, co na przykład dla $b = d = 10$ oznacza skrócenie czasu wykonywania około 50 000 razy.

Aby zaimplementować tę strategię, musimy śledzić dwie granice i utrzymywać dwie tablice stanów osiągniętych, a ponadto zapewnić sobie możliwość „rozumowania wstecz” — jeśli wiemy, że stan s' jest następnikiem stanu s przy wyszukiwaniu w przód, to jednocześnie musimy wiedzieć, że stan s jest następnikiem stanu s' przy wyszukiwaniu wstecz. Rozwiązanie otrzymamy, gdy zderzą się obie granice¹³.

Istnieje wiele różnych wersji wyszukiwania dwukierunkowego, tak samo jak istnieje wiele różnych algorytmów wyszukiwania jednokierunkowego. W tej sekcji opisujemy dwukierunkowy wariant wyszukiwania „najpierw najlepszy”. Chociaż istnieją dwie oddzielne granice, węzeł, który ma zostać rozwinięty jako następny, jest zawsze węzłem z minimalną wartością funkcji ewaluacyjnej, na którejkolwiek z tych granic. Gdy wartością funkcji ewaluacyjnej jest koszt ścieżki, otrzymujemy dwukierunkowy wariant wyszukiwania przy jednolitym koszcie, a jeśli koszt optymalnej ścieżki wynosi C^* , nie będzie rozwijany żaden węzeł o koszcie większym niż $\frac{C^*}{2}$. Możemy dzięki temu uzyskać znaczne przyspieszenie.

Ogólny algorytm wyszukiwania „najpierw najlepszy” w wersji dwukierunkowej widoczny jest na listingu 3.4. Przekazujemy dwie wersje problemu i funkcji ewaluacyjnej, dla wyszukiwania w przód i wyszukiwania wstecz, identyfikowane przyrostkami (odpowiednio) $_F$ i $_B$. Kiedy wynikiem funkcji ewaluacyjnej jest koszt ścieżki, wiemy, że pierwsze znalezione rozwiązanie będzie rozwiązaniem optymalnym, ale nie musi to być prawda w sytuacji, gdy mamy dwie różne funkcje ewaluacyjne. Musimy zatem śledzić najlepsze rozwiązanie znalezione do tej pory i wielokrotnie je aktualizować, dopóki funkcja TERMINATED nie zwróci wartości *prawda* oznaczającej, że nie istnieje już lepsze rozwiązanie.

LISTING 3.4. Dwukierunkowa odmiana wyszukiwania „najpierw najlepszy” wykorzystuje dwie kolejki priorytetowe implementujące granice i dwie tabele przeglądowe rejestrujące stany osiągnięte i reprezentujące je węzły. Gdy ścieżka reprezentowana przez jedną ze wspomnianych granic osiąga stan, który został już osiągnięty na ścieżce reprezentowanej przez drugą granicę, obie ścieżki zostają połączone (za pomocą funkcji JOIN-NODES) w jedno rozwiązanie. Pierwsze otrzymane rozwiązanie niekoniecznie jest rozwiązaniem optymalnym, więc kontynuowane jest poszukiwanie następnych aż do momentu, gdy funkcja TERMINATED zasygnalizuje brak takowych, zwracając wartość *prawda*

```
function BIBF-SEARCH(problem_F, f_F, problem_B, f_B) returns węzeł będący rozwiązaniem
                                     albo sygnalizacja niepowodzenia

    // węzeł reprezentujący stan początkowy
    węzeł_F ← NODE(problem_F.INITIAL)
```

¹³ Przeszukiwanie dwukierunkowe kończy się, gdy spotykają się granice dla obu kierunków (*granica_F* i *granica_B*), czyli gdy okazuje się, że mają one wspólny węzeł. Tak się jednak składa, że kolejka priorytetowa (bo w takiej postaci implementowane są obie granice) nie udostępnia funkcji sprawdzania, czy dany węzeł jest jej elementem; funkcję taką oferuje natomiast tablica przeglądowa rejestrująca stany osiągnięte. Zatem, choć koncepcyjnie testujemy, czy dwie granice mają niepustą część wspólną, w praktyce zmuszeni jesteśmy robić to w sposób okrężny, badając wspomniane tablice przeglądowe. Naszą implementację można rozszerzyć na przypadek wielu węzłów docelowych, należy wówczas uczynić je wszystkimi wartościami początkową kolejki *granica_B* i tablicy *osiągnięte_B*.

```

// węzeł reprezentujący stan końcowy
węzeł_B ← NODE(problem_B.INITIAL)
granica_F ← kolejka priorytetowa z funkcją ewaluacyjną f_F,
           początkowo z węzłem węzeł_F jako jedynym elementem
granica_B ← kolejka priorytetowa z funkcją ewaluacyjną f_B,
           początkowo z węzłem węzeł_B jako jedynym elementem

osiągnięte_F ← tabela przeglądowa, początkowo z jedną pozycją
              (węzeł_F.STATE, węzeł_F)
osiągnięte_B ← tabela przeglądowa, początkowo z jedną pozycją
              (węzeł_B.STATE, węzeł_B)
rozwiązanie ← niepowodzenie

while not TERMINATED(rozwiazanie, granica_F, granica_B) do
  if f_F(TOP(granica_F)) < f_B(TOP(granica_B)) then
    rozwiqzanie ← PROCEED(F, problem_F, granica_F,
                          osiągnięte_F, osiągnięte_B, rozwiązanie)
  else rozwiązanie ← PROCEED(B, problem_B, granica_B,
                              osiągnięte_B, osiągnięte_F, rozwiązanie)
  return rozwiązanie

function PROCEED(kierunek, problem, granica, osiągnięte,
                osiągnięte2, rozwiązanie) returns rozwiązanie
// Rozwinięcie czołowego węzła granicy; konfrontacja z inną granicą w osiągnięte2.
// Zmienna kierunek określa kierunek wyszukiwania:
//   F oznacza wyszukiwanie w przód, B wyszukiwanie wstecz.
węzeł ← POP(granica)
for each potomny in EXPAND(problem, węzeł) do
  s ← potomny.STATE
  if s not in osiągnięte or PATH-COST(potomny) < PATH-
    COST(osiągnięte[s])
  then
    osiągnięte[s] ← potomny
    dodaj potomny do granica
    if s in osiągnięte2 then
      rozwiązanie2 ← JOIN-NODES(kierunek, potomny,
                                osiągnięte2[s]))
      if PATH-COST(rozwiazanie2) < PATH-COST(rozwiazanie)
      then
        rozwiązanie ← rozwiązanie2
return rozwiązanie

```

3.4.6. Porównanie algorytmów wyszukiwania niedoinformowanego

Porównanie algorytmów wyszukiwania niedoinformowanego pod kątem czterech kryteriów omawianych w sekcji 3.3.4 przedstawiamy w tabeli 3.1. Porównanie to dotyczy drzewiastych wersji wyszukiwania, czyli wersji bez sprawdzania powtarzających się stanów. Podstawowa różnica między nimi, a algorytmami w wersji grafowej (które taką kontrolę wykonują) polega na tym, że te ostatnie są zupełne w skończonych przestrzeniach stanów, a ich złożoność — czasowa i pamięciowa — jest proporcjonalna do rozmiaru przestrzeni (mierzonego sumą liczby wierzchołków i krawędzi grafu $|V|+|E|$).

TABELA 3.1. Porównanie podstawowych cech algorytmów wyszukiwania: b oznacza czynnik rozgałęzienia, m jest maksymalną głębokością węzła w drzewie wyszukiwawczym, d jest głębokością najpłytszego rozwiązania, λ jest ograniczeniem głębokości.

Kryterium	Wyszukiwanie wszerek	Wyszukiwanie przy jednolitym koszcie	Wyszukiwanie w głąb	Wyszukiwanie z ograniczeniem głębokości	Iteracyjne zagłębianie	Dwukierunkowość (jeśli możliwa)
Zupełność	Tak ¹	Tak ^{1, 2}	Nie	Nie	Tak ¹	Tak ^{1,4}
Optymalność kosztowa	Tak ³	Tak	Nie	Nie	Tak ³	Tak ^{3,4}
Złożoność czasowa	$O(b^d)$	$O(b^{1+\lceil C^*/\epsilon \rceil})$	$O(b^m)$	$O(b^\lambda)$	$O(b^d)$	$O(b^{d/2})$
Złożoność pamięciowa	$O(b^d)$	$O(b^{1+\lceil C^*/\epsilon \rceil})$	$O(bm)$	$O(b\lambda)$	$O(bd)$	$O(b^{d/2})$

¹ Zupełny, jeśli b jest skończony, a przestrzeń stanów zawiera rozwiązanie lub jest skończona.
² Zupełny, jeśli koszt każdej akcji jest większy od pewnego dodatniego ϵ .
³ Optymalny kosztowo, jeśli wszystkie akcje mają taki sam koszt.
⁴ Jeśli w obu kierunkach prowadzone jest wyszukiwanie wszerek lub wyszukiwanie przy jednolitym koszcie.

3.5. Strategie wyszukiwania poinformowanego (heurystycznego)

W tym podrozdziale pokażemy, jak działa strategia **wyszukiwania poinformowanego** (ang. *informed search*), zwanego także **wyszukiwaniem sterowanym wiedzą** — wykorzystując zewnętrzne, specyficzne dla problemu wskazówki dotyczące lokalizacji celów, może ona znajdować rozwiązania bardziej efektywnie niż strategia niedoinformowana. Wspomniane „wskazówki” mają formę **funkcji heurystycznej**, oznaczanej przez $h(n)$ i mającej następujące znaczenie¹⁴:

$h(n)$ to szacowany koszt najtańszej ścieżki prowadzącej od stanu reprezentowanego przez węzeł n do stanu docelowego.

¹⁴ Może się wydawać dziwne, że argumentem funkcji heurystycznej jest **węzeł**, skoro jej zadaniem jest estymowanie odległości między *stanami*; otóż notacja $h(n)$ zamiast $h(s)$ podyktowana jest tradycją, a konkretnie względami zgodności z funkcjami ewaluacyjnymi $f(n)$ i funkcjami kosztu ścieżki $g(n)$.

Na przykład, w problemie znajdowania trasy odległość między stanem bieżącym a stanem docelowym może być przybliżana przez *odległość w linii prostej* między punktami na mapie reprezentującymi miasta odpowiadające tym stanom. Heurystykami i ich genezą zajmimy się dokładniej w podrozdziale 3.6.

3.5.1. Zachłanne wyszukiwanie „najpierw najlepszy”

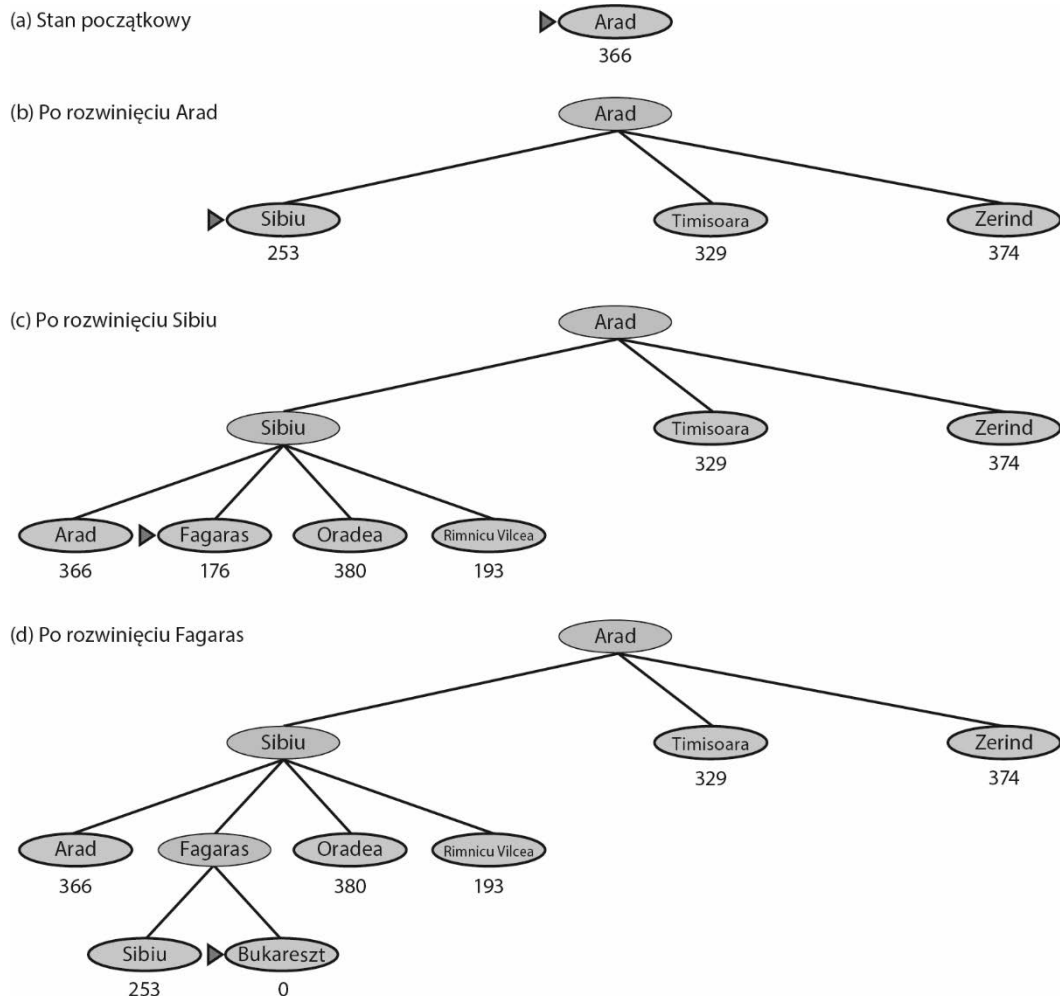
Algorytm wyszukiwania „najpierw najlepszy” w wersji **zachłannej** (ang. *greedy*) wybiera do rozwijania ten węzeł granicy, który wydaje się być położony najbliżej (w sensie funkcji $h(n)$) węzła docelowego — bo prawdopodobnie taka strategia powinna szybko doprowadzić do rozwiązania; innymi słowy, funkcja h jest w tym przypadku funkcją ewaluacyjną ($f(n) = h(n)$).

Zobaczmy, jak strategia ta sprawdza się w przypadku problemu poszukiwania drogi w Rumunii, gdzie jako wartości funkcji heurystycznej przyjmujemy odległość węzła od celu w linii prostej, dlatego oznaczmy tę funkcję przez h_{SLD} , od ang. *straight-line distance*. Wartości tej funkcji dla poszczególnych miast — czyli ich odległości w linii prostej (w milach) od Bukaresztu — zebrane są w tabeli 3.2, na przykład $h_{\text{SLD}}(\text{Arad}) = 366$. Zauważmy, że wartości tych nie sposób wydedukować ze specyfikacji problemu, czyli nie można ich wyliczyć przy użyciu funkcji ACTIONS i RESULT. Zauważmy także, iż użyteczność takiej właśnie heurystyki wynika z faktu, że odległość w linii prostej między miastami skorelowana jest ściśle z ich odległością liczoną wzdłuż dostępnej drogi.

TABELA 3.2. Wartości funkcji h_{SLD} — odległości poszczególnych miast od Bukaresztu, w linii prostej

Arad	366	Mehadia	241
Bukareszt	0	Neamt	234
Craiova	160	Oradea	380
Drobeta	242	Pitesti	100
Eforie	161	Rimnicu Vilcea	193
Fagaras	176	Sibiu	253
Giurgiu	77	Timisoara	329
Hirsova	151	Urziceni	80
Iasi	226	Vaslui	199
Lugoj	244	Zerind	374

Postęp opisanego wyszukiwania zilustrowany jest na rysunku 3.11. Węzłem początkowym (korzeniem) jest *Arad*; jego węzłem potomnym, który zostanie rozwinięty jako pierwszy, jest *Sibiu*, ponieważ (w świetle przyjętej heurystyki h_{SLD}) znajduje się on bliżej Bukaresztu niż węzły *Zerind* i *Timisoara*. Przechodzimy więc do węzła *Sibiu* i (zgodnie z heurystyką) wybieramy do rozwinięcia węzeł *Fagaras*. Na tej samej zasadzie następnym do rozwinięcia jest docelowy węzeł *Bukareszt*. W tym konkretnym przypadku rozwiązanie znalezione zostało bez potrzeby jakiegokolwiek węzła nieznajdującego się na ścieżce będącej rozwiązaniem. Znaleziona trasa *Arad* – *Sibiu* – *Fagaras* – *Bukareszt* nie jest jednak optymalna, bo trasa *Arad* – *Sibiu* – *Rimnicu Vilcea* – *Pitesti* – *Bukareszt* jest od niej krótsza o 32 mile. Przymiotnik „zachłanny” dla opisaney strategii oznacza, że w każdym kroku stara się ona zbliżyć jak najbardziej do celu; niniejszy przykład jest dowodem na to, że zachłanność nie zawsze popłaca, czyli nie zawsze prowadzi do rozwiązania optymalnego.



RYSUNEK 3.11. Kolejne etapy znajdowania drogi z Arad do Bukaresztu metodą zachłannego wyszukiwania drzewiastego „najpierw najlepszy” przy zastosowaniu heurystyki h_{SLD} . Dla poszczególnych węzłów podano wartości tej heurystyki

Wyszukiwanie grafowe „najpierw najlepszy” w wersji zachłannej jest algorytmem zupełnym w skończonych przestrzeniach stanów, ale nie w nieskończonych. Jego złożoność czasowa i pamięciowa wynosi w najgorszym przypadku $O(|V|)$; jednak przy wyborze odpowiedniej heurystyki można ją znacząco zredukować, w przypadku niektórych problemów nawet do $O(bm)$.

3.5.2. Wyszukiwanie A*

Najbardziej rozpowszechnionym algorytmem wyszukiwania poinformowanego jest **A*** („A z gwiazdką”, ang. *A-star*). Jest to w istocie wyszukiwanie „najpierw najlepszy” z funkcją ewaluacyjną określoną jako

$$f(n) = g(n) + h(n)$$

gdzie $g(n)$ jest kosztem ścieżki prowadzącej od węzła początkowego do węzła n , natomiast $h(n)$ oznacza szacowany koszt najkrótszej ścieżki prowadzącej od węzła n do węzła docelowego. W konsekwencji mamy więc

$f(n)$ = szacowany koszt najlepszej ścieżki biegnącej od węzła początkowego, przez węzeł n , do węzła docelowego.

Na rysunku 3.12 pokazujemy postępowanie wyszukiwania A* jako znajdowania ścieżki z Arad do Bukaresztu. Wartości funkcji $g(n)$ można łatwo obliczyć na podstawie kosztów akcji podanych na rysunku 3.1, wartości funkcji $h(n)$ podano explicite w tabeli 3.2. Zauważmy, że *Bukareszt* po raz pierwszy pojawia się na granicy w kroku (e), ale nie jest wybierany do rozwinięcia (i tym samym nie jest brany pod uwagę jako rozwiązanie), ponieważ przy wartości $f = 450$ nie jest to najtańszy węzeł w granicy — powinien być nim *Pitesti* ($f = 417$). Można to zinterpretować w ten sposób, że może istnieć rozwiązanie prowadzące przez *Pitesti*, kosztujące niewiele więcej niż 417, więc algorytm nie zadowolony się rozwiązaniem, którego koszt wynosi 450. Gdy w kroku (f) rozwiniemy węzeł *Pitesti*, najtańszym z wygenerowanych węzłów potomnych jest *Bukareszt* ($f = 418$), jest więc wykrywany i zwracany jako rozwiązanie optymalne.

Wyszukiwanie A* jest algorytmem zupełnym¹⁵. To, czy jest ono optymalne kosztowo, zależy od pewnych własności przyjętej heurystyki, z których najważniejszą jest **dopuszczalność** (ang. *admissibility*): **heurystyka dopuszczalna** to taka, która nigdy nie przeszacowuje kosztu osiągnięcia celu, czyli jest heurystyką *optymistyczną*. Wyszukiwanie A* z heurystyką dopuszczalną jest optymalne kosztowo, co możemy udowodnić metodą nie wprost, czyli przez sprowadzenie do sprzeczności. Załóżmy mianowicie, że koszt optymalnej ścieżki wynosi C^* , ale algorytm zwraca jako rozwiązanie ścieżkę o koszcie $C > C^*$. Aby tak się stało, musi na optymalnej ścieżce istnieć jakiś *nierozwinięty* węzeł n — gdyby wszystkie węzły na ścieżce optymalnej zostały rozwinięte, algorytm właśnie ją zwróciłby jako rozwiązanie. Oznaczając przez $g^*(n)$ koszt odcinka ścieżki optymalnej od węzła początkowego do węzła n , a przez $h^*(n)$ koszt odcinka ścieżki optymalnej od węzła n do najbliższego węzła docelowego, otrzymujemy:

$f(n) > C^*$ (bo w przeciwnym razie węzeł n musiałby być rozwinięty)

$f(n) = g(n) + h(n)$ (z definicji)

$f(n) = g^*(n) + h(n)$ (bo n znajduje się na optymalnej ścieżce)

$f(n) \leq g^*(n) + h^*(n)$ (bo dla heurystyki dopuszczalnej zachodzi $h(n) \leq h^*(n)$)

$f(n) \leq C^*$ (bo z definicji $C^* = g^*(n) + h^*(n)$)

Jak widać, pierwsza i ostatnia z powyższych równości przeczą sobie nawzajem, czego powodem jest przyjęcie założenia, że algorytm wyszukiwania A* z heurystyką dopuszczalną może zwrócić nieoptymalną ścieżkę jako rozwiązanie.

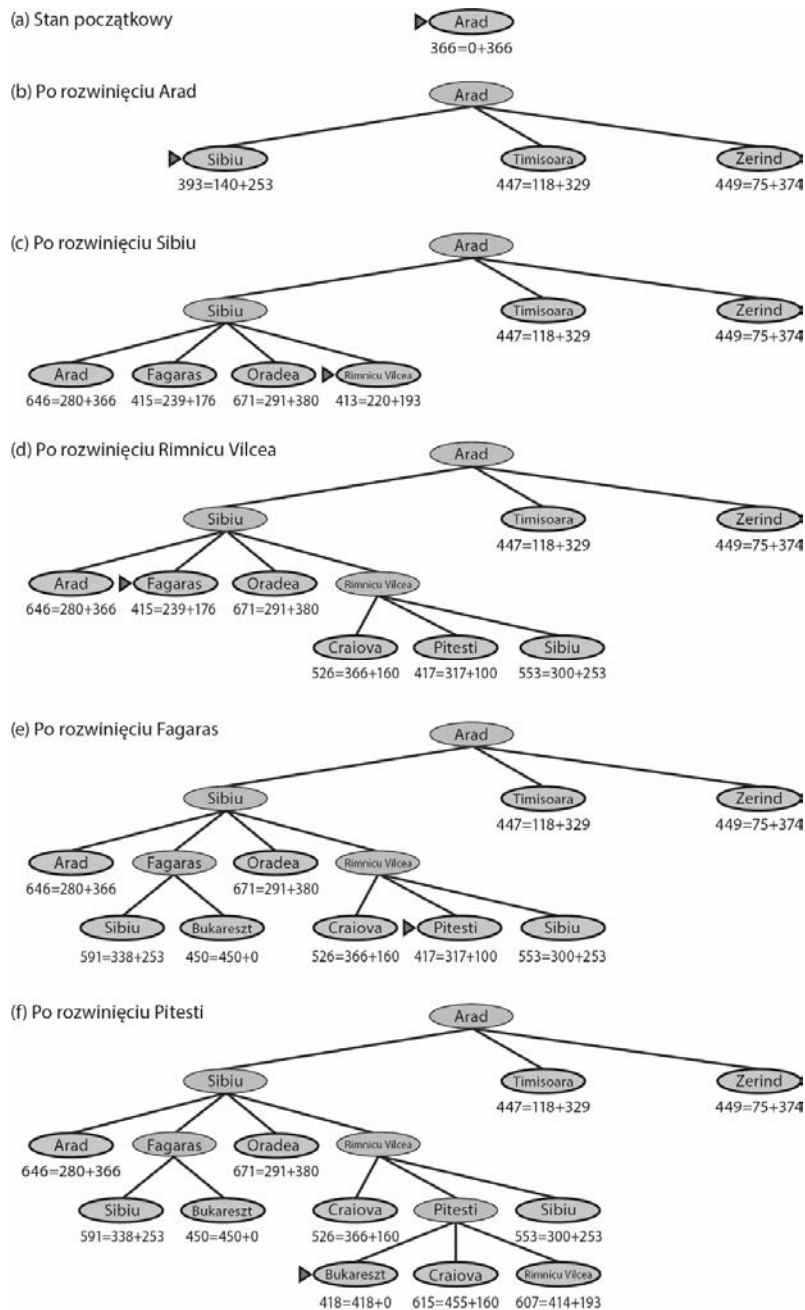
Drugą istotną, nieco silniejszą własnością heurystyki jest jej **spójność** (ang. *consistency*): heurystyka jest **spójna**, jeżeli dla każdego węzła n i każdego jego następnika n' generowanego przez akcję a zachodzi

$$h(n) \leq c(n, a, n') + h(n')$$

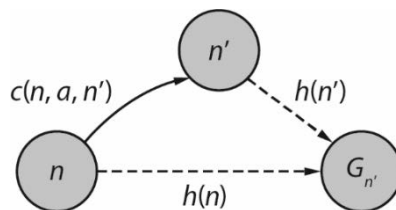
(powyższy warunek jest przykładem **nierówności trójkąta**, na mocy której długość boku trójkąta nie może być większa od sumy długości dwóch pozostałych jego boków¹⁶ (co ilustruje rysunek 3.13). Definiowana wcześniej heurystyka h_{SLD} jest przykładem spójnej heurystyki.

¹⁵ Ponownie zakładamy, że koszt każdej akcji jest większy od pewnego dodatniego ϵ , a przestrzeń stanów jest skończona lub zawiera rozwiązanie.

¹⁶ Klasyczna nierówność trójkąta jest nierównością *ostrą* — długość dowolnego boku jest *mniejsza* od sumy długości dwóch pozostałych; jeśli dopuścimy nierówność nieostrą, zgadzamy się traktować jako „trójkąty” układy trzech współliniowych punktów — *przypr. tłum.*



RYСУNEK 3.12. Kolejne etapy wyszukiwania A* znajdującego drogę do Bukaresztu. Dla poszczególnych węzłów podano wartości $f = g + h$, gdzie h jest odległością w linii prostej od celu, odczytaną z tabeli 3.2



RYSUNEK 3.13. Nierówność trójkąta ilustrująca spójność heurystyki: dla węzła n i jego następnika n' generowanego przez akcję a wartość $h(n)$ nie może być większa od sumy $h(n')$ i kosztu akcji a

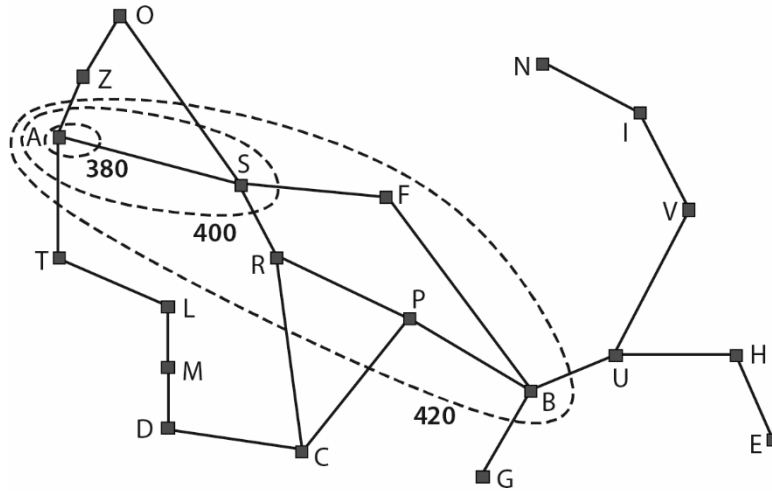
Każda spójna heurystyka jest dopuszczalna (ale nie odwrotnie), więc przy spójnej heurystyce wyszukiwanie A^* jest optymalne kosztowo. Ponadto, przy spójnej heurystyce, gdy po raz pierwszy osiągniemy jakiś stan położony na optymalnej ścieżce, nigdy nie będziemy ponownie dodawać tego stanu do granicy ani zmieniać odpowiadającej mu pozycji w tabeli przeglądowej stanów osiągniętych. Nie jest to prawdą w przypadku heurystyki niespójnej — do danego stanu może prowadzić wiele ścieżek, wskutek czego każdorazowo po napotkaniu ścieżki prowadzącej do tego stanu, tańszej niż poprzednia, będziemy ponownie dodawać do granicy węzeł reprezentujący ten stan i modyfikować wspomnianą wcześniej tabelę przeglądową, co wymaga dodatkowego czasu i pamięci. Z tego względu niektóre implementacje A^* unikają, w przypadku znalezienia tańszej ścieżki, powtórnego dodawania tego samego stanu do granicy, w zamian modyfikując jego węzły potomne — co oczywiście wymaga utrzymywania, w strukturach reprezentujących węzły, wskaźników zarówno do węzłów potomnych, jak i węzłów macierzystych. Komplikacje te skłaniają autorów algorytmów A^* do unikania niespójnych heurystyk, choć Falner i inni (2011) starają się rozwiewać obawy w tym względzie, twierdząc że opisane niepożądane efekty raczej rzadko zdarzają się w praktyce.

Jeśli heurystyka nie jest dopuszczalna, algorytm A^* może nie być optymalny kosztowo, ale jest taki przy spełnieniu pewnych warunków. Po pierwsze, jeśli istnieje choćby jedna ścieżka optymalna kosztowo, i dla każdego węzła n na tej ścieżce wartość $h(n)$ nie jest zawyżona, to ścieżka ta zostanie zwrócona jako rozwiązanie, bez względu na wartość h dla węzłów leżących poza tą ścieżką. Po drugie, jeśli optymalna ścieżka ma koszt C^* i znaleziona zostanie inna ścieżka o koszcie $C_2 > C^*$, to nawet jeśli heurystyka zawyża koszty, ale nie więcej niż o $C_2 - C^*$, istnieje gwarancja, że zwrócone zostanie rozwiązanie o koszcie C^* .

3.5.3. Kontury wyszukiwania

Użytecznym sposobem wizualizacji wyszukiwania są **kontury** naniesione na graf przestrzeni stanów, stanowiące analogię poziomom (warstwom) na mapie topograficznej. Podstawa podziału grafu na kontury są wartości funkcji f — wewnątrz konturu etykietowanego wartością w znajdują się wszystkie wierzchołki v , dla których $f(v) = g(v) + h(v) \leq w$. Na rysunku 3.14 widoczne są trzy kontury dla $w = 380, 400$ i 420 . Ponieważ algorytm A^* wybiera węzły do rozwinięcia w kolejności rosnących kosztów, budowę drzewa wyszukiwawczego można rozpatrywać jako budowanie koncentrycznych kręgów o zwiększającej się wartości w .

Koncepcja konturów ma zastosowanie także do wyszukiwania przy jednolitym koszcie, jednak podstawą etykietowania konturów jest wartość funkcji g , nie $g+h$. W przeciwieństwie do wyszukiwania A^* , gdzie kontury wykazują tendencję do rozszerzania się w kierunku celu, w wyszukiwaniu przy jednolitym koszcie rozszerzanie konturów ma charakter izotropowy — koncentryczne okręgi rozszerzają się równomiernie we wszystkich kierunkach. Przy dobrej heurystyce kontury w wyszukiwaniu A^* mogą przybierać kształty wydłużonych elips koncentrujących się wokół optymalnej ścieżki.



RYСУNEK 3.14. Mapa Rumunii z rysunku 3.1 z uwidocznionymi konturami dla $f = 380$, $f = 400$ i $f = 420$, z węzłem początkowym Arad. Dla węzłów znajdujących się wewnątrz konturu etykietowanego wartością w zachodzi zależność: $f+g \leq w$

Ponieważ koszty wszystkich akcji są dodatnie, to oczywiście wartość funkcji g rośnie w miarę wydłużania ścieżki — funkcja g jest więc **monotoniczna**¹⁷. Linie poszczególnych konturów nie przecinają się więc i zawsze można wstawić linię konturu między dwa dowolne węzły na dowolnej ścieżce (pod warunkiem rysowania konturów wystarczająco cienką kreską).

Nie jest natomiast oczywiste, czy monotoniczna jest także funkcja $f = g + h$. Gdy przedłużamy ścieżkę od węzła n do węzła n' , koszt tej ścieżki zmienia się z $g(n) + h(n)$ na $g(n) + c(n, a, n') + h(n')$, czyli jest monotoniczny (niemalejący) wtedy i tylko wtedy, gdy $h(n) \leq c(n, a, n') + h(n')$, czyli gdy heurystyka h jest spójna¹⁸. Należy jednak uwzględnić fakt, że na ścieżce może pojawić się kilka kolejnych węzłów z identyczną wartością $g(n) + h(n)$ — jest tak wtedy, gdy koszt właśnie podjętej akcji równa się spadkowi funkcji f przy przejściu do nowego węzła. We wcześniejszym problemie z agentem poruszającym się po gridzie jest tak w sytuacji, gdy komórka docelowa znajduje się w tym samym rzędzie co komórka, w której aktualnie znajduje się agent i agent ten robi jeden krok w kierunku celu — funkcja g zwiększa się wówczas o 1 i o tyle samo zmniejsza się funkcja h . Jeśli C^* jest kosztem optymalnego rozwiązania, możemy stwierdzić, że

- A^* rozwija wszystkie węzły, które mogą zostać osiągnięte ze stanu początkowego po ścieżce, na której dla każdego węzła n zachodzi $f(n) < C^*$; takie węzły nazywamy **niewątpliwie rozwiniętymi** (ang. *surely expanded*).
- A^* może wobec tego rozwijać niektóre z węzłów znajdujących się w „konturze docelowym” (czyli węzłów, dla których $f = C^*$) przed wybraniem węzła docelowego.
- A^* nie rozwija węzłów, dla których $f > C^*$.

¹⁷ W zasadzie powinniśmy napisać „ściśle monotoniczna” w odniesieniu do funkcji rosnącej, bo „monotoniczny” oznacza tyle, co „niemalejący”, czyli mogący również zachowywać stałą wartość.

¹⁸ Monotoniczność funkcji h ma więc ścisły związek ze spójnością heurystyki i w rezultacie synonimem „heurystyki spójnej” jest „heurystyka monotoniczna”, choć oba te terminy zostały sformułowane niezależnie od siebie (Pearl, 1984).

Wyszukiwanie A^* ze spójną heurystyką jest **optymalnie wydajne** w tym sensie, że każdy algorytm, który rozszerza ścieżki wyszukiwania ze stanu początkowego i używa tej samej heurystyki, musi rozwinąć wszystkie węzły, które są niewątpliwie rozwijane przez A^* (ponieważ każdy tych węzłów mógłby być częścią optymalnego rozwiązania). Jeśli chodzi o węzły, dla których $f = C^*$, to jeden algorytm może szczęśliwym trafem wybrać węzeł położony na optymalnej ścieżce, inny natomiast może mieć mniej szczęścia i w pierwszej kolejności wybrać inny — ta różnica jest nieistotna z perspektywy definicji optymalnej wydajności.

Wyszukiwanie A^* jest wydajne, ponieważ **przycina** (ang. *prunes*) drzewo wyszukiwawcze, czyli eliminuje z niego węzły, które nie są konieczne do znalezienia rozwiązania. Na rysunku 3.12 w części (b) widzimy węzły Timisoara ($f = 447$) i Zerind ($f = 449$). Są one potomkami korzenia i byłyby rozwijane w pierwszej kolejności w algorytmie z jednolitym kosztem i algorytmie wyszukiwania wszerek, ale algorytm A^* nie rozwinie ich nigdy, ponieważ wcześniej znajdzie rozwiązanie o koszcie $f = 418$.

Koncepcja przycinania — czyli eliminowania a priori pewnych możliwości, bez ich sprawdzania, jako nieistotnych dla rozwiązania — jest ważna w wielu obszarach sztucznej inteligencji.

To, że wyszukiwanie A^* jest zupełne, optymalne kosztowo i optymalnie wydajne spośród wszystkich takich algorytmów, jest zdecydowanie dobrą wiadomością. Niestety, nie oznacza to wcale, że stanowi ono odpowiedź na wszelkie problemy wyszukiwania: złą wiadomością jest to, że w wielu problemach liczba rozwiniętych węzłów rośnie wykładniczo do długości ścieżki będącej rozwiązaniem. Powróćmy do (wielokrotnie cytowanego) przykładu z odkurzaczem i założmy tym razem, że odkurzacz ten, znajdując się w jakimś kwadracie, ma możliwość odkurzenia *dowolnego kwadratu, bez przemieszczania się do niego*; koszt takiej akcji równy jest 1. Wobec początkowo zakurzonych N kwadratów możliwych jest 2^N stanów, z których każdy charakteryzuje się pewnym podzbiorem kwadratów już odkurzonych. Każdy z tych stanów znajduje się na ścieżce stanowiącej optymalne rozwiązanie, a ponieważ dla każdego z nich mamy $f < C^*$, wszystkie zostaną przeanalizowane w ramach A^* .

3.5.4. Wyszukiwanie satysfakcjonujące: heurystyki niedopuszczalne i ważne A^*

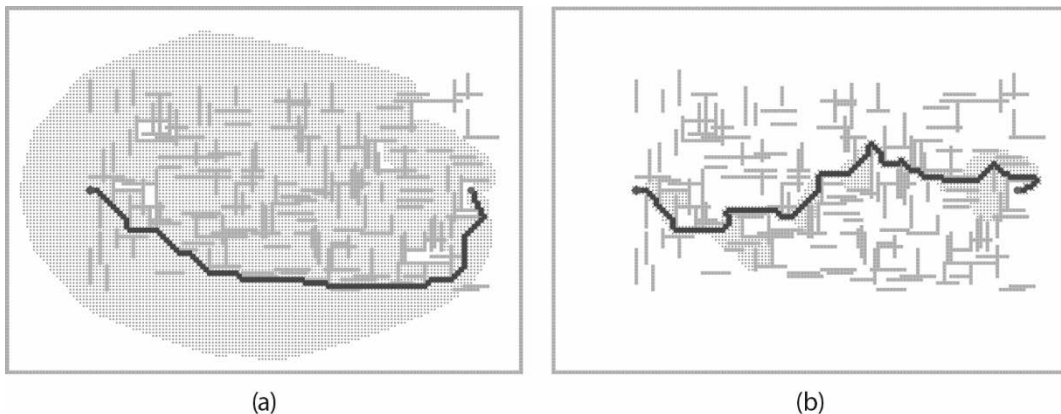
Wyszukiwanie A^* ma wiele zalet, ale jego głównym mankamentem jest duża liczba rozwijanych węzłów (dla niektórych problemów). Możemy zredukować liczbę rozwijanych węzłów, i tym samym zaoszczędzić czas i pamięć, jeżeli zrezygnujemy z poszukiwania rozwiązania absolutnie optymalnego i zadowolimy się rozwiązaniem suboptymalnym, które choć nie najlepsze, okazuje się „wystarczająco dobre” w konkretnym zastosowaniu. Takie suboptymalne rozwiązania określamy mianem **satysfakcjonujących** lub **zadowalających**. Jeżeli w wyszukiwaniu A^* użyjemy heurystyki, która **nie jest dopuszczalna**, — czyli może zawyżać oszacowania kosztu — ryzykujemy „zgubienie” przez algorytm optymalnego rozwiązania, jednakże heurystyka taka może być potencjalnie dokładniejsza dla danego problemu, co samo z siebie może powodować zredukowanie liczby rozwijanych węzłów. W inżynierii drogowej znane jest pojęcie **współczynnika objazdu** (ang. *detour index*) — jest to stosunek najkrótszej trasy między dwiema lokalizacjami do odległości między nimi w linii prostej; współczynnik ten odzwierciedla stopień „zakrzywienia dróg” w danym regionie. Jeżeli na przykład w danym regionie współczynnik ten wynosi 1,3, a według mapy odległość w linii prostej między dwoma miastami wynosi 10 mil, to można oczekiwać, że jadąc najkrótszą trasą między tymi miastami, trzeba będzie przejechać około 13 mil. Dla większości dróg współczynnik ten waha się w granicach od 1,2 do 1,6.

Tę ideę można uogólnić na dowolny problem — niekoniecznie związany z drogownictwem — w postaci podejścia zwanego **ważonym wyszukiwaniem A^*** (ang. *weighted A^* search*). Jego istotą jest różnicowanie wkładu funkcji g i h do funkcji f , według formuły

$$f(n) = g(n) + W \times h(n)$$

dla pewnego $W > 1$, zwanego *wagą heurystyki*.

Rysunek 3.15 przedstawia problem wyszukiwania w gridzie. W części (a) widać, jak algorytm A^* znajduje optymalne rozwiązanie, badając w tym celu znaczną liczbę węzłów. W części (b) do rozwiązania tego samego problemu użyto ważonego algorytmu A^* ; otrzymane rozwiązanie ma co prawda koszt trochę większy niż to z części (a), ale znalezione zostało znacznie szybciej. Nietrudno zauważyć, że w wyszukiwaniu ważonym kontury osiągniętych celów rozciągnięte są wyraźnie w kierunku celu. Oznacza to, że analizowanych jest znacznie mniej stanów, ale jeśli optymalna ścieżka kiedykolwiek zboczy poza kontury wyszukiwania ważonego (jak w tym przypadku), otrzymane rozwiązanie może nie być rozwiązaniem optymalnym. Ogólnie, jeśli koszt rozwiązania optymalnego wynosi C^* , ważne wyszukiwanie znajduje rozwiązanie o koszcie pomiędzy C^* a $W \times C^*$, przy czym w praktyce jest ono położone znacznie bliżej C^* niż $W \times C^*$.



RYСУNEK 3.15. Rezultaty dwóch wyszukiwań na tym samym gridzie: (a) klasyczne wyszukiwanie A* (b) ważone wyszukiwanie A* z wagą $W = 2$. Szare paski to przeszkody między komórkami, purpurowa linia to ścieżka od stanu początkowego (zielony punkt) do stanu docelowego (czerwony punkt). Małe kropki reprezentują stany osiągnięte w procesie wyszukiwania. W tym konkretnym przypadku ważone wyszukiwanie A* przeanalizowało 7-krotnie mniej stanów i zwróciło rozwiązanie droższe o 5% w porównaniu z klasycznym wyszukiwaniem A*.

Wnikliwi Czytelnicy zauważyli zapewne, że ważone wyszukiwanie A^* jest algorytmem dość ogólnym i niektóre inne opisywane w tym rozdziale algorytmy można uważać za szczególne przypadki kombinacji $f = g + W \times h$, dla różnych wartości W , jak w tabeli 3.3.

TABELA 3.3. Algorytmy wyszukiwania jako szczególne przypadki ważonego wyszukiwania A*

Algorytm	Funkcja ewaluacyjna	Waga
Klasyczne wyszukiwanie A*	$g(n)+h(n)$	$W = 1$
Wyszukiwanie przy jednolitym koszcie	$g(n)$	$W = 0$
Zachłanne wyszukiwanie „najpierw najlepszy”	$h(n)$	$W = \infty$
Ważne wyszukiwanie A*	$g(n)+W \times h(n)$	$1 < W < \infty$

Ważone wyszukiwanie A^* ma w sobie element zachłanności — przypomina zachłanny algorytm „najpierw najlepszy” w tym sensie, że koncentruje się na podążaniu do celu; nie ignoruje jednak całkowicie kosztu ścieżki i może zawiesić podążanie nią, jeśli zbyt mały postęp okupiony jest zbyt dużymi kosztami.

Istnieje wiele suboptymalnych algorytmów wyszukiwania, które można określić jako „wystarczająco dobre” i scharakteryzować za pomocą kryteriów określających, co jest uważane za „wystarczająco dobre” w świetle różnych kryteriów. **Wyszukiwanie z ograniczoną nieoptymalnością** (ang. *bounded suboptimal search*) daje gwarancję, że koszt zwróconego rozwiązania będzie wyższy od kosztu optymalnego rozwiązania o nie więcej niż stały czynnik ($c \leq W \cdot C^*$ dla ustalonego $W \geq 1$); wyszukiwanie A^* jest właśnie algorytmem tej kategorii. Po **wyszukiwaniu z ograniczeniami kosztu** (ang. *bounded-cost search*) spodziewamy się otrzymać rozwiązanie, którego koszt jest nie większy od ustalonego limitu C . Wreszcie, gdy czynnikiem krytycznym jest czas znajdowania rozwiązania, akceptowalny staje się każdy algorytm **wyszukiwania bez ograniczenia kosztu** (ang. *inbounded-cost search*), o ile gwarantuje zwrócenie rozwiązania w krótkim czasie.

Przykładem algorytmu tej ostatniej kategorii jest **szybkie wyszukiwanie** (ang. *speedy search*), stanowiące odmianę zachłannego wyszukiwania „najpierw najlepszy”; w tym przypadku funkcja heurystyczna dla danego węzła zwraca szacowaną *liczbę akcji* niezbędnych do osiągnięcia celu z poziomu tegoż węzła, niezależnie od kosztu tych akcji. W razie wystąpienia problemów, w których wszystkie akcje mają identyczny koszt, algorytm ten jest identyczny z oryginalną wersją zachłannego wyszukiwania „najpierw najlepszy”, w której jako wartość heurystyki przyjmuje się *koszt* osiągnięcia celu.

3.5.5. Wyszukiwanie z ograniczeniami pamięciowymi

Podstawowym mankamentem wyszukiwania A^* jest jego pamięciożerność. W tej sekcji omówimy zarówno kilka trików implementacyjnych, dzięki którym udaje się zmniejszyć zapotrzebowanie na pamięć, jak i kilka zupełnie nowych algorytmów, które z natury korzystają z pamięci bardziej oszczędnie.

Główne źródła zapotrzebowania na pamięć to kolejka przechowująca węzły tworzące *granice* oraz tabela przeglądowa ewidencjonująca stany *osiągnięte*. W przedstawionej przez nas implementacji wyszukiwania „najpierw najlepszy” (patrz listing 3.1) każdy węzeł wchodzący w skład granicy ewidencjonowany jest dwukrotnie: raz w kolejce implementującej tę granicę (dzięki której możemy wybierać węzły do rozwijania) i drugi raz w tablicy stanów osiągniętych (dzięki której wiemy, czy już wcześniej analizowaliśmy dany stan). W wielu przypadkach (takich, jak problem eksploatacji gridu) dublowanie to nie jest problemem, bo rozmiar *granicy* jest znacznie mniejszy niż rozmiar tablicy stanów *osiągniętych*; w innych przypadkach pożądane jest zredukowanie (w miarę możliwości) zapotrzebowania implementacji algorytmu na pamięć, za pomocą wspomnianych na wstępie trików — ze świadomością, iż skomplikuje to kod i (prawdopodobnie) wydłuży czas poszukiwania rozwiązania.

Jeden z takich trików polega na „przerzedzaniu” tablicy stanów osiągniętych, czyli usuwaniu z niej węzłów, co do których mamy pewność, że nie są już potrzebne; zwolnione w ten sposób pozycje mogą być ponownie wykorzystywane. W przypadku niektórych problemów możemy wykorzystywać własność separacji (patrz rysunek 3.6) w połączeniu z zakazem wykonywania akcji powodujących zawracanie na ścieżce (ang. *U-turn*) w celu zapewnienia, że każda akcja prowadzi bądź to na zewnątrz granicy, bądź do innego stanu tworzącego granicę. Możemy wówczas ograniczyć się do kontrolowania granicy na okoliczność redundantnych ścieżek i tablica ewidencjonująca stany *osiągnięte* nie będzie nam potrzebna.

W przypadku wielu problemów może być użyteczne kontrolowanie **liczby odwołań** do każdego osiągniętego stanu i usuwanie danego stanu z tabeli przeglądowej, gdy liczba ta *osiągnie* limit wynikający z natury problemu. Na przykład w prostokątnym gridzie każda komórka może być osiągnięta wyłącznie wskutek przemieszczenia się do niej z jednej z czterech komórek sąsiednich, więc do stanu odpowiadającego danej komórce algorytm może się odwołać co najwyżej cztery razy; przy czwartym odwołaniu można ów stan usunąć z tabeli przeglądowej.

Przeanalizujmy zatem pokrótce nowe algorytmy zaprojektowane z myślą o oszczędniejszym wykorzystaniu pamięci.

Istotą **wyszukiwania skupionego**, zwanego także **wyszukiwaniem wiązkowym** (ang. *beam search*), jest redukcja zapotrzebowania na pamięć poprzez ograniczanie liczebności granic — zamiast utrzymywać komplet węzłów stanowiących wynik rozwinięcia, ograniczamy się tylko do ich podzbioru („wiązki”) złożonego z k najlepszych, czyli tych, dla których funkcja ewaluacyjna f zwraca najmniejsze wartości. Wskutek tego zabiegu algorytm traci co prawda własność zupełności i optymalności kosztowej, ale za to zmniejsza się jego zapotrzebowanie na pamięć (co może nawet warunkować samą możliwość jego uruchomienia) i oczywiście skraca się czas jego wykonania, bo analizowane jest znacznie mniej stanów. Wybór odpowiedniej wartości k jest swego rodzaju sztuką kompromisu między wymaganiami pamięciowymi a jakością rozwiązania; dla wielu problemów możliwe jest uzyskiwanie rozwiązań bliskich optymalnym, przy jednoczesnej znaczącej redukcji zapotrzebowania na pamięć. Powracając do koncepcji konturów wyszukiwania z sekcji 3.5.3: można sobie wyobrazić wyszukiwanie A^* przy jednolitym koszcie jako promieniowanie (ang. *beaming*) wszystkich węzłów objętych konturem na zewnątrz tego konturu, a wyszukiwanie skupione — jak promieniowanie tylko „wiązki” k najlepszych kandydatów.

W odmianie tego algorytmu liczebność granicy (k) nie jest sztywno ustalona a priori, wiązkę tworzą te węzły, dla których funkcja ewaluacyjna zwraca wartość różniącą się o nie więcej niż ustalone δ od wartości minimalnej funkcji f . Liczebność podzbioru rozwijanych węzłów zależy zatem od „polaryzacji” granicy, czyli liczby węzłów wyraźnie odróżniających się (w sensie funkcji f) od reszty; nawet gdy liczebność ta jest duża, może gwałtownie zmaleć, gdy pojawi się odpowiednio „silny” węzeł.

Wyszukiwanie A^* z iteracyjnym zagłębianiem (ang. IDA^* — *Iterative-deepening A^* search*) ma się tak do klasycznego wyszukiwania A^* , jak iteracyjne zagłębianie do wyszukiwania w głąb: korzyści IDA^* biorą się z braku konieczności przechowywania w pamięci wszystkich osiągniętych stanów, za cenę wielokrotnego analizowania niektórych stanów. Algorytm ten ma duże znaczenie, i jest powszechnie wykorzystywany, w rozwiązywaniu problemów, których wymagania pamięciowe wykluczają użycie innych algorytmów.

W klasycznym zagłębianiu iteracyjnym limit głębokości powiększany jest o jeden w każdej iteracji; w algorytmie IDA^* limit głębokości w kolejnej iteracji wyznaczony zostaje przez węzeł, dla którego wartość funkcji f (równa sumie wartości funkcji g i h) jest najmniejsza i jednocześnie większa od tej wyznaczającej limit w poprzedniej iteracji. Innymi słowy, algorytm IDA^* w każdej iteracji dokonuje wyczerpującego przeszukiwania f -konturu, znajduje zewnętrzny węzeł położony najbliżej tego konturu i traktuje funkcję f dla tego węzła jako wartość nowego konturu. W przypadku problemów takich jak puzzle-8 (patrz rysunek 3.3), w których funkcja f przyjmuje tylko wartości całkowite, strategia ta działa znakomicie, konsekwentnie przybliżając rozwiązanie wraz z kolejnymi iteracjami. Jeśli więc koszt optymalnego rozwiązania wynosi C^* , liczba iteracji wynosi co najwyżej C^* (31 w najtrudniejszej konfiguracji puzzle-8); w przypadku, gdy funkcja f zwraca inną wartość (niekoniecznie całkowitą) dla każdego węzła, każdy węzeł wyznacza nowy kontur i liczba iteracji równa jest liczbie stanów.

Rekurencyjne wyszukiwanie „najpierw najlepszy” (ang. RBFS — *Recursive Best-First Search* — listing 3.5) naśladuje klasyczną wersję wyszukiwania „najpierw najlepsze”, lecz w przeciwieństwie do niej charakteryzuje się linowym zapotrzebowaniem na pamięć. RBFS przypomina nieco rekurencyjne wyszukiwanie w głąb, ale zamiast posuwać się w nieskończoność na bieżącej ścieżce, kieruje się wartością zmiennej (*limit_f*) zapamiętującej wartość kosztu (funkcji f) najlepszej *alternatywnej* ścieżki prowadzącej od dowolnego przodka¹⁹ bieżącego węzła. Gdy wartość funkcji f dla bieżącego węzła przekracza wartość zmiennej *limit_f*, rekurencja powraca do ścieżki alternatywnej i cofając się, zastępuje wartość f każdego węzła na ścieżce wartością najlepszą spośród jego węzłów potomnych, **zarchiwizowaną** w zmiennej *limit_f*. W rezultacie RBFS zapamiętuje najlepszą wartość f wśród liści „zapomnianego” poddrzewa i na tej podstawie podejmuje decyzję, czy drzewo to warte jest późniejszego ponownego rozwijania. Rysunek 3.16 pokazuje, jak RBFS znajduje drogę z Arad do Bukaresztu.

¹⁹ „Przodek” to przechodnie domknięcie (niepuste) relacji „węzeł macierzysty”. Niech M będzie węzłem macierzystym węzła X ; węzeł Y jest przodkiem węzła X , jeśli $Y = M$ lub Y jest przodkiem M . Mówiąc potocznie, przodkami węzła X są: jego węzeł macierzysty, węzeł macierzysty jego węzła macierzystego itd. — *przyp. tłum.*

LISTING 3.5. Algorytm rekurencyjnego wyszukiwania „najpierw najlepszy”

```

function RECURSIVE-BEST-FIRST-SEARCH(problem) returns rozwiązanie albo sygnalizacja niepowodzenia
    rozwiązanie, wartość_f  $\leftarrow$  RBFS(problem, NODE(problem.INITIAL),  $\infty$ )
    return rozwiązanie

function RBFS(problem, węzeł, limit_f) returns rozwiązanie albo sygnalizacja niepowodzenia
                                                oraz nowy koszt limit_f

    if problem.IS-GOAL(węzeł.STATE) then return węzeł
    następni  $\leftarrow$  LIST(EXPAND(węzeł))
    if następni jest pustą listą then return niepowodzenie,  $\infty$ 
    for each s in następni do // aktualizuj wartość z poprzedniego wyszukiwania
        s.f  $\leftarrow$  max(s.PATH-COST + h(s), węzeł.f)

    while prawda do
        najlepszy  $\leftarrow$  węzeł o najmniejszej wartości f w liście następni
        if najlepszy.f > limit_f then return niepowodzenie, najlepszy.f
        alternatywna  $\leftarrow$  druga najmniejsza wartość f w liście następni
        wynik, najlepszy.f  $\leftarrow$  RBFS(problem, najlepszy, min(limit_f, alternatywna))
        if wynik  $\neq$  niepowodzenie then return wynik, najlepszy.f

```

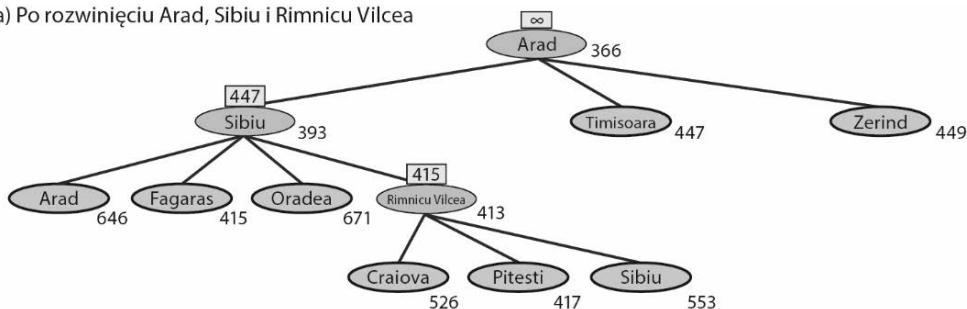
Algorytm RBFS jest nieco wydajniejszy niż IDA*, ale nadal jego bolączką jest nadmierna regeneracja węzłów. W przykładzie na rysunku 3.16, RBFS podąża ścieżką przez Rimnicu Vilcea, by wkrótce rozmyślić się i spróbować drogi przez Fagaras, po czym znowu zmienić zdanie. To swoiste niezdecydowanie bierze się stąd, że gdy aktualnie najlepsza ścieżka jest wydłużana, jej koszt prawdopodobnie rośnie — funkcja *h* jest zwykle mniej optymistyczna dla węzłów położonych blisko celu. W takim przypadku druga co do jakości ścieżka może stać się tą najlepszą, a więc wyszukiwanie musi się cofnąć, by nią podążyć. Każda taka zmiana decyzji odpowiada pojedynczej iteracji IDA* i może wymagać ponownej eksploracji „zapomnianych” węzłów w celu odtworzenia najlepszej ścieżki i rozszerzenia jej o kolejny węzeł.

RBFS jest algorytmem optymalnym kosztowo, o ile heurystyka *h* jest heurystyką dopuszczalną. Jego wymagania pamięciowe rosną proporcjonalnie do najgłębszego optymalnego rozwiązania, natomiast jego złożoność czasowa jest raczej trudna do sprecyzowania, bo zależy (między innymi) od adekwatności funkcji *h* oraz częstotliwości zmian optymalnej ścieżki w miarę rozwijania węzłów. Węzły rozwijane są w kolejności rosnącej funkcji *f*, nawet jeśli nie jest ona monotoniczna.

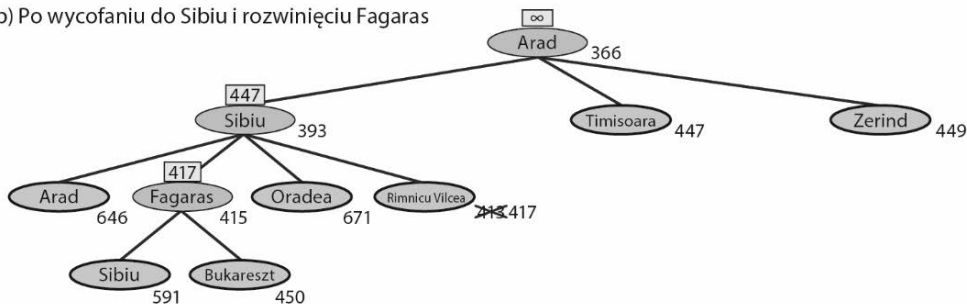
Oszczędność wykorzystywania pamięci, skądinąd ogólnie bardzo pożądana, w przypadku algorytmów IDA* i RBFS jest poniekąd ich słabą stroną. Pomiedzy kolejnymi iteracjami IDA* przechowywana jest tylko jedna liczba — limit kosztu (funkcji *f*). RBFS zapamiętuje co prawda więcej informacji, ale wykorzystuje pamięć jedynie w stopniu proporcjonalnym („liniowym”) do zagłębienia ścieżki i nawet gdyby pamięci było znacznie więcej, to i tak nie może zrobić z niej użytku. Przyrodzona obu algorytmom amnezja co do wcześniejszych działań skazuje je na wielokrotne analizowanie tych samych stanów.

Rozsądne wydaje się zatem żądanie, by można było nakazać algorytmowi wykorzystywanie pamięci w jawnie zadanej wielkości, a być może i całej dostępnej pamięci. Dwa algorytmy, które mają taką możliwość, to MA* (ang. *Memory-Bounded A** — A* ograniczone pamięciowo) i SMA* (ang. *simplified MA** — uproszczone MA*); opiszemy tu SMA* jako (zgodnie z nazwą) mniej skomplikowany. SMA* działa podobnie jak A*, rozwijając najlepszy liść aż do wyczerpania dostępnej pamięci — i tym samym niemożności dodania nowego węzła do drzewa wyszukiwawczego, bez usunięcia któregoś z istniejących węzłów; SMA* usuwa zatem najgorszy liść — ten, dla którego funkcja *f* zwraca największą wartość — i jednocześnie archiwizuje tę wartość w jego węźle macierzystym. Ów węzeł macierzysty „zapomnianego” poddrzewa posiada informację na temat najlepszej ścieżki w tym poddrzewie. Dzięki tej informacji SMA* regeneruje wspomniane poddrzewo tylko wtedy, kiedy inne ścieżki okazują się gorsze od tej „zapomnianej”. Innymi słowy, jeśli zapomniane zostaną wszystkie węzły potomne węzła *n*, nie wiadomo, dokąd podążać dalej z tego węzła.

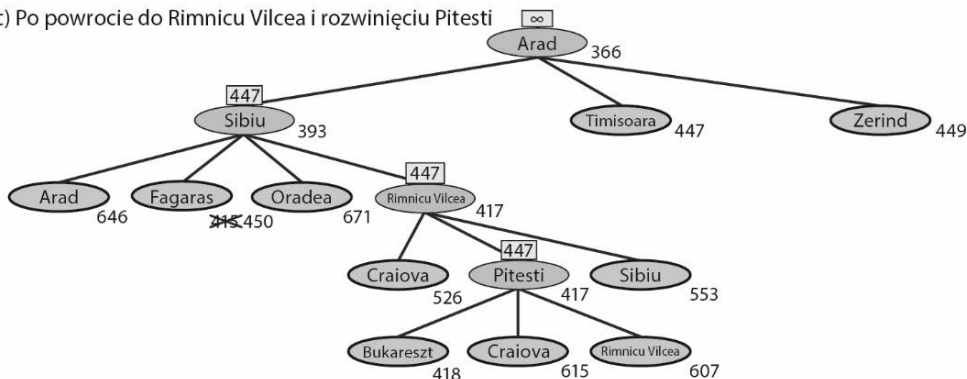
(a) Po rozwinięciu Arad, Sibiu i Rimnicu Vilcea



(b) Po wycofaniu do Sibiu i rozwinięciu Fagaras



(c) Po powrocie do Rimnicu Vilcea i rozwinięciu Pitesti



RYСУNEK 3.16. Wyszukiwanie RBFS najkrótszej ścieżki do Bukaresztu. Nad bieżącymi węzłami widoczne są wartości zmiennej `limit_f` zwracanej przez ostatnie wywołanie rekurencyjne, a każdy węzeł etykietowany jest wartością funkcji f . Etap (a): ścieżka przez Rimnicu Vilcea jest rozwijana aż do wystąpienia sytuacji, gdy najlepszy liść (Pitesti) ma wartość gorszą od najlepszej ścieżki alternatywnej (Fagaras). Etap (b): rekurencja cofa się i wartość f najlepszego liścia „zapomnianego” poddrzewa (417) jest archiwizowana w węzle Rimnicu Vilcea. Następnie rozwijany jest węzeł Fagaras, co prowadzi do wykrycia wartości 450 jako najlepszej wśród liści. Etap (c): rekurencja cofa się i wartość najlepszego liścia „zapomnianego” poddrzewa (450) jest archiwizowana w węzle Fagaras, po czym rozwijany jest węzeł Rimnicu Vilcea. Ponieważ w tym momencie koszt najlepszej ścieżki alternatywnej (przez Timisoara) jest równy co najmniej 447 (czyli więcej niż 417), rozwinięcie węzła Rimnicu Vilcea prowadzi do Bukaresztu

Kompletny opis algorytmu znajduje się w repozytorium online z przykładowym kodem do niniejszej książki, w tym miejscu chcielibyśmy zwrócić uwagę na pewną subtelność. Otóż, jak wyjaśniliśmy, SMA* rozwija *najlepszy* liść, jednocześnie usuwając z pamięci liść *najgorszy* — co się więc stanie, gdy liść najlepszy będzie jednocześnie najgorszym, bo funkcja f zwraca *jednakową wartość* dla wszystkich liści? SMA* rozstrzyga ten problem w ten sposób, że rozwija *najnowszy* z najlepszych liści i usuwa *najstarszy* z najgorszych, co wyklucza możliwość wybrania tego samego węzła jednocześnie do rozwinięcia i do usunięcia. Oczywiście wciąż nie rozwiązuje to problemu w sytuacji, gdy istnieje *tylko jeden liść*, lecz w takim przypadku bieżące drzewo wyszukiawcze musi mieć formę pojedynczej ścieżki, której węzły zajmują *całą dostępną pamięć*. Jeśli ów samotny liść nie jest celem, to nawet gdy znajduje się on na ścieżce do optymalnego rozwiązania, rozwiązanie to jest nieosiągalne z powodu *braku wolnej pamięci*. Samotny liść zostanie więc usunięty z pamięci dokładnie tak, jak gdyby nie posiadał węzłów potomnych.

SMA* jest algorytmem zupełnym, jeżeli istnieje osiągalne rozwiązanie — czyli gdy najpłytszy węzeł docelowy znajduje się na głębokości d takiej, że dostępna pamięć może pomieścić co najmniej $d+1$ węzłów. Algorytm jest optymalny kosztowo, jeżeli osiągalne jest dowolne rozwiązanie optymalne, w przeciwnym razie zwraca on najlepsze spośród osiągalnych rozwiązań. Z praktycznego punktu widzenia SMA* jest dość dobrym wyborem dla znajdowania optymalnych rozwiązań w sytuacji, gdy przestrzeń stanów tworzy graf, akcje mają zróżnicowane koszty, a generowanie węzłów jest kosztowne w porównaniu z kosztami zarządzania granicą i tablicą stanów osiągniętych.

Jednak w przypadku szczególnie trudnych problemów może się zdarzyć, że SMA* zmuszony będzie przełączać się nieustannie między wieloma ścieżkami, stanowiącymi kandydatury na rozwiązanie, i tylko niewielki podzbiór tych ścieżek zmieści się w dostępnej pamięci. Wystąpi wówczas, dobrze znane podsystemom zarządzania pamięcią wirtualną, zjawisko **migotania stron** (ang. *thrashing*) polegające na tym, że system komputerowy większość czasu spędza na nieustannym przeczuciu treści między pamięcią RAM a plikiem wymiany, zamiast wykonywać użyteczne obliczenia. I tak oto *problem teoretycznie rozwiązywalny przy użyciu algorytmu A^* i dostępności nieograniczonej pamięci, może okazać się nierozwiązywalny ze względu na nieakceptowalnie duży czas wykonania algorytmu wynikający z ograniczeń pamięciowych*. Tak to już jest, że oszczędność pamięci okupiona jest zwiększonym czasem obliczeń (i vice versa) i nie istnieje żadna teoria pozwalająca na wypracowanie kompromisu w tym względzie; opisane zjawisko jest nieuniknione i wydaje się, że jedynym sposobem jego złagodzenia jest rezygnacja z wygórowanych wymagań dotyczących optymalności rozwiązania.

3.5.6. Dwukierunkowe wyszukiwanie heurystyczne

Przy jednokierunkowym wyszukiwaniu „najpierw najlepszy” użycie funkcji ewaluacyjnej $f(n) = g(n) + h(n)$ prowadzi do wyszukiwania A^* , gwarantującego znalezienie rozwiązań optymalnych kosztowo (przy założeniu, że h jest heurystyką dopuszczalną) przy jednocześnie optymalnej wydajności w sensie liczby generowanych węzłów.

W przypadku wyszukiwania „najpierw najlepszy” w wersji dwukierunkowej moglibyśmy użyć takiej samej funkcji ewaluacyjnej, jednak nie daje to gwarancji ani optymalności kosztowej rozwiązania, ani optymalnej wydajności, nawet jeśli heurystyka h jest dopuszczalna. Jest zrozumiałe, że przy dwukierunkowym wyszukiwaniu analizujemy nie pojedyncze węzły, lecz raczej *pary węzłów* (po jednym z każdej granicy) co do których można udowodnić, że są niewątpliwie rozwinięte, więc każdy dowód, że wydajność algorytmu jest optymalna, musi uwzględniać właśnie pary węzłów (Eckerle i in., 2017).

W związku z dwukierunkowością wyszukiwania zdefiniujemy nową notację: oznaczenia związane z kierunkiem w przód opatrywać będziemy indeksem F , zaś oznaczenia związane z kierunkiem w tył — indeksem B . Zatem funkcja ewaluacyjna dla kierunku w przód będzie miała postać $f_F(n) = g_F(n) + h_F(n)$, a dla kierunku w tył — $f_B(n) = g_B(n) + h_B(n)$. Mimo iż oba kierunki związane są z rozwiązywaniem tego samego problemu, dla każdego z nich obowiązuje inna funkcja ewaluacyjna, chociażby ze względu na różnicę obu heurystyk: jedna odzwierciedla dążenie do węzła docelowego, druga — do węzła początkowego. Zakładamy, że obie heurystyki są dopuszczalne.

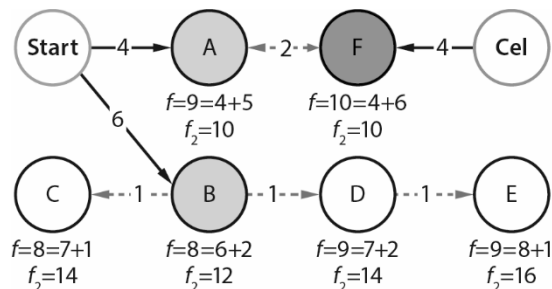
Weźmy pod uwagę dwie ścieżki: jedną od węzła początkowego do węzła m i drugą od węzła docelowego do węzła n . Zdefiniujemy dolne ograniczenie (lb , od ang. *lower bound*) kosztu rozwiązania, które podąża najpierw pierwszą ścieżką od węzła początkowego do węzła m , następnie dociera do węzła n i podąża drugą ścieżką do węzła docelowego:

$$lb(m, n) = \max(g_F(m) + g_B(n), f_F(m), f_B(n))$$

Innymi słowy, koszt takiej ścieżki musi być co najmniej tak duży, jak suma kosztów g obu ścieżek częściowych (ponieważ koszt odcinka łączącego obie części musi być nieujemny), a ponadto nie może być mniejszy od kosztu f którejkolwiek ze ścieżek częściowych (bo wartości heurystyki są optymistyczne). Można udowodnić, że dla dowolnej pary węzłów (m, n) spełniającej warunek $lb(m, n) < C^*$ musimy rozwinąć oba te węzły, ponieważ przechodząca przez nie ścieżka może być potencjalnym rozwiązaniem optymalnym. Sęk w tym, że nie mamy pewności, który węzeł najlepiej rozwinąć, zatem żaden algorytm wyszukiwania dwukierunkowego nie może gwarantować optymalnej wydajności — w skrajnym wypadku, gdy z każdej pary wybrany zostanie niewłaściwy węzeł do rozwinięcia jako pierwszy, liczba wygenerowanych węzłów może być dwukrotnie większa od minimalnej. Niektóre dwukierunkowe algorytmy wyszukiwania heurystycznego w sposób jawny zarządzają kolejką par (m, n) , my jednak pozostaniemy przy wyszukiwaniu „najpierw najlepszy” z listingu 3.4 implementującym granice w postaci kolejek priorytetowych i definiującym funkcję ewaluacyjną zgodnie z kryterium dolnego ograniczenia:

$$f_2(n) = \max(2g(n), g(n) + h(n))$$

Jako następny do rozwinięcia wybierany jest węzeł o najmniejszej wartości f_2 (niezależnie od tego, do której granicy należy). Wyklucza to możliwość wybrania do rozwinięcia węzła n dla którego zachodzi $g(n) > \frac{C^*}{2}$. Mówimy, że dwie połówki wyszukiwania „spotykają się pośrodku” w tym sensie, że w żadnej z granic nie istnieje węzeł, dla którego koszt ścieżki przekraczałby wspomniane $\frac{C^*}{2}$. Przykładowe wyszukiwanie dwukierunkowe zilustrowaliśmy schematycznie na rysunku 3.17.



RYСУNEK 3.17. Dwukierunkowe wyszukiwanie utrzymuje dwie kolejki węzłów. Z lewej strony węzły A i B są następnikami węzła początkowego $Start$; z prawej strony węzeł F jest wstecznym następnikiem węzła Cel . Dla każdego węzła podano wartość funkcji ewaluacyjnej $f = g + h$ i wartość $f_2 = \max(2g, g + h)$. Wartością funkcji g jest suma kosztów akcji, uwidocznionych przy poszczególnych strzałkach, wartości heurystyki h przyjęto dowolnie (nie da się ich wydedukować z rysunku). Koszt optymalnego rozwiązania $Start-A-F-Cel$ wynosi $C^* = 4 + 2 + 4 = 10$, co oznacza, że dwukierunkowy algorytm „spotkania pośrodku” nie powinien rozwijać żadnego węzła, dla którego $g > \frac{C^*}{2} = 5$, więc następnym rozwijanym węzłem na drodze do optymalnego rozwiązania powinien być A albo F (dla obu $g = 4$). Gdybyśmy kierowali się najmniejszą wartością funkcji f , wybralibyśmy do rozwinięcia węzły B i C ($f = 8$) zamiast A ($f = 9$), następnie D i E zostałyby związane z A . Dla nich wszystkich jednak $g > \frac{C^*}{2}$, więc żaden z nich nie zostanie rozwinięty, gdy funkcją ewaluacyjną będzie f_2 .

W opisanym algorytmie heurystyka h_F szacuje odległość danego węzła od celu (lub, gdy problem ma wiele stanów docelowych, odległość od najbliższego celu), analogicznie h_B szacuje odległość od węzła początkowego. Nazywa się to **wyszukiwaniem od czoła do celu** (ang. *front-to-end*). Alternatywa, zwana **wyszukiwaniem od czoła do czoła** (ang. *front-to-front*) wykorzystuje heurystykę mierzącą odległość węzła do drugiej granicy. Oczywiście, gdy granica zawiera miliony węzłów, obliczanie heurystyki dla nich wszystkich (w celu znalezienia tego z najmniejszą wartością) byłoby nieefektywne; rozwiązaniem zastępczym jest wówczas ograniczenie się do niewielkiej *próbki węzłów* losowo wybieranych z granicy. W przypadku pewnych specyficznych domen problemowych możliwe jest sumowanie granicy — na przykład w problemie przeszukiwania gridu możemy przyrostowo obliczać „obwódnię” granicy i w charakterze heurystyki przyjmować odległość węzła od tejże obwódni.

Relacje między wyszukiwaniem dwukierunkowym a jednokierunkowym są różne — dwukierunkowe jest niekiedy bardziej wydajne, niekiedy wręcz przeciwnie. Generalnie, jeśli mamy bardzo dobrze zdefiniowaną heurystykę, to wyszukiwanie A^* generuje kontury koncentrujące się w kierunku celu i dodanie wyszukiwania dwukierunkowego niewiele może pomóc. W przypadku przeciętnej heurystyki wyszukiwanie dwukierunkowe ze spotkaniem granic pośrodku cechuje się mniejszą liczbą rozwijanych węzłów i jako takie jest preferowane. W przypadku najgorszym, czyli przy źle zdefiniowanej heurystyce, wyszukiwanie nie wykazuje tendencji koncentrowania się na celu i wariant dwukierunkowy ma taką samą asymptotyczną złożoność jak A^* . Wyszukiwanie dwukierunkowe z funkcją ewaluacyjną f_2 i heurystyką dopuszczalną jest algorytmem zupełnym i optymalnym kosztowo.

3.6. Funkcje heurystyczne

W tym podrozdziale postaramy się pokazać, jak adekwatność zdefiniowanej heurystyki warunkuje wydajność wyszukiwania oraz jak w ogóle wynajdywane są funkcje heurystyczne. W charakterze przykładu użyjemy gry puzzle-8, której planszę przedstawiliśmy po raz pierwszy na rysunku 3.3 i reprodukuje ponownie na rysunku 3.18. Przypomnijmy — na planszy o wymiarach 3×3 można przesuwac (poziomo lub pionowo, ale nie na skos) 8 kwadratowych kafelków o rozmiarach 1×1 każdy, aby otrzymać pewną predefiniowaną konfigurację docelową; najczęściej w tej roli przyjmowana jest konfiguracja widoczna w prawej części rysunku 3.18.

7	2	4
5		6
8	3	1

Stan początkowy

	1	2
3	4	5
6	7	8

Stan docelowy

RYSunEK 3.18. Jedna z typowych instancji gry puzzle-8, najkrótsze rozwiązanie wymaga 26 akcji

Liczba osiągalnych stanów w tej grze wynosi $9!/2 = 181\,400$ i przechowywanie ich wszystkich w pamięci jest sprawą banalną. Ale w odmianie puzzle-15 — z 15 kafelkami na planszy 4×4 — liczebność przestrzeni stanów sięga $16!/2$, czyli ponad 10 bilionów, a jej przeszukiwanie wymaga wsparcia naprawdę dobrej heurystyki dopuszczalnej. Poszukiwanie takiej heurystyki ma długą historię, najczęściej stosowanymi heurystykami są dwie następujące:

- h_1 — równa jest liczbie kafelków znajdujących się na niewłaściwej pozycji; w części (a) rysunku 3.18 żaden z 8 kafelków nie znajduje się na swoim miejscu, zatem $h_1 = 8$. Heurystyka ta jest dopuszczalna, ponieważ każdy niewłaściwie ułożony kafelek wymaga co najmniej jednego ruchu w celu umieszczenia go na właściwym miejscu.
- h_2 — równa jest sumarycznej odległości dzielącej kafelki od ich pozycji docelowych; ponieważ kafelki nie mogą poruszać się po przekątnej, więc wspomniana „odległość” dla każdego kafelka równa jest sumie ruchów poziomych i pionowych koniecznych do jego sprowadzenia na docelową pozycję. Ograniczenie możliwości ruchu do dwóch prostopadłych kierunków przypomina układ ulic w dużych miastach, liczona w ten sposób odległość nosi nazwę **odległości miejskiej** (ang. *city-block distance*), osiedlowej lub **odległości Manhattanu** (ang. *Manhattan distance*). Gdy zsumujemy odległości kafelków w części (a) rysunku 3.18, otrzymamy

$$h_2 = 3+1+2+2+2+3+3+2 = 18$$

Jak można się było spodziewać, żadna z tych heurystyk nie zawyża rzeczywistego kosztu rozwiązania, który równy jest 26.

3.6.1. Dokładność heurystyki a wydajność algorytmu

Jednym z kryteriów oceny jakości heurystyki jest **efektywny czynnik rozgałęzienia**. W drzewie jednorodnym (czyli takim, w którym każdy węzeł niebędący liściem ma dokładnie b węzłów potomnych) o głębokości d liczba węzłów jest równa

$$W = 1 + b + b^2 + b^3 + \dots + b^d$$

Tę zależność możemy uogólnić na dowolne drzewo wyszukiwawcze. Jeśli w procesie wyszukiwania A^* generowanych jest N węzłów, a najgłębszy węzeł znajduje się na głębokości d , to istnieje wartość b^* spełniająca równanie

$$N + 1 = 1 + b^* + (b^*)^2 + (b^*)^3 + \dots + (b^*)^d$$

(dodajemy 1 do N , by uwzględnić węzeł początkowy). Ta właśnie wartość to wspomniany efektywny czynnik rozgałęzienia.

I tak na przykład, jeśli algorytm A^* znajdzie rozwiązanie o głębokości 5, generując 42 węzły, to wartość b^* równa jest około 1,92. Efektywny czynnik rozgałęzienia może być różny w różnych instancjach problemu, ale zwykle dla określonej domeny (na przykład gry puzzle-8) jest mniej więcej taki sam we wszystkich niebanalnych instancjach problemu. Wyniki eksperymentalnych pomiarów czynnika b^* na małym zbiorze problemów pewnej domeny mogą zatem stanowić dobry przewodnik dotyczący użyteczności zastosowanej heurystyki. Przy dobrze zaprojektowanej heurystyce współczynnik ten powinien mieć wartość bliską 1, co oznacza, że nawet złożone problemy mogą być rozwiązywane kosztem rozsądnej mocy obliczeniowej.

Korf i Reid (1998) twierdzą natomiast, że lepszym miernikiem obcinania w algorytmie A^* przy zastosowaniu określonej heurystyki h jest czynnik k_h , o jaki rzeczywista głębokość d drzewa wyszukiwawczego zredukowana zostaje do **głębokości efektywnej** d_e — w tym sensie, że całkowity koszt wyszukiwania zmniejsza się z $O(bd)$ do

$$O(b^{d_e}) = O(b^{d-k_h})$$

Ich eksperymenty z kostką Rubika i puzzlami- k (dla różnych wartości k) pokazują, że powyższa formuła umożliwia trafne przewidywanie całkowitego kosztu wyszukiwania dla badanych instancji problemów w szerokim zakresie długości zwracanych rozwiązań — przynajmniej tych większych od k_h .

W tabeli 3.4 przedstawiamy rozwiązania losowo wygenerowanych instancji problemu układanki puzzle-8, przy użyciu algorytmu BFS i algorytmu A^* z (opisywanymi wcześniej) heurystykami h_1 i h_2 . Dla poszczególnych długości rozwiązania podajemy średnią liczbę wygenerowanych węzłów i wartość efektywnego czynnika rozgałęzienia. Wyniki te wykazują zarówno przewagę heurystyki h_2 nad h_1 , jak i wyraźną korzyść stosowania heurystyki w ogóle.

TABELA 3.4. Porównanie kosztu wyszukiwania i efektywnego czynnika rozgałęzienia dla losowo wygenerowanych instancji problemu puzzle-8, przy użyciu algorytmów: „najpierw najlepszy” oraz A^* z dwiema heurystykami: h_1 (zwracającą liczbę niewłaściwie ułożonych kafelków) i h_2 (zwracającą sumaryczną odległość Manhattanu dla kafelków). Wyniki zostały uśrednione po 100 instancjach problemu dla każdej długości rozwiązania od 6 do 28

Długość rozwiązania	Koszt wyszukiwania (liczba wygenerowanych węzłów)			Efektywny czynnik rozgałęzienia		
	BFS	$A^*(h_1)$	$A^*(h_2)$	BFS	$A^*(h_1)$	$A^*(h_2)$
6	128	24	19	2,01	1,42	1,34
8	368	48	31	1,91	1,40	1,30
10	1033	116	48	1,85	1,43	1,27
12	2672	279	84	1,80	1,45	1,28
14	6783	678	174	1,77	1,47	1,31
16	17270	1683	364	1,74	1,48	1,32
18	41558	4102	751	1,72	1,49	1,34
20	91493	9905	1318	1,69	1,50	1,34
22	175921	22955	2548	1,66	1,50	1,34
24	290082	53039	5733	1,62	1,50	1,36
26	395355	110372	10080	1,58	1,50	1,35
28	463234	202565	22055	1,53	1,49	1,36

Można by zapytać, czy heurystyka h_2 jest zawsze lepsza od h_1 ?. Odpowiedź brzmi: „Zasadniczo tak” — a z definicji obu heurystyk wynika, że dla dowolnego węzła n zachodzi zależność $h_2(n) \geq h_1(n)$. Mówimy, że h_2 **dominuje** nad h_1 . Dominacja ta przekłada się bezpośrednio na efektywność algorytmu: A^* używający heurystyki h_2 nigdy nie rozwinie więcej węzłów niż A^* używający h_1 (z wyjątkiem oczywiście remisu między h_1 a h_2). Uzasadnienie tego stwierdzenia jest proste. Jak wcześniej wyjaśnialiśmy, każdy węzeł, dla którego funkcja ewaluacyjna f zwraca wartość mniejszą niż C^* , zostanie niewątpliwie rozwinięty; jest to równoważne stwierdzeniu, że każdy węzeł n , dla którego $h(n) < C^* - g(n)$, zostanie niewątpliwie rozwinięty, o ile heurystyka h jest spójna. Ale ponieważ dla każdego węzła n zachodzi zależność $h_2(n) \geq h_1(n)$, każdy węzeł niewątpliwie rozwinięty przez wyszukiwanie (A^*, h_2) zostałby niewątpliwie rozwinięty przez wyszukiwanie (A^*, h_1) , a ponadto (A^*, h_1) może rozwijać jeszcze inne węzły. Generalnie więc korzystnie jest używać heurystyk zwracających większe wartości, pod warunkiem wszakże, iż heurystyki te są spójne, oraz ich obliczanie nie trwa zbyt długo w porównaniu z obliczaniem innych (stąd zastrzeżenie na wstępie, że są one „zasadniczo” lepsze).

3.6.2. Generowanie heurystyk na podstawie rozluźnionych problemów

Gdy analizuje się heurystyki h_1 (liczbę źle ułożonych kafelków) i h_2 (odległość Manhattanu) dla problemu układanki puzzle-8 oraz przewagę drugiej z nich nad pierwszą, można się zastanawiać, skąd biorą się pomysły na tego typu heurystyki i — co ważniejsze — czy komputery są w stanie mechanicznie konstruować takowe?

Wartością $h_1(n)$ i $h_2(n)$ jest oszacowanie długości ścieżki na odcinku od węzła n do węzła docelowego; są to heurystyki adekwatne dla oryginalnego problemu puzzle-8, lecz także dla jego wersji rozszerzonej. Wyobraźmy sobie mianowicie, iż reguły gry rozluźnione zostały w ten sposób, iż kafelek może przesuwać się dokądkolwiek, nie tylko na sąsiednie puste pole; wtedy h_1 nadal powinna zwracać dokładną długość najkrótszego rozwiązania. Podobnie, jeśli dopuścimy, by kafelek mógł przesuwać się na sąsiednie pole w dowolnym kierunku, nawet na pole zajęte już przez inny kafelek, heurystyka h_2 nadal powinna zwracać dokładną długość najkrótszego rozwiązania. Ogólnie, problem będący wynikiem uproszczenia oryginalnego problemu, w sensie zmniejszenia ograniczeń nałożonych na możliwe akcje, nosi nawet **problemu rozluźnionego** (ang. *relaxed problem*). Graf odzwierciedlający przestrzeń stanów problemu rozluźnionego jest *supergrafem* w stosunku do grafu przestrzeni stanów problemu oryginalnego, ponieważ zawiera w stosunku do niego dodatkowe krawędzie będące efektem zniesienia poszczególnych ograniczeń.

Ponieważ graf problemu rozluźnionego różni się od grafu problemu oryginalnego tylko dodatkowymi krawędziami, przeto optymalne rozwiązanie oryginalnego problemu jest — z definicji — rozwiązaniem także problemu rozluźnionego, jednakże problem rozluźniony może mieć lepsze rozwiązanie, jeśli dodatkowe krawędzie jego grafu wnoszą do wynikowej ścieżki jakieś skróty. Zatem *koszt optymalnego rozwiązania problemu rozluźnionego stanowi dopuszczalną heurystykę dla problemu oryginalnego*. Co więcej, ponieważ otrzymana w ten sposób heurystyka odzwierciedla dokładny koszt rozluźnionego problemu, musi ona spełniać nierówność trójkąta, jest więc heurystyką dopuszczalną.

Możliwości rozluźniania oryginalnego problemu stają się bardziej widoczne, jeśli zapiszemy ów problem w jakimś języku programowania²⁰. Dla oryginalnej postaci puzzle-8 pojedynczy ruch kafelka można zapisać mniej więcej w taki sposób:

Kafelek można przesunąć z kwadratu X na kwadrat Y pod warunkiem, że
(kwadraty X i Y sąsiadują ze sobą) i (kwadrat Y jest pusty)

Eliminując selektywnie jeden z warunków ujętych w nawiasy lub oba warunki, możemy uzyskać trzy rozluźnione problemy:

(a)

Kafelek można przesunąć z kwadratu X na kwadrat Y pod warunkiem, że
(kwadraty X i Y sąsiadują ze sobą)

(b)

Kafelek można przesunąć z kwadratu X na kwadrat Y pod warunkiem, że
(kwadrat Y jest pusty)

(c)

Kafelek można przesunąć z kwadratu X na kwadrat Y

Problem (a) jest źródłem heurystyki h_2 (odległości Manhattanu) zgodnie z wymaganiem, że heurystyka ta powinna pozostawać adekwatna, gdy przesuniemy jakiś kafelek w kierunku jego docelowej lokalizacji. Heurystyka wyprowadzona z problemu (b) jest przedmiotem ćwiczenia online 3.GASC. Z problemu (c) można wyprowadzić heurystykę h_1 (której wartością jest liczba niewłaściwie ułożonych kafelków) — powinna ona pozostawać adekwatna w sytuacji, jeśli kafelki będą mogły przemieszczać się do swych docelowych kwadratów w jednym ruchu (czyli w pojedynczej akcji).

²⁰ W rozdziałach 8. i 11. opisujemy formalne języki przydatne do takiego zadania — zapis problemu w takim języku umożliwia mechaniczne wyprowadzanie z niego problemów rozluźnionych. Obecnie wystarczy nam zapis w języku potocznym.

Zauważmy, że rozluźnione problemy generowane w opisany sposób mogą być rozwiązywane *bez wyszukiwania*, ponieważ rozluźnienie reguł pozwala na dekompozycję problemu na osiem niezależnych podproblemów. Jeśli rozluźniony problem jest trudno rozwiązywalny, obliczanie wartości wyprowadzonej z niego heurystyki jest kosztowne.

Możliwe jest automatyczne generowanie heurystyk na podstawie definicji problemu i wywodzących się z niego problemów rozluźnionych, czego dowodem jest program ABSOLVER, wykorzystujący jeszcze inne techniki pomocnicze (Prieditis, 1993). Zespół twórców tego programu wygenerował za jego pomocą nową heurystykę dla układanki puzzle-8, lepszą niż wszystkie dotąd znane; znalazł także pierwszą użyteczną heurystykę dla kostki Rubika.

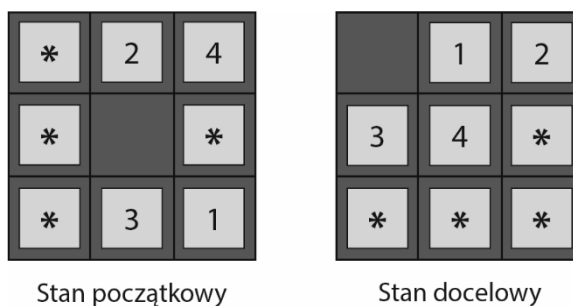
Jeśli dla danego problemu dostępny jest zbiór heurystyk dopuszczalnych $h_1 \dots h_m$ i żadna z nich nie jest wyraźnie lepsza od pozostałych, to którą należałoby wybrać? Okazuje się, że wcale nie musimy wybierać — możemy spożytkować zalety wszystkich, definiując nową:

$$h(n) = \max(h_1(n), \dots, h_m(n))$$

Ta złożona heurystyka wybiera funkcję, która jest najbardziej adekwatna dla danego węzła. Ponieważ składowe heurystyki h_i są heurystykami dopuszczalnymi, h również jest takową; jeśli wszystkie są spójne, h jest również spójna. Z definicji h dominuje nad wszystkimi swymi składowymi. Mankament takiej złożonej heurystyki jest dłuższy czas jej obliczania, wynikający (zazwyczaj) z konieczności obliczenia wszystkich heurystyk składowych; jeśli staje się on problemem, można go zniwelować, wybierając do obliczenia *tylko jedną* z heurystyk składowych — losowo lub (lepiej) na podstawie algorytmu uczenia maszynowego przewidującego, która będzie najlepsza dla danego węzła. Może to skutkować utratą spójności wynikowej heurystyki, ale zwykle prowadzi do szybszego rozwiązywania problemów.

3.6.3. Generowanie heurystyk na podstawie podproblemów — bazy wzorców

Heurystyki (dopuszczalne) dla danego problemu można także wyprowadzać na podstawie kosztów rozwiązań jego **podproblemów**. Na rysunku 3.19 widzimy instancję podproblemu układanki puzzle-8 z rysunku 3.18, polegającego na umieszczeniu we właściwych miejscach kafelków 1, 2, 3, 4 i pustego kwadratu. Rzecz jasna, koszt optymalnego rozwiązania tego podproblemu wyznacza dolne ograniczenie dla kosztu optymalnego rozwiązania pełnego problemu; w niektórych przypadkach okazuje się to oszacowaniem dokładniejszym niż heurystyka odległości Manhattanu.



RYСУNEK 3.19. Instancja podproblemu układanki puzzle-8 z rysunku 3.18 polegającego na sprowadzeniu na docelowe lokalizacje kafelków 1, 2, 3, 4 i pustego miejsca, bez zajmowania się innymi kafelkami

Pomysł **baz wzorców** (ang. *pattern databases*) polega na przechowywaniu dokładnych kosztów rozwiązania dla każdej możliwej instancji podproblemu — w naszym przykładzie, każdej możliwej konfiguracji wymienionych czterech kafelków i pustego miejsca (w bazie będzie wówczas $9 \times 8 \times 7 \times 6 \times 5 = 15\,120$ wzorców; położenia pozostałych czterech kafelków są nieistotne z perspektywy rozwiązywania tego konkretnego podproblemu, ale ich ruchy mogą mieć wpływ na koszt jego rozwiązania). Następnie obliczamy heurystykę h_{DB} dla każdego stanu występującego w wyszukiwaniu, po prostu odczytując jej wartość z bazy dla określonej konfiguracji. Sama baza konstruowana jest poprzez wyszukiwanie wstecz począwszy od stanu docelowego i rejestrowanie kosztów każdego nowo napotkanego wzorca²¹. Koszty takiego wyszukiwania amortyzują się w kolejnych instancjach problemu, ma więc ono sens, jeśli raz zbudowana baza wzorców będzie często wykorzystywana do rozwiązywania wielu problemów.

Wybór kafelków 1-2-3-4 jest całkowicie dowolny, równie dobrze moglibyśmy wybrać kafelki 5-6-7-8 lub 2-4-6-8. Każda baza wzorców dostarcza pewnej heurystyki, a poszczególne heurystyki można ze sobą łączyć (wybierając maksymalną, jak wcześniej wyjaśnialiśmy). Skonstruowana w ten sposób złożona heurystyka jest dokładniejsza niż heurystyka odległości Manhattanu — w przypadku odmiany układanki z 15 kafelkami na planszy 4×4 liczba generowanych węzłów zmniejsza się 1000-krotnie. Każda nowa baza wzorców wiąże się jednak z dodatkowymi wymogami pamięciowymi i dodatkowym czasem obliczeń, co zmniejsza zysk netto kosztu rozwiązania problemu.

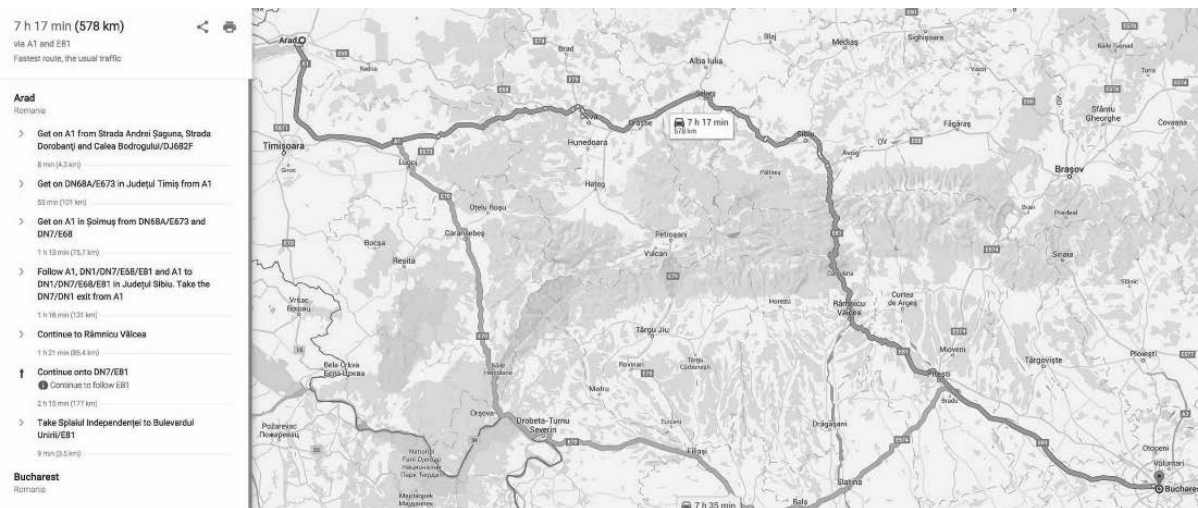
Można się zastanawiać, czy heurystyki uzyskane z baz wzorców dla konfiguracji 1-2-3-4 i 5-6-7-8 można do siebie dodawać, skoro te dwa podproblemy nie pokrywają się — czy wynikowa heurystyka byłaby dopuszczalna? Odpowiedź brzmi „nie”, ponieważ oba wymienione podproblemy prawie na pewno będą miały „część wspólną”, czyli wspólne ruchy dla danego stanu — jest wręcz niemożliwe przesunięcie kafelków 1-2-3-4 na ich miejsca docelowe bez ruszania kafelków 5-6-7-8, i vice versa. Ale co się stanie, jeśli nie uwzględnimy ruchów kafelków 5-6-7-8 — jeżeli, zamiast abstrahować te kafelki do postaci gwiazdek, sprawimy, że po prostu znikną z planszy? Będzie to oznaczać, że nie rejestrujemy całkowitego rozwiązania podproblemu 1-2-3-4, ale jedynie zliczamy ruchy tych właśnie kafelków. Wtedy obie bazy wzorców będą **rozłączne**, suma obu kosztów — liczby ruchów kafelków 1-2-3-4 i kosztu rozwiązania podproblemu 5-6-7-8 — nadal będzie dolnym ograniczeniem kosztu rozwiązania całego problemu. Dzięki takiej rozłączności rozwiązanie problemu układanki puzzle-15 skróci się 10 000 razy w porównaniu ze stosowaniem heurystyki odległości Manhattanu, i trwać będzie zaledwie kilka milisekund. W przypadku 24 kafelków na planszy 5×5 rozwiązywanie problemu skróci się milion razy. Rozłączne bazy wzorców sprawdzają się w przypadku problemu przesuwanych kafelków, ponieważ problem ten daje się podzielić na podproblemy w taki sposób, że jedna akcja związana jest z dokładnie jednym podproblemem (w jednym ruchu przesuwamy dokładnie jeden kafelek).

3.6.4. Generowanie heurystyk w oparciu o punkty orientacyjne

Istnieją usługi online, które udostępniają mapy z dziesiątkami milionów lokalizacji i w ciągu milisekund znajdują optymalne kosztowo trasy przejazdu między nimi (rysunek 3.20). Jak to jest możliwe, wszak algorytmy wyszukiwania, które opisywaliśmy do tej pory, są miliony razy wolniejsze? Jest to wynik wielu trików, ale podstawowa tajemnica kryje się w **preobliczeniach** (ang. *precomputations*), czyli wstępnych obliczeniach optymalnych kosztów ścieżek. Są to obliczenia niewątpliwie czasochłonne, ale wykonywane są tylko raz i czasochłonność ta amortyzuje się niemal natychmiast wobec miliardów żądań napływających od użytkowników.

Moglibyśmy wygenerować doskonałą heurystykę przez wstępne obliczenie i przechowywanie kosztu optymalnej ścieżki między każdą parą wierzchołków. Jeśli potraktujemy mapę jako graf, którego wierzchołki V są lokalizacjami, a krawędzie E drogami łączącymi te lokalizacje, to obliczenie takiej heurystyki wymagałoby $O(|V|^2)$ pamięci i $O(|E|^3)$ czasu — rzecz wykonalna w praktyce dla grafu liczącego 10 000 wierzchołków, ale nie 10 milionów!

²¹ Dzięki wyszukiwaniu wstecz począwszy od celu otrzymujemy natychmiast dokładny koszt rozwiązania każdej napotkanej instancji. Jest to przykład **programowania dynamicznego**, którym zajmujemy się dokładniej w rozdziale 17.



RSYNEK 3.20. Usługa webowa umożliwiająca wytyczanie tras, funkcjonująca w oparciu o algorytmy wyszukiwania

Podejście bardziej praktyczne polega więc na wybraniu kilku (być może 10 lub 20) **punktów orientacyjnych** (ang. *landmarks*)²² spośród lokalizacji (wierzchołków), po czym dla każdego punktu orientacyjnego L należy obliczyć dokładne koszty optymalnych ścieżek $C^*(v, L)$ łączących go z każdym z pozostałych wierzchołków v (będziemy także potrzebować wartości $C^*(L, v)$ — w grafie nieskierowanym wartości te byłyby identyczne, ale w grafie skierowanym, koniecznym choćby dla uwzględnienia dróg jednokierunkowych, należy obliczać je oddzielnie). Dysponując tablicą wartości C^* możemy łatwo skonstruować następującą efektywną (choć niekoniecznie dopuszczalną) heurystykę dla węzła n : dla wszystkich punktów orientacyjnych L obliczamy optymalny koszt ścieżki prowadzącej od węzła n do celu poprzez L i wybieramy najmniejszą z tych wartości:

$$h_l(n) = \min_{L \in \text{punkty orientacyjne}} C^*(n, L) + C^*(L, \text{cel})$$

Jeśli optymalna ścieżka przebiegać będzie przez punkt orientacyjny, heurystyka ta będzie dokładna; jeśli nie, heurystyka ta nie będzie dopuszczalna, bo koszt celu będzie zawyżony. Jeśli w wyszukiwaniu A^* używamy dokładnej heurystyki, to za każdym razem, gdy osiągniemy węzeł nieleżący na optymalnej ścieżce, każdy węzeł wygenerowany wskutek jego rozwinięcia będzie leżał na optymalnej ścieżce. Wyobraźmy sobie kontury tworzące się wokół tej optymalnej ścieżki: wyszukiwanie będzie nią podążać, dodając w każdej iteracji akcje o koszcie c i przechodząc tym samym do nowego stanu, dla którego wartość h będzie mniejsza o c od tej dla stanu bieżącego. Oznacza to, że wartość funkcji $f = g + h$ powinna wynosić niezmiennie C^* wzdłuż całej ścieżki.

Niektóre algorytmy znajdowania tras wykonują się jeszcze szybciej, dzięki **skrótom** (ang. *shortcuts*) — dodanym do grafu sztucznym krawędziom, z których każda reprezentuje optymalną ścieżkę złożoną z wielu akcji. Wyobraźmy sobie, że dysponujemy takimi skrótami między każdą parą 100 największych miast USA; jeśli chcemy dojechać z kampusu Uniwersytetu Berkeley w Kalifornii do gmachu Uniwersytetu w Nowym Jorku, to możemy wykorzystać skrót między Sacramento a Manhattanem, przemierzając 90% trasy w ramach jednej akcji.

²² Zwanych niekiedy także „centrami” (ang. *pivots*) lub „kotwicami” (ang. *anchors*).

Heurystyka $h_L(n)$ jest efektywna, ale nie jest dopuszczalna. Po krótkim zastanowieniu się możemy jednak sformułować podobną heurystykę, która jest zarówno efektywna, jak i dopuszczalna.

$$h_{DH}(n) = \max_{L \in \text{punkty orientacyjne}} |C^*(n, L) - C^*(cel, L)|$$

Heurystyka na nosi nazwę **różnicowej** (ang. *differential heuristic*) ze względu na występujące w niej odejmowanie. Jeżeli cel znajduje się na optymalnej ścieżce od n do L , to widać wyraźnie, że odległość od n do celu szacowana jest przez h_{DH} jako różnica kosztu odcinków (n, L) i (cel, L) . Jeżeli cel leży nieco poza optymalną ścieżką od n do L , h_{DH} nie jest już heurystyką dokładną, ale wciąż jest dopuszczalna. Punkt orientacyjny, który znajduje się na optymalnej ścieżce od n do $celu$, jest raczej bezużyteczny; gdy znajduje się on w połowie drogi między n a $celem$, wartością $h_{DH}(n)$ jest zero.

Punkty orientacyjne można wybierać na wiele sposobów. Najprostszy z nich, wybór losowy, jest szybki, ale mało użyteczny, gdy wiele punktów orientacyjnych znajdzie się zbyt bliskie siebie. Można minimalizować prawdopodobieństwo takich sytuacji, stosując *zachłanną* strategię wyboru: pierwszy punkt orientacyjny wybieramy losowo, a następnie w każdej z kolejnych iteracji ustalamy punkt orientacyjny, który znajduje się jak najdalej od najbliższego z istniejących już punktów orientacyjnych (przeczytaj to jeszcze raz uważnie). Najbardziej elastyczną strategią jest dobór — i okresowe jego aktualizowanie — punktów orientacyjnych na podstawie historii żądań użytkowników. Z perspektywy heurystyki różnicowej najkorzystniejszym układem są punkty orientacyjne równomiernie rozłożone blisko krańców mapy; zbudowanie takiego układu wymaga wykonania następujących kroków:

1. Wyliczamy współrzędne (cx, cy) „środka ciężkości” (centroidu) wszystkich lokalizacji na mapie

$$cx = \frac{1}{N} \sum_{i=1}^N x_i$$

$$cy = \frac{1}{N} \sum_{i=1}^N y_i$$

gdzie (x_i, y_i) , $i = 1, \dots, N$ są współrzędnymi lokalizacji.

2. Ustalamy liczbę punktów orientacyjnych (K) i z punktu (cx, cy) wyprowadzamy K półprostych, w miarę równomiernie rozłożonych kątowo.
3. Jako punkty orientacyjne wybieramy lokalizacje położone jak najbliżej tych półprostych i jednocześnie jak najdalej od centroidu.

Strategia punktów orientacyjnych sprawdza się bardzo dobrze w problemach wyszukiwania tras również ze względu na projektowanie układu dróg na świecie: ponieważ większość użytkowników preferuje wybór trasy właśnie w oparciu o punkty orientacyjne, inżynierowie dążą do projektowania najszerzych i najszybszych tras właśnie między tymi punktami. Wyszukiwanie wykorzystujące heurystyki bazujące na punktach orientacyjnych preferuje wybór takich właśnie tras.

3.6.5. Uczenie maszynowe dla lepszego wyszukiwania

Przedstawiane dotychczas strategie wyszukiwania — wyszukiwanie wszerek, wyszukiwanie A^* itd. — to ustalone algorytmy starannie zaprojektowane i zaprogramowane przez informatyków. Powstaje w tym momencie pytanie, czy agent potrafiłby *nauczyć się* lepszego wyszukiwania? Odpowiedź jest twierdząca, a sposób, w jaki agent mógłby się uczyć, opiera się na ważnej koncepcji **metapoziomowej przestrzeni stanów** (ang. *metalevel state space*). Każdy stan w tej metapoziomowej przestrzeni reprezentuje wewnętrzny (obliczeniowy) stan programu przeszukującego „zwykłą” przestrzeń stanów, taką jak na przykład mapa Rumunii. Aby odróżnić te dwa typy przestrzeni, będziemy tę „zwykłą”

nazywać **obiektołą przestrzenią stanów** (ang. *object-level state space*). Na przykład aktualny stan wewnętrzny algorytmu A* można utożsamiać z bieżącą postacią drzewa wyszukiwawczego; każda akcja w tej metapoziomowej przestrzeni jest krokiem algorytmu zmieniającym stan wewnętrzny — w algorytmie A* taki krok polega na rozwinieciu liścia i dodaniu do drzewa wygenerowanych węzłów potomnych. Rysunek 3.12, przedstawiający kolejne stadia rozwoju drzewa wyszukiwawczego, może być traktowany jako ilustracja *budowania ścieżki w przestrzeni metapoziomowej* — każdy stan na tej ścieżce odpowiada konkretnej postaci drzewa wyszukiwawczego w przestrzeni obiektowej.

Wspomniana „ścieżka” z rysunku 3.12 składa się z pięciu etapów, z których jeden — rozwiniecie przez Fagaras — nie jest szczególnie użyteczny. W przypadku trudniejszych problemów wyszukiwania zdarza się wiele takich chybionych etapów, a agent, wykorzystując algorytm **uczenia metapoziomowego** (ang. *meta-level learning*), zyskuje pole do popisu pod względem umiejętności unikania takich etapów, czyli rezygnacji z eksplorowania mało obiecujących poddrzew. Techniki edukowania agenta w tym kierunku zostały opisane w rozdziale 22., a ich zadaniem jest minimalizowanie **całkowitego kosztu** rozwiązywania problemów, droga kompromisów między kosztem ścieżki a złożonością obliczeniową.

3.6.6. Wynajdywanie heurystyk przez doświadczenie

Opisaliśmy wcześniej jeden ze sposobów konstruowania użytecznych heurystyk — rozwiązywanie „rozluźnionych” wersji oryginalnego problemu, dla których można łatwo znajdować optymalne rozwiązania. Alternatywnym sposobem jest wynajdywanie heurystyk drogą uczenia się przez doświadczenia — zilustrujemy tę koncepcję na przykładzie układanki puzzle-8. Zdobywanie „doświadczenia” polega na rozwiązywaniu wielu takich układanek. Każde optymalne rozwiązanie układanki jest przykładową parą (cel, ścieżka); algorytm uczący może zostać wykorzystany do skonstruowania, na podstawie kolekcji takich par, funkcji heurystycznej h , od której można oczekiwać przybliżania prawdziwego kosztu ścieżki dla innych stanów, które mogą się pojawiać podczas wyszukiwania. Wiele podejść do tego tematu skutkuje niedoskonałą aproksymacją, wskutek czego otrzymana heurystyka może nie być dopuszczalna. Prowadzi to do konieczności rozstrzygania nieuniknionych kompromisów między czasem uczenia się, czasem wyszukiwania i kosztem otrzymywanych rozwiązań. Techniki uczenia maszynowego prezentujemy w rozdziale 19. Uczenie ze wzmacnianiem, opisywane w rozdziale 22., również ma zastosowanie w wyszukiwaniu.

Niektóre techniki uczenia maszynowego spisują się lepiej, gdy opisowi każdego stanu towarzyszy zestaw **cech** (ang. *features*) pomocnych w przewidywaniu wartości heurystyki dla tego stanu, na przykład cecha „liczba źle ułożonych kafelków” może być pomocna w szacowaniu odległości aktualnego stanu układanki od stanu docelowego; oznaczmy tę cechę przez $x_1(n)$. Moglibyśmy wygenerować losowo 100 różnych konfiguracji układanki i zebrać statystykę rzeczywistych kosztów ich rozwiązań. I może na przykład okazać się, że przy wartości $x_1(n) = 5$ średni koszt rozwiązania wynosi około 14, itp.; oczywiście możemy w opisach stanów wykorzystywać wiele cech — drugą pomocną cechą stanu może być „liczba par kafelków, które sąsiadują ze sobą w bieżącym stanie i nie sąsiadują w stanie docelowym”; oznaczmy tę cechę przez $x_2(n)$. Jest zrozumiałe, że w celu otrzymania użytecznej heurystyki powinniśmy *jakoś* połączyć obie te cechy. Najczęściej pod słowem „jakoś” kryje się kombinacja liniowa cech:

$$h(n) = c_1 x_1(n) + c_2 x_2(n)$$

Wartości stałych c_1 i c_2 należy dobrać eksperymentalnie, drogą najlepszego dostosowania funkcji h do danych eksperymentalnych otrzymanych w losowo wygenerowanych konfiguracjach. Można się spodziewać, że c_1 i c_2 będą dodatnie, z prostego względu — większe wartości cech x_1 i x_2 czynią problem trudniejszym do rozwiązania. Zauważmy, że heurystyka h zwraca wartość 0 dla węzłów docelowych, ale niekoniecznie musi być dopuszczalna ani spójna.

Podsumowanie

W tym rozdziale przedstawiliśmy algorytmy wyszukiwania, które agent stosować może do wybierania sekwencji działań w wielu różnych środowiskach — pod warunkiem, że są to środowiska epizodyczne, jednoagentowe, w pełni obserwowalne, deterministyczne, statyczne, dyskretne i rozpoznane. Wyszukiwanie wiąże się nieuchronnie ze znajdowaniem kompromisów między czasem wykonywania algorytmów, ich zapotrzebowaniem na pamięć i jakością zwracanych rozwiązań. Możemy zwiększyć efektywność wyszukiwania, jeżeli wykorzystamy wiedzę zależną od domeny problemu i zakodujemy ją w formie funkcji heurystycznej, której wartość dla danego stanu jest oszacowaniem odległości tego stanu od stanu docelowego. Innym sposobem na przyspieszenie wyszukiwania jest dokonanie pewnych obliczeń wstępnych (związanych z punktami orientacyjnymi) i wykorzystywanie ich wyników.

- Wyszukiwanie można rozpocząć dopiero po ścisłym sformułowaniu definicji **problemu**.
- Definicja problemu zawiera pięć elementów: **stan początkowy**, zbiór **akcji**, **model przejścia** (opisujący efekty zastosowania konkretnych akcji do konkretnych stanów), zbiór **stanów docelowych** i funkcję określającą **koszt poszczególnych akcji**.
- Środowisko problemu reprezentowane jest przez **graf przestrzeni stanów**. **Ścieżka** (ciąg akcji) w przestrzeni stanów, od stanu początkowego do stanu docelowego, stanowi **rozwiązanie** problemu.
- Z perspektywy algorytmów wyszukiwania stany i akcje są (generalnie) **niepodzielne** — nieistotna jest ich wewnętrzna struktura — chociaż można przypisywać poszczególnym stanom cechy, przydatne w procesie uczenia się agenta.
- Podstawowymi kryteriami oceny algorytmów wyszukiwania są: **zupełność**, **optymalność kosztowa**, **złożoność czasowa** i **złożoność pamięciowa**.
- **Wyszukiwanie niedoinformowane** to algorytm, dla którego jedynym źródłem wiedzy jest definicja rozwiązywanego problemu. Działanie algorytmu wyszukiwania opiera się na budowaniu, krok po kroku, drzewa wyszukiwawczego. Ze względu na strategię owego budowania — czyli kryterium wyboru węzłów do rozwijania — rozróżniamy następujące metody wyszukiwania:
 - „**najpierw najlepszy**” — do rozwinięcia wybierany jest węzeł, dla którego funkcja ewaluacyjna zwraca najmniejszą wartość;
 - **wyszukiwanie wszerz** rozwija w pierwszej kolejności węzły o najmniejszej głębokości; wyszukiwanie to jest zupełne, optymalne pod względem liczby akcji w rozwiązaniu, lecz posiada wykładniczą złożoność pamięciową;
 - **wyszukiwanie przy jednolitym koszcie** rozwija w pierwszej kolejności węzły o najmniejszym koszcie ścieżki ($g(n)$), jest optymalne pod względem całkowitego kosztu rozwiązania;
 - **wyszukiwanie w głąb** wybiera do rozwinięcia węzły położone najgłębiej w drzewie (te, które jeszcze nie zostały rozwinięte); nie jest ani zupełne, ani optymalne kosztowo, lecz jego zaletą jest liniowa złożoność pamięciowa. **Ograniczone wyszukiwanie w głąb** ignoruje węzły, których głębokość w drzewie przekracza pewien ustalony limit;
 - **wyszukiwanie z iteracyjnym zagłębianiem** to iteracyjne powtarzanie ograniczonego wyszukiwania w głąb, z sukcesywnie zwiększanym, w każdej iteracji, ograniczeniem głębokości — aż do osiągnięcia węzła docelowego. Jest algorytmem zupełnym, a jeśli przeprowadzana jest kompletna eliminacja cykli, jest optymalny pod względem liczby akcji w rozwiązaniu. Jego złożoność czasowa jest podobna jak przy wyszukiwaniu wszerz, natomiast jego złożoność pamięciowa jest liniowa.
 - **wyszukiwanie dwukierunkowe** działa w dwóch przeciwnych kierunkach — od węzła początkowego i od węzła docelowego — budując dwie niezależne granice, których spotkanie się kończy wykonywanie algorytmu.

- Metody **wyszukiwania poinformowanego** wykorzystują **funkcje heurystyczne**; wartość takiej funkcji $h(n)$ jest oszacowaniem częściowego kosztu rozwiązania, od węzła n do węzła docelowego. Metody te mogą również wykorzystywać inne źródła informacji, na przykład bazy wzorców zawierające koszty rozwiązań wybranych problemów. Ważniejsze z metod tej kategorii to:
 - **zachłanna wersja wyszukiwania „najpierw najlepszy”** — wybiera do rozwijania węzły z najmniejszą wartością h . Nie jest optymalna kosztowo, lecz zwykle jest efektywna;
 - **wyszukiwanie A^*** rozwija w pierwszej kolejności węzły z najmniejszą wartością funkcji ewaluacyjnej $f(n) = g(n) + h(n)$. Jest algorytmem zupełnym i optymalnym kosztowo, pod warunkiem, że heurystyka h jest heurystyką dopuszczalną. Duża złożoność pamięciowa tej metody wciąż stanowi poważne wyzwanie w przypadku większości problemów;
 - **dwukierunkowe wyszukiwanie A^*** jest niekiedy bardziej efektywne niż wersja jednokierunkowa;
 - **A^* z iteracyjnym zagłębianiem (IDA^*)** to połączenie klasycznego wyszukiwania A^* z iteracyjnym zagłębianiem, cechuje się mniejszymi wymaganiami pamięciowymi niż klasyczne A^* ;
 - **rekurencyjne wyszukiwanie „najpierw najlepszy” (RBFS) i A^* ograniczone pamięciowo (SMA*)** to solidne, optymalne kosztowo algorytmy zdolne do efektywnego wykorzystywania pamięci w granicach jej dostępności: przydatne do rozwiązywania problemów, dla których wymagania pamięciowe klasycznego A^* są dyskwalifikujące;
 - **wyszukiwanie skupione (wiązkowe)** narzuca górny limit na liczbę węzłów tworzących granicę. Nie jest algorytmem zupełnym ani optymalnym kosztowo, zwracane rozwiązania są jednak zwykle bliskie optymalności, poza tym czas wykonania jest znacząco krótszy w porównaniu z algorytmami zupełnymi;
 - **ważone wyszukiwanie A^*** koncentruje się na szybkim zmierzaniu do celu, nie gwarantując optymalności kosztowej.

Wydajność algorytmów wyszukiwania wykorzystujących heurystykę zależy przede wszystkim od jakości funkcji heurystycznej. Źródłami dobrych heurystyk mogą być: optymalne rozwiązania rozluźnionych wersji oryginalnego problemu, obliczone wcześniej i zapamiętane w bazach wzorców wartości kosztów rozwiązań podproblemów, oraz uczenie maszynowe przez doświadczenie w rozwiązywaniu problemów danej klasy.

Bibliografia i uwagi historyczne

Temat wyszukiwania w przestrzeni stanów pojawił się w historii sztucznej inteligencji dość wcześnie. Prace Newella i Simona dotyczące teorii logiki (1957) i GPS (1961) doprowadziły do ustanowienia algorytmów wyszukiwania jako podstawowego narzędzia dla badaczy z lat 60. ubiegłego wieku, a rozwiązywanie problemów uznane zostało za kanoniczne zadanie sztucznej inteligencji. Praca Richarda Bellmana (1957) z dziedziny badań operacyjnych wykazała znaczenie addytywności kosztów ścieżek w upraszczaniu algorytmów optymalizacji. Tekst Nilsa Nilssona (1971) dostarczył temu obszarowi solidnych podstaw teoretycznych.

Układanka puzzle-8 jest „uboższym kuzynem” układanki puzzle-15, której historię szczegółowo opisują Slocum i Sonneveld (2006). W 1880 roku puzzle-15 zwróciły uwagę opinii publicznej i matematyków (Johnson i Story, 1879; Tait, 1880) — redaktorzy „American Journal of Mathematics” stwierdzili: „układanka puzzle-15 w ciągu ostatnich kilku tygodni stała się tematem nr 1 w opinii publicznej, zyskując sobie uznanie dziewięciu na dziesięciu Amerykanów, niezależni od płci, wieku i pozycji społecznej”. 12 marca 1880 roku, na łamach tygodnika „Weekly News-Democrat” w Emporia (w stanie Kansas), porównano wręcz tę układankę do „epidemii, która ogarnęła cały kraj”.

Sam Loyd, słynny amerykański projektant gier, ogłosił się samozwańczo wynalazcą puzzle-15 (Loyd, 1959), faktycznie jednak wymyślił ją w latach 70. XIX wieku Nouyes Chapman, listonosz z wioski Canastota w stanie Nowy Jork; oficjalny patent na łamigłówkę z przesuwanymi kafelkami przyznano jednak w 1878 roku Ernestowi Kinseyowi. Ratner i Warmuth (1986) wykazali, że ogólna wersja puzzli na planszy $n \times n$ (puzzle- k , gdzie $k = n^2 - 1$) należy do klasy problemów NP-zupełnych (patrz Dodatek A).

Ernö Rubik wymyślił swą słynną kostkę w 1974 roku, odkrywając również algorytm jej dobrych, ale nie optymalnych, rozwiązań. Korf (1997) znalazł optymalne rozwiązania dla niektórych losowo wygenerowanych instancji problemu, używając wyszukiwania IDA* w połączeniu z bazami wzorców. Rokicki i in. (2014) udowodnili, że każdą konfigurację kostki Rubika można rozwiązać w maksymalnie 26 ruchach, jeśli obrót o 180 stopni potraktujemy jako dwa ruchy (w 20 ruchach, jeśli potraktujemy ów obrót jako jeden ruch); znalezienie tego dowodu pochłonęło 35 lat (!) pracy procesora, lecz nie prowadziło bezpośrednio do wydajnego algorytmu. Agostinelli i in. (2019) wykorzystali uczenie maszynowe ze wzmacnianiem, głębokie uczenie i przeszukiwanie drzewa metodą Monte Carlo²³ do „wyuczenia” znacznie wydajniejszego *solvera*, dającego rozwiązanie w czasie krótszym niż sekunda, lecz bez gwarancji optymalności.

Każdy z przytaczanych w tym rozdziale przykładów wyszukiwania, zaczerpniętych z rzeczywistego świata, był przedmiotem wielu prac badawczych. Metody zestawiania optymalnych tras przelotów chronione są w większości tajemnicą i prawami autorskimi, niemniej jednak Carl de Marcken wykazał, że (wobec ogromnych komplikacji taryf i wielu innych ograniczeń) problem konstruowania optymalnego lotu jest formalnie nierozstrzygalny; swój wynik uzyskał przez redukcję oryginalnego problemu do (nierozstrzygalnych) diofantycznych problemów decyzyjnych (Robinson, 2002). Problem komiwojażera (TSP) jest jednym ze standardowych problemów kombinatorycznych w informatyce teoretycznej (Lawler i in., 1992). Karp (1972) udowodnił, że problem decyzyjny TSP²⁴ jest NP-trudny (patrz Dodatek A), ale opracowano skuteczne metody heurystycznej aproksymacji dokładnego rozwiązania (Lin i Kernighan, 1973). Arora (1998) opracował schemat (o złożoności wielomianowej) aproksymacji rozwiązania TSP w sytuacji, gdy koszty dróg są ich długościami w przestrzeni euklidesowej. Problem optymalnych topografii układów VLSI badany był przez LaPaugh (2010), a liczne artykuły poświęcone tej problematyce pojawiają się w specjalistycznych czasopismach. Nawigację robotów omawiamy w rozdziale 26. Automatyczne sekwencjonowanie montażu zostało po raz pierwszy zademonstrowane przez robota FREDDY (Michie, 1972), a obszernej recenzji dostarczają Bahubalendruni i Biswal (2016).

Niedoinformowane algorytmy wyszukiwania są jednym z głównych obszarów informatyki (Cormen i in., 2009) i badań operacyjnych (Dreyfus, 1969). Wyszukiwanie wszcz zostało sformułowane przez Moore’a (1959) w celu rozwiązywania labiryntów. Metoda programowania dynamicznego (Bellman, 1957; Bellman i Dreyfus, 1962), która systematycznie rejestruje rozwiązania dla wszystkich podproblemów o rosnącej długości, może być postrzegana jako odmiana wyszukiwania wszcz.

Algorytm Dijkstry — w formie, w jakiej jest zazwyczaj przedstawiany (Dijkstra, 1959) — ma zastosowanie do jawnych grafów skończonych. Nilsson (1971) przedstawił wersję algorytmu Dijkstry, którą nazwał „wyszukiwaniem przy jednolitym koszcie” (według jego sformułowania „rozciga się wzdłuż konturów o jednakowych kosztach ścieżki”); wersja ta umożliwia wyszukiwane również w grafach niejawnie definiowanych i (lub) nieskończonych. W pracy Nilssona znajduje się wyjaśnienie koncepcji zamkniętych i otwartych list oraz „wyszukiwania grafowego”. Nazwa BEST-FIRST-SEARCH pojawiła się po raz pierwszy w *Handbook of AI* (Barr i Feigenbaum, 1981). Algorytmy Floyda-Warshalla (Floyd, 1962) i Bellmana-Forda (Bellman, 1958; Ford, 1956) dopuszczają ujemne koszty akcji, o ile w grafie przestrzeni stanów nie występują cykle o ujemnym koszcie.

Pewna odmiana iteracyjnego zagłębiania, zaprojektowana w celu efektywnego wykorzystywania zegara szachowego, została po raz pierwszy zastosowana przez Slate’a i Atkina (1977) w programie CHESS 4.5. W algorytmie *B* Martellego (1977) również zauważyć można ślady takiego zagłębiania. Samo sformułowanie techniki iteracyjnego zagłębiania zawdzięczamy Bertramowi Raphaelowi (1976), a wyeksponowana ona została w pracy Korfa (1985a).

²³ Patrz https://pl.wikipedia.org/wiki/Monte-Carlo_Tree_Search — przyp. tłum.

²⁴ Czyli odpowiedź na pytanie o istnienie trasy komiwojażera o koszcie mniejszym niż ustalone K — przyp. tłum.

Wykorzystanie heurystyk w rozwiązywaniu problemów pojawia się we wczesnej pracy Simona i Newella (1958), ale sformułowanie „wyszukiwanie heurystyczne” i przykłady wykorzystywania funkcji heurystycznych szacujących odległość do celu, pojawiły się nieco później (Newell i Ernst, 1965; Lin, 1965). Zakrojone na szeroką skalę eksperymenty w zakresie wyszukiwania heurystycznego przeprowadzone zostały przez Dorana i Michiego (1966). Chociaż przeanalizowali długość ścieżki i jej „penetrancję” (czyli stosunek jej długości do całkowitej liczby odwiedzonych dotychczas węzłów), wydaje się, że zignorowali informacje dostarczone przez koszt ścieżki (funkcję $g(n)$). Algorytm A*, uwzględniający ten koszt w wyszukiwaniu heurystycznym, został opracowany przez Harta, Nilssona i Raphaela (1968). Dechter i Pearl (1985) badali warunki, pod jakimi A* jest optymalnie wydajny (pod względem liczby rozwiniętych węzłów).

W oryginalnej pracy poświęconej A* (Hart i in., 1968) sformułowany został warunek spójności dla funkcji heurystycznych. Warunek monotoniczności został wprowadzony przez Pohla (1977) jako prostszy zamiennik spójności, ale Pearl (1984) wykazał, że oba warunki są sobie równoważne.

Pohl (1977) był pionierem w badaniu związku między błędami funkcji heurystycznych a złożonością czasową wyszukiwania A* — podstawowe wyniki w tym zakresie uzyskano dla wyszukiwania drzewiastego z jednostkowymi kosztami akcji i zarówno jednym stanem docelowym (Pohl, 1977; Gaschnig, 1979; Huyn i in., 1980; Pearl, 1984), jak i wieloma stanami docelowymi (Dinh i in., 2007). Korf i Reid (1998) pokazali, jak można przewidzieć dokładną liczbę rozwiniętych węzłów, a nie tylko jej asymptotyczne przybliżenie, w różnych rzeczywistych domenach problemowych. Koncepcja „efektywnego czynnika rozgałęzienia” została zaproponowana przez Nilssona (1971) jako empiryczna miara wydajności algorytmu. Helmert i Röoer (2008) zauważyli, że przy wyszukiwaniu grafowym, w przypadku kilku dobrze znanych problemów, w których liczba węzłów na optymalnej ścieżce rośnie wykładniczo do rozmiaru problemu, wyszukiwanie A* wykazuje wykładniczą złożoność czasową.

Istnieje wiele odmian algorytmu A*. Pohl (1970) wprowadził wyszukiwanie ważone A*, a później jego wersję dynamiczną (1973), w której waga zmienia się wraz z głębokością drzewa. Ebendt i Drechsler (2009) dokonali syntezy znanych wyników i przeprowadzili analizę niektórych zastosowań. Hatem i Ruml (2014) zaprezentowali uproszczoną i ulepszoną wersję ważonego A*, łatwiejszą w implementacji. Wilt i Ruml (2014) zaproponowali alternatywę dla zachłannego wyszukiwania — szybkie wyszukiwanie koncentrujące się na minimalizacji czasu wyszukiwania, a w swej późniejszej pracy (2016) udowodnili, że najlepsze heurystyki dla wyszukiwania satysfakcjonującego różnią się od tych dla wyszukiwania optymalnego. Kilka przydatnych sztuczek ułatwiających implementowanie szybkiego wyszukiwania znajdują programiści w pracy Burnsa i in. (2012), a Felner (2018) pokazuje jak zmienia się kod implementujący szybkie wyszukiwanie w związku z wcześniejszym testowaniem węzła na okoliczność bycia węzłem docelowym.

Pomysłodawcą wyszukiwania dwukierunkowego jest Pohl (1971). Holte i in. (2016) opisują wersję wyszukiwania dwukierunkowego, która gwarantuje spotkanie granic pośrodku, dzięki czemu ma szersze zastosowanie. Eckerle i in. (2017) wprowadzają pojęcie węzła niewątpliwie rozwiniętego, opisują zbiór par takich węzłów i udowadniając, że żadne wyszukiwanie dwukierunkowe nie gwarantuje optymalnej wydajności. Algorytm NBS (Chen i in., 2017) jawnie wykorzystuje kolejkę par węzłów.

Kombinacja dwukierunkowego wyszukiwania A* i koncepcji punktów orientacyjnych stanowi podstawę usługi online umożliwiającej efektywne wyszukiwanie tras, udostępnianej przez Microsoft (Goldberg i in., 2006). Po wstępnym obliczeniu i zapamiętaniu kosztów dróg łączących punkty orientacyjne, usługa potrafi znaleźć optymalną trasę między każdą parą spośród 24 milionów punktów na mapie USA, przeszukując zaledwie 0,1% ich populacji. Korf (1987) demonstruje kilka dodatkowych technik — podcele, makrooperatory i abstrakcje — które pozwalają uzyskać jeszcze większe przyspieszenie podstawowych mechanizmów. Delling i in. (2009) pokazują, jak używać wyszukiwania dwukierunkowego, punktów orientacyjnych, struktury hierarchicznej i innych sztuczek do znajdowania optymalnych tras samochodowych. Anderson i in. (2008) opisują pokrewną metodę, zwaną **wyszukiwaniem precyzowanym** (ang. *coarse-to-fine search*), której istotą jest wykorzystywanie punktów orientacyjnych na różnych poziomach hierarchicznej abstrakcji. Korf (1987) opisuje warunki, przy spełnieniu których metoda ta zapewnia przyspieszenie w stopniu wykładniczym. Wyniki eksperymentów i analiz badających ilościowe aspekty wyszukiwania hierarchicznego prezentuje Knoblock w swej pracy (1991).

A* i inne algorytmy wyszukiwania w przestrzeniach stanów są ściśle związane z technikami **podziału i ograniczeń** (ang. *branch-and-bound*), stosowanymi powszechnie w badaniach operacyjnych (Lawler i Wood, 1966; Rayward-Smith i in., 1996). Kumar i Kanal (1988) podejmują próbę „wielkiej unifikacji” wyszukiwania heurystycznego, programowania dynamicznego oraz techniki podziału i ograniczeń, opatrując swą teorię akronimem CDP, od ang. *composite decision process* — „złożony proces decyzyjny”.

Ponieważ typowy komputer w latach 60. ubiegłego wieku posiadał co najwyżej kilka tysięcy słów pamięci, obiektem intensywnych badań stały się wówczas metody wyszukiwania z ograniczeniami pamięciowymi. Program *Graph Traverser* (Doran i Michie, 1966), jeden z pierwszych tego rodzaju, wybiera najlepszą akcję znaną przez wyszukiwanie „najpierw najlepszy” prowadzone aż do wyczerpania dostępnej pamięci. Pierwszym powszechnie stosowanym, optymalnym pod względem długości ścieżki wynikowej, heurystycznym algorytmem wyszukiwania z ograniczeniami pamięciowymi był IDA* (Korf, 1985b), nic więc dziwnego w tym, że opracowano sporo jego wariantów. Analiza efektywności algorytmu IDA* i jego trudności związanych z wykorzystywaniem *funkcji heurystycznych zwracających liczby rzeczywiste* pojawia się w pracy Patrick i in. (1992).

Oryginalna wersja algorytmu RBFS (Korf, 1993) jest faktycznie nieco bardziej skomplikowana niż algorytm pokazany na listingu 3.5, któremu w rzeczywistości bliżej jest do (opracowanego niezależnie) algorytmu zwanego **iteracyjnym rozwijaniem** (ang. IE — *iterative expansion*) (Russell, 1992). RBFS używa zarówno dolnego, jak i górnego ograniczenia, oba algorytmy zachowują się identycznie w przypadku heurystyki dopuszczalnej, ale RBFS rozwija węzły począwszy od najlepszych, nawet przy niedopuszczalnej heurystyce. Pomysł śledzenia najlepszej ścieżki alternatywnej pojawił się wcześniej w eleganckiej implementacji A* w języku Prolog (Bratko, 2009) oraz w algorytmie DTA* (Russell i Wefald, 1991) — w tej ostatniej pracy omawiane są także metapoziomowe przestrzenie stanów i metapoziomowe uczenie maszynowe.

Algorytm MA* pojawił się w pracy Chakrabarti i in. (1989); jego uproszczona wersja — SMA* — jest wynikiem wysiłków zmierzających do jego zaimplementowania (Russell, 1992). Kaindl i Khorsand (1994) zastosowali SMA* do stworzenia dwukierunkowego algorytmu wyszukiwania, który okazał się znacznie szybszy od poprzedników. Korf i Zhang (2000) opisują podejście „dziel i zwyciężaj”, a Zhou i Hansen (2002) wprowadzają ograniczone pamięciowo grafowe wyszukiwanie A* oraz strategię jego przełączania się na wyszukiwanie wszędy w celu bardziej efektywnego wykorzystywania pamięci (Zhou i Hansen, 2006).

Pomysł wyprowadzania dopuszczalnych heurystyk z rozluźnionych wersji oryginalnego problemu pojawił się w przełomowej pracy Helda i Karpa (1970), którzy zastosowali heurystykę minimalnego drzewa rozpinającego do rozwiązania problemu komiwojażera (patrz Ćwiczenie online 3.MSTR.). Automatyzacja procesu „rozluźniania” problemów została pomyślnie zaimplementowana przez Frieditisa (1993). Pojawia się coraz więcej literatury na temat zastosowania uczenia maszynowego do konstruowania heurystyk (Samadi i in., 2008; Arfaee i in., 2010; Thayer i in., 2011; Lelis i in., 2012).

Wykorzystanie baz wzorców do wyprowadzenia dopuszczalnych heurystyk jest zasługą Gassera (1995) oraz Culbersona i Schaeffera (1996, 1998); problem rozłączności baz wzorców opisywany jest przez Korfa i Felnera (2002). Wykorzystywanie wzorców symbolicznych jest pomysłem Edelkampa (2009). Felner i in. (2007) omawiają *kompresję* baz wzorców w celu oszczędniejszego wykorzystywania pamięci. Probabilistyczną interpretację heurystyki badali Pearl (1984) oraz Hansson i Mayer (1989).

Wśród podręczników traktujących o heurystykach w algorytmach wyszukiwania na szczególną uwagę zasługują *Heuristics* Pearl’a (1984) i *Heuristic Search* Edelkampa i Schrödl’a (2012). Publikacje na temat nowych algorytmów wyszukiwania pojawiają się pod egidą *International Symposium on Combinatorial Search* (SoCS) i *International Conference on Automated Planning and Scheduling* (ICAPS), a także w ramach wielu konferencji poświęconych sztucznej inteligencji, między innymi AAAI oraz IJCAI. Można je także znaleźć na łamach znanych czasopism, takich jak „Artificial Intelligence” i „Journal ACM”.

SKOROWIDZ

A

- abstrahowanie stanów
 - w planowaniu, 401
- abstrakcja, 86
- adaptacyjne proponowanie rozkładów, 579
- addytywna
 - funkcja wartościująca, 602
 - nagroda niedyskontowana, 626
 - nagroda zdyskontowana, 625
- agent, 20, 55, 80
 - bazujący na wiedzy, 243
 - decyzyjny, 657
 - hybrydowy, 277
 - inteligentny, 53, 242
 - komponenty programu, 78
 - logiczny, 273
 - odruchowy, 69, 81
 - oparty na modelach, 72, 81
 - online, 167
 - planowania online, 426
 - program, 67
 - programowy, 63
 - racjonalny, 20, 59, 616
 - rozwiązujący problem, 83
 - uczący się, 76
 - ukierunkowany na cele, 73, 81
 - uległy, 661
 - w częściowo obserwowalnym środowisku, 161
 - z funkcją użyteczności, 74
 - zbierający informacje, 609
- AI Index, 46
- AI, Artificial Intelligence, 17, 35
- akcje, 85, 87, 156
 - nieodwracalne, 166
 - pierwotne, 403
 - preferowane, 400
 - relewantne, 395
- stosowalne, 389, 395
- wspólne, 665
- wynik wykonania, 389
- wysokiego poziomu, 403, 426
- aksjomat, 243, 304
 - wykluczenia akcji, 282, 666
 - efektu, 275
 - Kołmogorowa, 442
 - następnika stanu, 276, 283
 - prawdopodobieństwa, 441
 - ramy, 275
 - warunków wstępnych, 282
- aksjomatyka Peano, 305
- aktualizacja, 159
 - Bellmana, 633
 - rozkładów Gaussa, 534
- aktuator, 55, 62, 245
- algebra Robbinsa, 355
- algorytm, 25
 - AC-3, 220
 - agenta hybrydowego, 278
 - cząsteczkowego MCMC, 554
 - Davisa-Putnama, 268
 - dekodowania turbo, 510
 - Dijkstry, 100
 - dla problemów
 - MDP, 633
 - POMDP, 651
 - eliminacji zmiennych, 480
 - enumeracji, 479
 - filtrowania cząstek, 552
 - iterowania
 - polityki, 637
 - wartości, 633, 654
 - kaskady cząstek, 554
 - klasteryzacji, 485
 - LRTA*, 170, 175
 - łańcuchowania, 265
 - progresywnego, 327, 351
 - regresywnego, 332, 351
- MA*, 119
- MAC, 228
- MCMC, 566, 573
- Metropolisa, 172
- Metropolisa-Hastingsa, 493
- MIN-CONFLICTS, 231
- minimax, 673
- Monte Carlo, 486
- NET-VISA, 569
- ONLINE-DFS-AGENT, 168
- online dla problemów
 - MDP, 639
 - POMDP, 655
- planowania
 - hierarchicznego, 409
 - klasycznego, 393
- próbkiowania, 487
- próbkiowania z odrzucaniem, 488
- przeszukiwania grafów AND-OR, 153
- RETE, 331
- rezolucji, 263
- rozstrzygania o wynikaniu, 256
- SATPLAN, 281
- simplex, 173
- SMA*, 119
- symulowanego wyżarzania, 142
- świadka, 659
- TREE-CSP-SOLVER, 234
- UCT, 641, 656
- unifikacji, 323
- Viterbiego, 526
- ważenia wiarygodności, 491
- węgierski, 573
- wnioskowania, 250
 - MCMC, 562
 - progresywnego, 266
- wyszukiwania, 83, 92, 116
 - A*, 111, 112, 133
 - alfa-beta, 206
 - kryteria, 132

algorytm
 wyszukiwania
 lokalnego, 138, 270
 minimaksowego, 180
 porównanie, 108
 przy jednolitym koszcie, 99
 przyrostowego, 158
 rekurencyjnego, 119
 w głąb, 101
 w grach, 204
 wspinaczkowego, 138, 141
 wszerz, 99, 404
 zanikowego MCMC, 554
 zupełny, 99, 328
 zupełny nawracania, 268
 algorytmy
 ewolucyjne, 143, 171
 genetyczne, 143–146
 łańcuchowania, 265
 macierzowe, 527
 niedoinformowane, 83, 98
 poinformowane, 83
 alokacja zasobów, 696
 alternatywa, 252
 wykluczająca, 254
 analityk, 618
 analiza
 decyzyjna, 618
 łańcuchów Markowa, 495
 przypadku
 ubezpieczenie samochodu, 475
 rodowodowa, 508
 składowych, 270
 wrażliwości, 611
 angeliczny niedeterminizm, 406
 aproksymacja, 408
 opisu optymistycznego, 408
 opisu pesymistycznego, 408
 polityki, 646
 zachowawcza, 279
 aproksymowane wnioskowanie, 500
 aproksymowanie zachowawcze, 415
 architektury poznawcze, 331
 argument przeciwnika, 166
 argumentowość, 295
 arność, 295, 328
 asercje, 303
 asercje probabilistyczne, 438

ASI, Artificial Superintelligence, 52
 atak sybilli, 563
 atrybut, 79
 autolokalizacja, 162
 robota, 529
 z jednoczesnym mapowaniem, 550
 automatyczne
 dowodzenie przez rezolucję, 351
 sekwencjonowanie montażu, 92
 wnioskowanie, 18
 autonomia, 61
 autonomiczne planowanie, 48

B

badania operacyjne, 27
 bandyta
 Bernoulliego, 645
 jednoręki, 643
 wieloręki, 641
 zero-jedynkowy, 645
 baza
 Herbranda, 346
 wiedzy, 242, 243, 254, 283
 debugowanie, 315
 ewaluacja, 315
 danych
 deduktywna, 332, 353
 wzorców, 128, 402
 rozłączna, 128
 behawioryzm, 30
 Big Data, 44
 bilard, 211
 binarny diagram decyzyjny, 428

C

cache'owanie, 279
 cechy, 187
 cel, 390
 ciało, 265
 ciąg percepcyjny, 56
 ciągłość, 588
 CSP, constraint satisfaction problem,
 212, 216, 218, 224
 cybernetyka, 33
 cyrkumskrypcja, 378, 382
 cyrkumskrypcja priorytetowa, 379

czas, 367
 mieszania, 520
 zatrzymania, 643
 częstotliwość, 458
 częstotliwość mutowania, 143
 czujniki, 55, 62, 245
 błędu, 503
 pewności, 41, 511
 reprodukcyjne, 143
 robota, 162
 czysty symbol, 268

D

decydent, 618, 661, 662
 decyzje
 minimaksowe, 180, 612
 optymalne, 181
 proste, 585
 strategiczne, 662
 w środowisku wieloagentowym, 661
 złożone, 621
 deduktywne bazy danych, 332, 353
 dekompozycja
 drzewiasta, 236, 238
 hierarchiczna, 402
 problemu, 401
 zupełna, 360
 dekompozycyjność, 588
 demodulacja, 348, 352, 354
 demoniczny niedeterminizm, 406
 diagram wpływów, 507, 585
 dobra deficytowe
 licytowanie, 697
 dobro
 społeczne, 672
 egalitarne, 672
 struktury koalicyjne, 694
 użytkitarne, 672
 wspólne, 700
 dokładność modelu HMM, 532
 domena
 elektroniczne układy cyfrowe, 311
 liczby, zbiory i listy, 305
 relacje rodzinne, 303
 świat blokowy, 392
 świat Wumpusa, 307
 transport lotniczy, 390
 wymiana koła, 391

domeny
 ciągłe, 216
 dyskretne, 216
 modelu, 292
 niedeterministyczne, 411
 zadaniowe, 706
 dominacja, 599
 stochastyczna, 600, 616
 ścisła, 599
 domknięcie rezolucyjne, 264, 346
 domniemanie, 378
 dopasowywanie
 cytowań, 567
 wzorców, 329
 dopuszczalność, 111
 doskonała
 informacja, 177, 206
 równowaga Nasha, 682
 dośrodkowanie negacji, 340
 dowodzenie twierdzeń, 256
 przez rezolucję, 260, 318, 342, 344
 drzewo, 93, 95
 binarne, 99, 102, 105
 gry, 178
 optymalne decyzje, 183
 połączeń, 485
 wyszukiwawcze, 96, 178
 AND-OR, 151
 przycinanie, 115
 d-separacja, 469
 dualizm, 23
 dualność kwantyfikatorów, 300
 duracja, 367, 421
 dwukierunkowe wyszukiwanie, 121
 dylemat więźnia, 669
 dynamiczne sieci
 bayesowskie, DBN, 540
 decyzyjne, 631, 657
 dyskretyzacja, 148, 471
 dziedziczenie, 360
 dziedziczenie wielokrotne, 374
 dziedzina, 292

E

efekt, 389
 horyzontu, 190
 obramowania, 597

pewności, 597
 warunkowy, 415
 zakotwiczenia, 598
 efektywny czynnik rozgałęzienia, 124, 125
 ekonomia, 26
 eksperymenty psychologiczne, 19
 eksploracja, 60
 ekspresywność, 79
 eksternalia, 701
 element
 uczący, 76
 wykonawczy, 76
 elementy PEAS, 63
 eliminacja
 implikacji, 339
 koniunkcji, 258
 zmiennych, 480, 506, 545
 elitaryzm, 144
 empiryzm, 23
 estymacja stanu, 200, 283, 518, 552
 estymacja stanu logicznego, 279
 etyka deontologiczna, 24
 etykieta, 154
 eugenika, 144
 ewaluacja
 bazy wiedzy, 315
 sieci decyzyjnych, 605
 ewolucja, 146
 ewolucja maszynowa, 39

F

fakt, 265
 faktoryzacja, 261, 341
 fałszywy alarm, 572
 FIFO, First-In-First-Out, 96
 filtr
 najbliższego sąsiada, 573
 zakładanej gęstości, 554
 filtrowanie, 162, 518, 519, 552
 cząstek, 547, 573
 cząstek z Rao-Blackwellizacją, 552
 Kalmana, 532, 533, 536, 538, 552
 rozszerzone Kalmana, 539
 równanie rekurencyjne, 649
 fizyczny system symboli, 37
 fizyka kwalitatywna, 363, 386

fizykalizm, 23
 fluent, 274, 283, 369
 formułowanie problemu, 86
 funkcja, 290
 agenta, 56, 80
 akcja-użyteczność, 629
 całkowita, 294
 celu, 138, 177
 charakterystyczna, 689
 dobra społecznego, 702
 gęstości prawdopodobieństwa, 439
 heurystyczna, 108
 hiperboliczna, 660
 jednostki, 362
 kosztu, 34
 kosztu akcji, 85
 liniowa ważona, 188
 logistyczna, 474
 nagród, 657
 pochodzenia, 564
 porządkująca, 590
 potencjału, 630
 przystosowania, 144
 Q, 629
 wartościująca, 590
 wyboru społecznego, 702
 wypłaty, 177, 668
 wypukła, 149
 funkcje
 ewaluacyjne, 116, 187, 206
 w grach losowych, 198
 heurystyczne, 123, 133
 obliczalne, 26
 Skolema, 340
 użyteczności, 74, 177, 586, 589, 598, 605

G

gałąź, 92
 gaussowski model błędu, 542
 generator problemów, 77
 generowanie
 heurystyk, 126–128
 podproblemy, 127
 problemy rozluźnione, 126
 punkty orientacyjne, 128
 wyjaśnień, 381

globalne maksimum, 138
 głębokie uczenie, 45
 głębokość efektywna, 124
 głosowanie, 702
 pluralistyczne, 703
 z natychmiastową dogrywką, 704
 zatwierdzające, 704
 Gödel Kurt, 347
 gra
 brydż, 210
 poker, 210
 Reversi, 210
 scrabble, 210
 Sudoku, 222
 szachy, 208
 tryktrak, 210
 ultimatum, 705
 w spinacze, 687
 warcaby, 209
 gracz jałowy, 692
 gradient
 empiryczny, 148
 funkcji, 148
 graf, 86
 AND-OR, 153
 dualny, 218
 egzystencjalny, 373
 eulerowski, 174
 moralny, 469
 ograniczeń, 214, 330
 planistyczny, 397
 problemu rozluźnionego, 126
 przestrzeni stanów, 132, 178
 struktury koalicyjnej, 694
 szerokość drzewiasta, 237
 grid, 87
 gromadzenie wiedzy, 312
 gruntowanie, 562
 gry, 176
 asystenckie, 53, 686
 częściowo obserwowalne, 208
 dwuosobowe, 177
 karciane, 203
 kolejność ruchów, 185
 kooperatywne, 663, 689
 imputacja, 690
 indywidualna racjonalność, 690

obliczenia, 693
 rdzeń, 691
 stabilność koalicji, 691
 strategie, 690
 struktury koalicyjne, 689, 694
 wektor wypłat, 690
 wkład marginalny, 692
 wyniki, 689, 690
 losowe, 198
 niekooperatywne, 667
 o sumie zerowej, 177, 673
 postać ekstensywna, 681, 684
 postać sekwencyjna, 685
 powtarzane, 677
 sekwencyjne, 681
 stochastyczne, 196
 wideo, 210
 wieloosobowe, 181
 z częściową obserwowalnością, 200
 z naprzemiennymi ruchami, 667
 akcje, 668
 czysta strategia, 669
 funkcja wypłat, 668
 gracze, 668
 koncepcja rozwiązań, 669
 macierz wypłat, 668
 najlepsza odpowiedź, 670
 problem z koordynacją, 671
 profil strategii, 669
 punkt centralny, 671
 równowaga dominujących strategii, 670
 równowaga Nasha, 670
 strategia, 669
 strategia dominująca, 670
 strategia mieszana, 669
 wynik gry, 669
 grzbiec, 140

H

harmonogramowanie, 48, 421
 rozwiązywanie problemów, 422
 hesjan, 149
 heurystyczne wyszukiwanie alfa-beta, 187
 heurystyczny minimax, 187

heurystyka
 dokładność, 124
 dopuszczalna, 111, 398
 ignorowania list usuwania, 399
 ignorowania warunków wstępnych, 398
 minimalnego marginesu, 424
 minimalnych konfliktów, 231, 238
 różnicowa, 130
 rzędu, 227, 238
 spójna, 111
 wartości najmniej ograniczającej, 227, 238
 zabójcy, 186
 zerowego ruchu, 208
 heurystyki
 generowanie, 126–128
 niedopuszczalne, 115
 w planowaniu, 398
 wynajdywanie przez doświadczenie, 131
 hierarchia
 abstrakcji, 429
 taksonomiczna, 360
 hierarchiczne sieci zadaniowe, HTN, 403, 426
 hipergraf, 241
 hipoteza
 progu spełnialności, 272
 Sapira-Whorfa, 288
 HMM, Hidden Markov Models, 527
 homeostataza, 34
 horyzont decyzyjny, 625
 humanoidy, 47

I

identyfikacja zapytań, 312
 iloczyn punktowy, 481
 imputacja, 690
 indeks
 alokacji dynamicznej, 644
 Gittinsa, 644
 indeksowanie, 270
 indeksowanie predykatowe, 324
 indukcja, 23
 matematyczna, 347
 wsteczna, 677

indywidualna racjonalność, 690
informacje
 implementacja agenta, 609
 doskonałe, 607, 681
 niedoskonałe, 683
 teoria wartości, 606
 własności wartości, 609
 zbieranie, 610

instancja, 146
instancja podproblemu, 127
instancja
 ogólna, 319, 351
 szczegółowa, 319, 351
instrukcja zależności, 559
instrukcje ilościowe, 564
inteligencja, 17
inteligencja zbiorowa, 712
interfejsy „mózg-maszyna”, 28
interpretacja, 295
 obrazów, 49
 zamierzona, 295
introspekcja, 19
inżynieria
 komputerowa, 31
 ontologii, 356
 wiedzy, 309
irracjonalność, 619
iteracyjne zagłębianie, 186
 wyszukiwania, 104
iterowanie
 polityki, 637, 657
 wartości, 633, 657
 zbieżność, 634

J

jakościowe sieci probabilistyczne, 512, 601
jądro
 ergodyczne, 496
 przejścia, 496
język
 CLASSIC, 377
 formalny, 290
 logiki pierwszego rzędu, 291
 myśli, 288
 naturalny, 290
 PDDL, 389

programowania Prolog, 334
 reprezentowania wiedzy, 243, 283
języki programowania
 probabilistycznego, PPL, 556, 575, 580

K

kalmanowska macierz wzmocnienia, 537
kategorie, 359
kategorie naturalne, 364
kierunek
 diagnostyczny, 449
 przyczynowy, 449
klasy
 referencyjne, 459
 równoważności, 187
klasyfikacja, 376
 tekstu, 452
 umiejętności gracza, 561
klasyfikator bayesowski, 452
klauzula, 261
 Horna, 265, 283
 jednostkowa, 261, 269
klauzule
 definitywne, 265
 definitywne pierwszego rzędu, 325
kłątwa zwycięzcy, 619
klimatologia, 49
kodowanie ogólnej wiedzy, 313
kognitywistyka, 31
kolejka
 FIFO, 96
 LIFO, 96
 priorytetowa, 96
kolejność
 symboli i wartości, 270
 zmiennych i wartości, 226
kolektywne podejmowanie decyzji, 696
kolorowanie mapy, 213
komórka nerwowa, 29
kompozycja obiektu, 361
kompozycyjność, 288
komunikat
 o nominacji, 697
 ogłaszający zadanie, 697

komutatywność, 225
koncepty rozwiązań, 708
koniunkcja, 252, 258
koniunkt, 252
konkluzja, 379
konsekwencjalizm, 24, 58
konstrukcja if ... then ... else, 560, 567
kontury wyszukiwania, 113
korzeń, 92
koszt
 akcji, 88, 132, 156
 alternatywny, 648
 wyszukiwania, 125
krajobraz
 przestrzeni stanów, 138
 przydatności, 173
krata subsumpcyjna, 324
krotka, 293
kryptarytmy, 217
kryterium tylnych drzwi, 505
krytyk, 76
krzyżowanie, 143, 171
k-spójność, 221
kwadratowe układy dynamiczne, 173
kwantyfikatory
 ogólny, $\forall$, 297
 szczegółowy, $\exists$, 298
kwantyfikatory zagnieżdżone, 299

L

las, 236
lemat o liftingu, 345, 346
licytacja
 angielska, 698
 kombinatoryczna, 702
 ujawniająca prawdę, 699
 z rosnącą ofertą, 698
licytowanie dóbr deficytowych, 697
liczniki Bordy, 703
LIFO, Last-In-First-Out, 96
lifting, 322
lingwistyka, 34
lingwistyka obliczeniowa, 34
Lisp, 37
lista
 dodawania, 389
 usuwania, 389

literał negatywny, 251
 literały komplementarne, 261
 logicyzm, 20
 logika, 248
 domniemań, 378, 379, 382
 formalna, 25
 indukcyjna, 459
 linearna temporalna, 373
 modalna, 370–372
 niemonotoniczna, 259, 378, 382
 opisowa, 376, 382
 pierwszego rzędu, 243, 287, 315
 asercje, 303
 inżynieria wiedzy, 309
 kwantyfikatory, 297
 normalna postać koniunkcyjna, 339
 semantyka, 292
 składnia, 292
 symbole, 294
 wnioskowanie, 318
 zapytania, 303
 preferowanych modeli, 378
 probabilistyczna, 580
 rozmyta, 248, 291, 511
 temporalna, 292
 wyższego rzędu, 292
 losowe restarty, 270
 LRТА*, 169
 lukr składniowy, 306

Ł

łańcuchowanie
 progresywne, 266, 283, 318, 329, 351, 353
 progresywne przyrostowe, 330
 regresywne, 266, 283, 318, 332, 351
 łańcuchy Markowa, 493, 495, 515, 552
 luk trwałości, 543

M

macierz
 Hessego, 149
 obserwacji, 527
 wypłat, 668
 wzmocnienia, 537

makijaż genetyczny, 144
 makrooperatory, 429
 maksimum
 lokalne, 140
 spodziewanej użyteczności, 586
 maksymalizacja funkcji celu, 676
 marginalizacja, 444
 margines, 422
 maszyna
 Moore’a, 710
 skończenie stanowa, 677
 mat
 gwarantowany, 201
 probabilistyczny, 202
 przypadkowy, 202
 matematyka, 25
 materializm, 23
 maximin, 673
 MCMC, Markov Chain Monte Carlo, 493, 546
 MDP, Markov Decision Process, 623, 657
 medycyna, 49
 megawęzeł, 486
 menedżer zadania, 697
 mereologia, 385
 metawnioskowanie, 205
 metoda
 Newtona-Raphsona, 149
 szukania liniowego, 148
 ścieżki krytycznej, 422
 metody próbkowania bezpośredniego, 487
 miara, 362
 miara wydajności, 58, 62, 80, 245
 migotanie stron, 121
 mikromort, 591
 minimalizacja logiczna, 362
 minimax, 180
 heurystyczny, 187
 oczekiwany, 197, 206
 minimum globalne, 138
 mnożenie punktowe, 481
 model, 248, 316, 371
 aktorski, 709
 czasu dyskretnego, 514
 czujników, 72
 expit, 474

koneksjonistyczny, 42
 konkurencji niedoskonałej, 685
 logiki pierwszego rzędu, 292
 obserwacyjny, 516
 poznawczy, 19
 probabilistyczny, 436, 575
 otwartego świata, OUPM, 580
 otwartego wszechświata, OUPM, 562, 563
 system NET-VISA, 568
 wnioskowanie, 566
 przejścia, 85, 88, 151, 156, 274, 515, 552, 657
 przejściowych usterek, 543
 sensoryczny, 515, 649
 standardowy, 21
 świata, 72
 trwałych awarii, 543
 ukryty Markowa, 526, 552
 wspomaganie, 620
 Modus Ponens, 258, 351
 momenty, 367
 monitorowanie, 162
 akcji, 418
 aktywności, 418
 celu, 418
 planu, 418
 przestrzegania traktatów
 antynuklearnych, 568
 ruchu drogowego, 574
 monotoniczne ustępstwa, 707
 monotoniczność, 259, 588
 logiki, 378
 preferencji, 592
 możliwy percept, 159
 multiplekser, 560
 multiplikatywna funkcja użyteczności, 603

N

nagradzany proces Markowa, 642
 nagroda, 78
 addytywna, 625
 potencjalna, 630
 średnia, 627
 naiwny model bayesowski, 451, 452, 458

najbardziej
prawdopodobne wyjaśnienie, 509
prawdopodobny ciąg, 524
wiarygodne wyjaśnienie, 518, 552
najlepsza odpowiedź
iterowana, 673
krótkowzroczna, 673
najmniejsza liczba pozostałych
wartości, 238, 329
największa oczekiwana użyteczność, 435
nastawienia propozycyjne, 370
następnik, 252
nasylenie, 346
nawigowanie robotów, 91
nawracanie
chronologiczne, 228
skokowe, 229
nawroty
dynamiczne, 240
inteligentne, 270
ukierunkowane na konflikty, 229,
238
w oparciu o zależności, 240
negocjowanie, 704
w domenach zadaniowych, 706
z naprzemiennymi ofertami, 705
neuron, 28
neuronauka, 28
niechęć do wieloznaczności, 597
niedoskonała informacja, 206
niekompetencja
praktyczna, 433
teoretyczna, 433
niemonotoniczność, 378
niepewność, 432, 433, 434
czyichś preferencji, 613
egzystencjalna, 563
relacyjna, 560
tożsamości, 563
wypląt, 686
nieporadność, 433
nieprzezroczystość referencyjna, 370
nierówność trójkąta, 111, 113
niespełnialność, 283
niezależne podproblemy, 233
niezależność, 446
bezwzględna, 457
kontekstowa, 470, 559

marginalna, 447
podproblemów, 401
preferencyjna, 602
preferencyjna wzajemna, 602
użytecznościowa, 603
użytecznościowa wzajemna, 603
warunkowa, 450, 458, 468, 506
norma maksymalna, 635
normalna postać koniunkcyjna, 262, 339
notacja
„dużego 0”, 83
Backusa-Naura, 252
przedrostkowa, 306
wrostkowa, 306
NP-zupełność, 26
numer Gödla, 347

O

obiekty, 290, 359, 365, 369
gwarantowane, 571
mentalne, 370
wewnętrzne, 366
zewewnętrzne, 366
złożone, 361
obliczalność efektywna, 26
obliczanie
indeksów Gittinsa, 644
punktów równowagi, 673
rozkładu prawdopodobieństw
a posteriori, 518
stanu przekonań, 518
wiarygodności, 521
obliczenia
autonomiczne, 81
kwantowe, 32
w grach kooperatywnych, 693
obrazowanie mózgu, 19
ocena stanu, 162, 279
oczekiwana
użyteczność, 75, 586, 594
użyteczność loterii, 589
wartość pieniężna, 592
odcinanie, 189
odmładzanie, 144
odwrotny model logitowy, 474
odwroćcie dowodów, 554
odwzorowanie zwięzające, 634

oferta, 697
ograniczenia
czasowe, 421
dwuargumentowe, 216
eliminujące symetrię, 237
globalne, 217, 221
jednoargumentowe, 216
pierwszeństwa, 215
preferencji, 218
racjonalnych preferencji, 587
trójargumentowe, 216
współbieżnych akcji, 666
ograniczona racjonalność, 21
ograniczoność zasobów, 421
ograniczony PlanSAT, 431
ontologia, 356, 358
domeny, 310
górna, 357, 382
operator do, 503
operatory modalne, 371
opis instancji problemu, 314
optogenetyka, 28
optymalizacja, 172
wypukła, 149, 171
z ograniczeniami, 149
optymalność
kosztowa, 97
w sensie Pareto, 672
optymizm w warunkach niepewności,
169
osadzenie, 251
osobliwość, 29

P

pamięć, 101
pamięć RAM, 121
paradoks Condorceta, 702
paramodulacja, 349, 352, 354
partner, 662
partycja, 360
PEAS, 63
pełne sformułowanie stanu, 139
percept, 56, 159, 412
perceptron, 39
pętla
otwarta, 85
zamknięta, 85

- plan warunkowy, 150, 151
- planowanie
 - bezczujnikowe, 411, 413
 - hierarchiczne, 402
 - hierarchiczne z wyprzedzeniem, 411
 - HTN, 404
 - klasyczne, 397
 - abstrahowanie stanów, 401
 - algorytmy, 393
 - heurystyki, 398
 - spełnialność boolowska, 396
 - klasyczne, 388
 - linearne, 427
 - Monte Carlo częściowo
 - obserwowalne, 656
 - online, 411, 418
 - oparte na przypadkach, 429
 - ponowne, 411
 - produkcji, 215, 421
 - przez wnioskowanie, 280
 - warunkowe, 411, 417
 - wieloagentowe, 663, 708
 - kooperacja, 667
 - koordynacja, 667
 - wielozłożonowe, 662
 - wieloeffektorowe, 662
 - z wyprzedzeniem, 83
 - zgodne, 411, 429
- PlanSAT, 431
- plany
 - bezczujnikowe, 426
 - częściowo uporządkowane, 397
 - warunkowe, 426
- płaskowyż, 140
- pobocze, 140
- podjmowanie decyzji
 - kolektywne, 696
 - prostych, 585
 - strategicznych, 662
 - w środowisku wieloagentowym, 661
 - złożonych, 621
- podejście
 - deklaratywne, 244
 - proceduralne, 244
- podgra, 682
- podgraf rodowy, 469
- podkategoria, 359
- podproblemy, 127
- podstawienie szczegółowe, 319
- pojazdy zrobotyzowane, 47
- pokrycie Markowa, 469
- pole średnie, 510
- polidrzewa, 484
- polityka, 192, 430
 - algorytm iterowania, 637
 - aproksymacja, 646
 - dominująca, 648
 - iterowanie, 657
 - optymalna, 623, 627
 - rozgrywek, 193
 - selekcji, 193
- połączone komponenty, 233
- poprzednik, 252
- porządek preferencji społecznych, 702
- postępowanie właściwe, 21
- posunięcie, 179
- poziom abstrakcji, 86
- pozyskiwanie wiedzy, 310
- pozytywizm logiczny, 23
- pożyteczna sztuczna inteligencja, 620
- prawda, 248
- prawdopodobieństwo, 20, 198, 436, 457
 - a posteriori, 457, 530
 - a priori, 437
 - akceptacji, 499
 - bezwartunkowe, 437, 457
 - marginosowe, 444
 - wartunkowe, 437, 457
- prawidłowość, 257
- prawo Moore'a, 32
- precyzowalność zstępująca, 406
- predykaty interwałów czasowych, 368
- preferencje
 - bez niepewności, 602
 - w warunkach niepewności, 603
- preobliczenia, 128
- probabilistyczne bazy danych, 581
- problem, 132
 - bandyty, 641, 659
 - bezczujnikowy, 155
 - binarny, 216
 - goryli, 52
 - gridowy, 87
 - jednorękiego bandyty, 643
 - komiwojażera, 91
 - króla Midasa, 52
- kwalifikacji, 277, 432
- najtańszej sekwencji testowej, 610
- optymalizacji z ograniczeniami, 218
- ośmiu hetmanów, 139, 145
- partycjonowania zbioru, 694
- planowania produkcji, 421
- pokrycia zbioru, 399
- poszukiwacza skarbów, 610
- projekcji czasowej, 286
- ramy, 275, 286
- rozluźnione, 126
- selekcji, 647
- spadku gradientowego, 138
- sporządzania mapy, 164
- stopu, 320
- topografii, 91
- wielorękiego bandyty, 621
- wspinaczki, 138
- wyłączalnego robota, 686
- wyrównywania wartości, 21
- wyszukiwania, 85
- zgodny, 155
- zliczeniowo P-zupełny, 485
- problemy
 - bezczujnikowe, 171
 - CSP, 232
 - z domenami ciągłymi, 216
 - z ograniczeniami nieliniowymi, 216
- decyzyjne
 - metapoziomowe, 647
 - sekwencyjne, 621
- eksploracyjne, 171
- MDP, 630
 - algorytmy, 633, 651
 - algorytmy online, 639, 655
 - częściowo obserwowalne, POMDP, 649
 - iterowanie wartości, 651, 654
 - relacyjne, 658
 - reprezentacje czynnikowe, 658
- optymalizacyjne, 137
- rozluźnione, 398
- SAT, 272
- słabo ograniczone, 272
- spełniania ograniczeń, CSP, 212, 238, 397
- standaryzowane, 87

- turystyczne, 91
 w częściowo obserwowalnych
 środowiskach, 160
 wyszukiwania online, 165
 procedura wnioskowania, 255
 proces
 decyzyjny Markowa, MDP, 623, 657
 inżynierii wiedzy, 309
 jednorodny czasowo, 516
 Markowa pierwszego rzędu, 515
 program
 agenta, 56, 67, 80
 General Problem Solver, 24
 programowanie
 dynamiczne, 128, 337, 623
 nieseryjne, 508
 w czasie rzeczywistym, RTDP, 640
 genetyczne, 39, 143, 173
 liniowe, LP, 149, 171, 173, 638
 probabilistyczne, 555
 interpretowanie tekstu, 575
 w logice, 333, 352
 w logice z ograniczeniami, 239, 338
 zorientowane obiektowo, OOP, 374
 programy
 generatywne, 576, 580
 modele Markowa, 578
 wnioskowanie, 579
 wyniki wnioskowania, 577
 logiki probabilistycznej, 580
 projektowanie
 agentów, 663
 mechanizmów, 663
 Prolog, 352
 ograniczenia, 338
 redundantne wnioskowania, 335
 semantyka bazodanowa, 337
 zapętlenie, 335
 propagacja
 ankiety, 285
 jednostek, 269
 ograniczeń, 218
 wsteczna, 193
 propozycjonalizacja, 320, 351, 562
 protokół monotonicznych ustępstw, 707
 próbkowanie
 adaptacyjne, 509
 bezpośrednie, 487
 blokowe, 498
 Gibbisa, 493, 496, 497
 istotnościowe, 490
 logiczne, 509
 Metropolisa-Hastingsa, 499
 Monte Carlo łańcuchami Markowa,
 493, 506
 Thompsona, 647
 ważne, 490
 z odrzucaniem, 488
 przechodniość, 588
 przechowywanie, 324
 przedział ufności, 646
 górna granica, 646
 przedziały, 367
 przeglądanie, 191
 przekleństwo optymalizacji, 596, 619
 przemianowanie, 327
 przeplatanie, 427
 przesłanka, 252
 przestrzenie bezpiecznie
 eksplorowalne, 166
 przestrzeń
 stanów, 85
 metapoziomowa, 130
 obiektowa, 131
 sześciowymiarowa, 147
 zdarzeń elementarnych, 436
 przetarg, 699
 drugiej ceny, 699
 pierwszej ceny, 699
 przetwarzanie języka naturalnego, 18,
 34
 przewidywanie, 159, 518, 519, 552
 przezroczystość referencyjna, 370
 przycinanie, 177
 alfa-beta, 182
 nieużytków, 208
 niezależne od domeny, 400
 wyprzedzające, 190
 przypadek warunkujący, 463
 przypisanie, 213, 502
 częściowe, 213
 pełne, 213
 spójne, 213
 przypuszczenia, 381
 przystawka proceduralna, 375
 psychologia, 30
 ewolucyjna, 598
 poznawcza, 30
 punkt
 krzyżowania, 143
 stały, 328
 punkty
 orientacyjne, 129
 wewnętrzne, 173
- ## Q
- Q-funkcja, 605
- ## R
- rachunek
 hedonistyczny, 617
 prawdopodobieństwa, 25, 292, 616
 predykatów pierwszego rzędu, 287
 sytuacyjny, 397
 wariacyjny, 172
 zdań, 243, 251, 283
 agent, 273
 dowodzenie twierdzeń, 256
 planowanie przez
 wnioskowanie, 280
 semantyka, 252
 składnia, 251
 sprawdzanie modeli, 268
 wnioskowanie, 318
 wyszukiwanie lokalne, 270
 zdarzeń, 366
 racjonalne
 decyzje, 434
 działanie, 53
 preferencje, 587, 589
 racjonalność, 17, 59
 ramki, 385
 ramki Minsky'ego, 41
 randomizacja, 71
 randomizowane badania kontrolne,
 506
 Rao-Blackwellizacja, 554, 573
 rating, 561
 rdzeń gry kooperatywnej, 691

redukcja
 do wnioskowania, 320
 słabszych ruchów, 191
 redukowanie symetrii, 400
 redundantne ścieżki, 96
 reguła, 252
 „warunek-akcja”, 70
 Bayesa, 448, 450, 457
 łańcuchowa, 465
 reguły wnioskowania, 258, 283
 reifikacja, 359
 rekombinacja, 143
 rekurencyjna estymacja, 519
 rekurencyjne równanie filtrowania, 649
 relacja, 213
 osiągalności, 371
 wynikania, 248, 283
 relacje
 binarne, 375
 między obiektami, 290
 zakłócanie, 470
 relacyjne modele probabilistyczne,
 556, 557, 580
 wnioskowanie, 561
 relewancja zmiennych, 483
 reprezentacja, 287
 atomowa, 78
 czynnikiowa, 79, 389, 438
 lokalistyczna, 80
 nieparametryczna, 471
 rozproszona, 80
 strukturalna, 79, 438
 wiedzy, 18, 34
 retrogradacja, 192
 rewizja przekonań, 380
 rezolucja, 339
 binarna, 341
 jednostkowa, 261, 349
 pełna, 261
 zaprzeczeniowo zupełna, 344
 zupełność, 344
 rezolwenta, 260
 robot, 18
 rozcięcie cyklu, 235
 rozczarowanie podecyzyjne, 594, 619
 rozgrywanie gier, 48
 rozgrywka, 192

rozkład
 Bernoullego, 438
 dyskretny logarytmiczno-normalny,
 564
 Gaussa
 aktualizowanie, 534
 liniowy, 533
 warunkowy, 473
 wielowymiarowy, 473, 533
 gruboogonowy, 172
 kanoniczny, 469
 Poissona, 564
 prawdopodobieństwa, 439
 wspólny, 440
 wspólny pełny, 441, 457
 rzędu wielkości, 564
 stacjonarny, 496, 520
 warunkowy, 469
 rozluźniony plan, 400
 rozpoznawanie
 mowy, 48
 planu, 667
 rozproszone spełnianie ograniczeń, 241
 rozrost, 193
 rozszerzenie, 379
 rozszerzone interpretacje, 316
 rozwiązanie
 cykliczne, 154
 optymalne, 86
 rozwiązywanie problemów
 harmonogramowania, 422
 rozproszone poprzez kooperację, 709
 w częściowo obserwowalnych
 środowiskach, 160
 za pomocą wyszukiwania, 83
 rozwijanie, 545, 562
 równania
 diofantyczne, 238
 korekty, 505
 równanie
 Bellmana, 628
 strukturalne, 502
 równość, 301, 348
 równość logiczna, 252
 równowaga
 Bayesa-Nasha, 686
 maksymimalna, 676

Nasha, 670, 709
 szczegółowa, 496
 równoważnik pewności, 593
 równoważność logiczna, 256
 różnica symetryczna, 254
 rzadkie modele temporalne, 540
 rzeczy, 365
 rzut, 192

S

SAT, 257
 schemat, 146
 schemat akcji, 389
 sekwencyjne
 metody Monte Carlo, 554
 problemy decyzyjne, 621
 próbkowanie istotnościowe, 546, 554
 z postarzaniem, 547
 selekcja, 143, 193
 semantyka, 248, 252, 283
 angeliczna, 406
 bazodanowa, 301, 302
 sfaktoryzowana granica, 554
 sieci bayesowskie, 44, 461, 462
 aproksymowane wnioskowanie, 486
 czasu ciągłego, 553
 dynamiczne, DBN, 540, 550, 552
 aproksymowane wnioskowanie,
 546
 konstruowanie, 541
 ściśle wnioskowanie, 545
 hybrydowe, 472, 506
 konstruowanie, 465
 próbkowanie Gibbsa, 493
 próbkowanie z odrzucaniem, 488
 reprezentowanie akcji, 503
 semantyka, 464
 ściśle wnioskowanie, 477
 uporządkowanie węzłów, 466
 warunkowa niezależność, 468
 wieloelementowe, 582
 ze zmiennymi ciągłymi, 471
 sieci
 decyzyjne, 585, 603, 616
 ewaluacja, 605
 kontraktowe, 696

- logiczne Markowa, 581
 - Markowa, 508
 - przekonań, 461
 - przyczynowe, 461, 501, 506
 - semantyczne, 373–375, 382
 - wielokrotnie połączone, 484
 - wkładów marginalnych, 693
 - skalowanie nagród, 629
 - skierowana spójność łukowa, 233
 - składnia, 283
 - skolemizacja, 319, 340
 - słabe metody, 40
 - słownik predykatów, funkcji i stałych, 312
 - softbot, 63
 - solidne decyzje, 611
 - sortowanie topologiczne, 233
 - spacer losowy, 168, 534
 - spamiętywane programowanie
 - w logice, 337, 353
 - spełnialność, 257
 - spełnialność boolowska, 396
 - spójnik logiczny, 297
 - spójność, 376
 - heurystyki, 111
 - łukowa, 219
 - ścieżkowa, 220
 - wierzchołkowa, 219
 - sprawdzanie
 - modeli, 250, 268, 283
 - w przód, 227
 - występowania, 323
 - stała Skolema, 319
 - stan, 87, 156, 389
 - agenta odkurzacza, 150
 - docelowy, 85, 88
 - końcowy, 177
 - początkowy, 85, 87, 132, 156, 206, 390
 - przekonań, 150, 171, 206, 279, 432, 606, 649, 683
 - standardowa loteria, 591
 - standaryzacja
 - nazw zmiennych, 395
 - zmiennych, 322, 340
 - statystyka, 25
 - statystyka porządkowa, 595
 - sterowanie
 - optymalne, 172
 - rozmyte, 511
 - stopień
 - potwierdzania, 459
 - prawdy, 291
 - przekonania, 433
 - strata polityki, 636
 - strategia, 150, 201, 669
 - dominująca, 698
 - DOVE, 678
 - HAWK, 678
 - TIT-FOR-TAT, 677
 - typu A, 186
 - typu B, 186
 - Zeuthena, 708
 - strategie
 - ewolucyjne, 143, 173
 - rezolucyjne, 349
 - w grach kooperatywnych, 690
 - struktura preferencji, 601
 - struktury danych, 95
 - stwierdzenia obserwacyjne, 23
 - stwierdzenie, 436
 - substancje, 365
 - subsumpcja, 376
 - sukces przypadkowy, 420
 - superproces bandyty, 647
 - sybille, 563
 - sygnatura typu, 558
 - sylogizm, 19
 - symbol, 251
 - symbol równości, 301
 - symbole
 - funkcyjne, 294
 - predykatowe, 294
 - stałe, 294
 - symetria wartości, 237
 - symulacja, 192, 193
 - symulacja łańcuchów Markowa, 493
 - symulowane wyżarzanie, 141, 171
 - synchronizacja doskonała, 665
 - synteza dedukcyjna, 351
 - system Vickreya, 699
 - systemy
 - eksperckie, 40
 - planowania
 - portfolio, 425
 - reaktywnego, 430
 - produkcyjne, 331, 353
 - rekomendacji, 48
 - rzadkie, 466
 - utrzymywania prawdy, 240, 380
 - przypuszczeniowe, 381
 - uzasadnieniowe, 381
 - wieloagentowe, 82, 661
 - zachowywania spójności logicznej, 380
 - szeregowalne podcele, 400
 - szerokość
 - drzewa, 238, 485
 - drzewa wynikowego, 237
 - drzewiasta grafu, 237
 - hiperdrzewa, 241
 - indukowana, 241
 - sztuczna
 - inteligencja, AI, 17, 35
 - superinteligencja, ASI, 52
 - sztuczne życie, 173
 - szybkość mieszania, 498
- Ś**
- ścieżka, 86, 132, 422
 - ścieżka krytyczna, 422
 - śląd wykonania, 576, 580
 - śledzenie wielocelowe, 570
 - ślepe zaułki, 165
 - średnica grafu, 103
 - środowisko, 245
 - częściowo obserwowalne, 171
 - deterministyczne, 64
 - dyskretne, 65
 - epizodyczne, 65
 - jednoagentowe, 64
 - kooperatywne, 64
 - niedeterministyczne, 171
 - półdynamiczne, 65
 - rozpoznane, 65
 - rywalizacji, 64, 176
 - statyczne, 65
 - stochastyczne, 622
 - wieloagentowe, 661
 - zadaniowe, 62, 63, 66, 80
 - świadeństwo, 437
 - świat
 - blokowy, 38
 - Wumpusa, 244, 307, 453

T

tabela prawdy, 253
 tablica
 prawdopodobieństwa warunkowego, 463
 rzadka, 630
 transpozycji, 186
 taksonomia, 360
 tautologia, 257
 technika obserwowanych literałów, 285
 teoria
 deskryptywna, 596
 dowodów, 511
 ewolucji, 146
 gier, 27, 176
 konceptje rozwiązań, 708
 kooperatywnych, 689, 709
 niekooperatywnych, 667, 709
 liczb naturalnych, 305
 mnogości, 306
 możliwości, 512
 normatywna, 596
 NP-zupełności, 26
 optymalizacji, 172
 podejmowania decyzji, 27, 435, 457, 596, 616
 potwierdzania, 23
 sterowania, 33
 syntaktyczna, 385
 użyteczności, 434, 587, 616
 wartości informacji, 606
 wieloatrybutowa użyteczności, 598
 wyboru społecznego, 702
 zbiorów, 306
 term, 296
 term podstawowy, 319
 test
 terminalny, 177, 206
 Turinga, 18
 testowanie celu, 156
 tłumaczenie maszynowe, 48
 tomografia sieciowa, 508
 totalny test Turinga, 18
 tragedia wspólnot, 701
 transpozycje, 186
 Turing Alan, 18
 twierdzenie, 304

Arrowa, 703
 Gibbarda-Satterthwaite'a, 704
 Gödla o niezupełności, 347
 Herbranda, 346
 o dedukcji, 257
 o kształtowaniu, 629
 o niezupełności, 25
 o równoważności przychodów, 700
 o zbieżności perceptronu, 39
 potoczne Nasha, 680
 zasadnicze o rezolucji, 264, 345

U

UCB1, 194
 uczenie, 251
 agenta online, 170
 hebbowskie, 35
 głębokie, 45
 maszynowe, 18, 42
 metapoziomowe, 131
 się, 60, 518
 klauzul konfliktujących, 270
 ograniczeń, 230, 238, 240
 przez wyjaśnianie, 429
 z odwrotnym wzmocnieniem, 53
 ze wzmocnieniem, 657, 711
 ujawnianie preferencji, 590
 układ
 konfliktowy, 705
 wielkiej skali integracji, 91
 ukryte modele Markowa, HMM, 43, 526, 552
 unifikacja, 318, 321, 322, 351, 352
 unifikacja równościowa, 349
 unifikator, 322
 unifikator najbardziej ogólny, 323
 uniwersum Herbranda, 345
 uogólnienie, 537
 Uogólniony Modus Ponens, 321, 351
 uporządkowanie koniunktów, 329
 urzeczowienie, 359
 usterka detekcji, 572
 utrzymywanie spójności łukowej, 228
 utylitaryzm, 24
 uzasadnienie, 379
 uzupełnienie kwadratu, 535

użyteczność, 24, 74, 589
 adaptacyjna, 620
 oczekiwana, 594
 pieniędzy, 592
 skalowanie, 590
 stanów, 627
 szacowanie, 590
 wielokryterialna, 601, 619
 znormalizowana, 591

W

waga, 490
 wariacyjne metody aproksymacji, 510
 wartości
 domyślne, 375
 logiczne, 283
 wartość
 informacji, 606, 617
 informacji doskonałej, 607
 materialna, 188
 minimaksowa, 179
 najmniej ograniczająca, 227, 238
 oczekiwana, 188
 Shapleya, 692
 warunek
 dwustronny, 252
 wstępny, 379, 389
 warunkowanie
 przekrojów, 234, 238
 rozcięć, 235
 ważenie
 ograniczeń, 232
 wiarygodności, 491, 506
 ważone zliczanie modeli, 485
 wektor wypłat, 690
 weryfikacja celu
 późna, 99
 wczesna, 98
 węzeł, 92, 95
 AND, 151
 OR, 151
 węzły
 decyzyjne, 604
 deterministyczne, 470
 losowe, 196, 206, 604
 macierzyste, 92, 95

- niewątpliwie rozwinięte, 114
- osiągnięte, 92
- potomne, 92
- przecieku, 470
- użyteczności, 604
- wiarygodność, 492
- wiązanie rekordów, 581
- wiązka, 362
- widzenie komputerowe, 18
- wiedza podstawowa, 243
- wielkość kroku, 148
- wielokryterialna teoria użyteczności, 616
- wieloużywalność, 421
- wielozadaniowość, 647
- większość bezwzględna, 704
- wkład marginalny gracza, 692
- właściwości, 290
- wnioskowanie, 238, 243, 255, 258, 283
 - aproksymowane, 486, 500, 546
 - dla kategorii, 373
 - domyślne, 511
 - logiczne, 242, 250
 - na podstawie
 - domniemań, 378
 - dowodów, 450
 - od celu do faktów, 267
 - od faktów do celu, 266
 - probabilistyczne, 42, 461
 - probabilistyczne w czasie, 513
 - progresywne, 266
 - przestrzenne, 387
 - przez
 - dziedziczenie, 374
 - enumerację, 478
 - rezolucję, 341
 - symulację łańcuchów Markowa, 493
 - psychologiczne, 387
 - redundantne, 335
 - regresywne, 267
 - sterowane danymi, 267
 - ściśle, 477, 484
 - w sieciach DBN, 545
 - ukierunkowane na cel, 267
 - w logice pierwszego rzędu, 318
 - w modelach OUPM, 566
 - w modelach temporalnych, 518
 - w programach generatywnych, 579
 - w rachunku zdań, 318
 - w relacyjnych modelach
 - probabilistycznych, 561
 - z pełnych wspólnych rozkładów, 444
- wspinaczka
 - pierwszego wyboru, 141
 - stochastyczna, 141
 - z losowymi restartami, 141
- wspólny plan, 665
- współbieżność, 664
- współbieżność prawdziwa, 665
- współczynnik
 - dyskonta, 625, 706
 - Giniego, 672
 - konkurencyjności, 165
 - objazdu, 115
- wszechwiedza, 60
- wszechwiedza logiczna, 372
- wybór zwykłą większością, 703
- wydajność algorytmu, 124
- wygładzanie, 518, 521, 552
 - dwukierunkowe, 523
 - z ustalonym opóźnieniem, 524
- wykonywanie naprzemienne, 664
- wymiatanie priorytetowe, 658
- wymuszenie stanu, 155
- wynik, 206
 - gry, 689, 690
 - społeczny, 702
- wynikanie, 283
- wyrażenie lambda, 296
- wysumowywanie, 444
- wyszukiwania, 83, 85, 324
 - kontury, 113
 - porównanie algorytmów, 108
 - struktury danych, 95
 - uczenie maszynowe, 130
- wyszukiwanie
 - „najpierw najlepszy”, 94, 109, 118
 - A*, 111, 112, 133
 - dwukierunkowe, 133
 - optymalnie wydajne, 115
 - ważone, 115, 133
 - z iteracyjnym zagłębianiem, 118, 133
 - z uczeniem w czasie rzeczywistym, 169
- alfa-beta, 185, 206, 207
- AND-OR, 171
- antagonistyczne, 176
- bez obserwacji środowiska, 155
- bez ograniczenia kosztu, 117
- drzewiaste, 97
- drzewiaste Monte Carlo, 206
- dwukierunkowe, 106, 132
- grafowe, 96
- heurystyczne alfa-beta, 187
- heurystyczne dwukierunkowe, 121
- hierarchiczne, 405
- lokalne, 137
 - na usługach CSP, 230
 - online, 168
 - skupione, 142
 - w przestrzeniach ciągłych, 147
 - wiązkowe, 142
 - zachłanne, 139
- minimaksowe, 179, 180
- Monte Carlo, 192
- najkrótszej ścieżki, 120
- niedoinformowane, 98
- od czoła do celu, 123
- od czoła do czoła, 123
- online, 164
- poinformowane, 108, 133
- przy jednolitym koszcie, 100, 132
- przyrostowe, 171
- RBFS, 120
- rekurencyjne, 118
- rekurencyjne „najpierw najlepszy”, 133
- satysfakcjonujące, 115
- skupione, 118, 133
- sterowane wiedzą, 108
- stochastyczne skupione, 143
- szybkie, 117
- tabu, 172
- uspokajające, 189
- w czasie rzeczywistym, 175
- w głąb, 101, 102, 132
- w głąb ograniczone, 132
- w przód w planowaniu, 393
- w środowiskach częściowo obserwowalnych, 155
- wiązkowe, 118, 133
- wspinaczkowe, 138, 168, 171

wyszukiwanie
 wstecz w planowaniu, 395
 wszcz, 98, 132
 z częściowymi obserwacjami, 159
 z iteracyjnym zagłębianiem, 103, 132
 z nawrotami, 103, 224
 z niedeterministycznymi akcjami, 150
 z odcinaniem, 189
 z ograniczeniami
 kosztu, 117
 pamięciowymi, 117
 głębokości, 103
 z ograniczoną nieoptymalnością, 117
 z wnioskowaniem, 227
 z zagłębianiem, 105
 zachłanne, 109
 zachłanne „najpierw najlepszy”, 133, 132
 wyżarzanie, 141
 wznawiane MDP, 645
 wzniesienie, 138

Z

zabójcze ruchy, 186
 zaburzenia, 503
 zachłanne wyszukiwanie drzewiaste, 110
 zadania, 421
 zakodowanie opisu, 314
 zakres, 438

założenie
 Markowa, 515, 552
 o otwartości świata, 301, 413
 o unikalności nazw, 301
 o zamkniętości domen, 301, 556
 o zamkniętości świata, 301, 382, 413
 sensoryczne Markowa, 516
 zapętlenie, 335
 zapętłona propagacja przekonań, 510
 zaprzeczenie, 251
 zapytania, 266, 303, 444
 zapytania dla procedury wnioskowania, 314
 zasada
 obojętności, 459
 racji niedostatecznej, 459
 rezolucji, 260, 283
 ujawniania, 699
 włączania-wykluczania, 441
 zasoby, 426
 zastępowalność, 319, 588
 zbieranie informacji, 60
 zbieżność iterowania wartości, 634
 zbiór
 konfliktowy, 229
 magiczny, 332, 353
 negocjacyjny, 705
 osiągalności, 406
 rozmyty, 511
 zadań, 421
 zdań, 243

zdania, 283
 atomowe, 251, 296
 złożone, 251, 297
 zdarzenia, 366, 436
 zejście gradientowe, 141
 złożoność
 algorytmów, 83, 100, 101
 czasowa, 97, 103, 104, 108
 danych, 329
 pamięciowa, 97, 103, 108
 próbkowania Gibbsa, 497
 ścisłego wnioskowania, 484
 zmienne, 79
 ciągłe, 471
 dowodowe, 514
 ilościowe, 565
 losowe, 558, 565
 losowe indeksowane, 581
 niemodelowane, 503
 ponadczasowe, 274
 relewancja, 483
 stanu, 514
 ukryte, 475
 zmowa cenowa, 698
 znakowanie wsteczne, 240
 zobowiązania
 epistemiczne, 291, 292, 315, 434
 ontologiczne, 291, 315
 zupełność, 97, 250, 588
 zupełność rezolucji, 264, 344

PROGRAM PARTNERSKI

— GRUPY HELION —

- 
1. ZAREJESTRUJ SIĘ
 2. PREZENTUJ KSIĄŻKI
 3. ZBIERAJ PROWIZJĘ

Zmień swoją stronę WWW w działający bankomat!

Dowiedz się więcej i dołącz już dzisiaj!

<http://program-partnerski.helion.pl>

GRUPA
Helion 

Sztuczna inteligencja: to się staje na naszych oczach!

Sztuczna inteligencja budzi zachwyt i kontrowersje. W porównaniu z innymi gałęziami nauki jest stosunkowo młoda: liczy około siedemdziesięciu lat, mimo że czerpie ze znacznie starszych idei. Jednak błyskawiczny rozwój sztucznej inteligencji i przeobrażanie osiągnięć nauki w działające technologie sprawiają, że wyrobienie poglądu na całokształt tej dziedziny jest trudnym zadaniem. Warto więc spojrzeć na historię rozwoju sztucznej inteligencji z perspektywy jej współczesnych osiągnięć i dzięki temu lepiej zrozumieć, czym ta nauka jest w swojej istocie i dokąd podąża.

Oto pierwszy tom dzieła, które stanowi inspirujące spojrzenie na sztuczną inteligencję. Jego zrozumienie nie wymaga wybitnej znajomości informatyki i matematyki. Książka jest wspaniałą syntezą wczesnych i późniejszych koncepcji, a także technik, przeprowadzoną we frameworku idei, metod i technologii. Zawiera ogrom rzetelnej wiedzy przekazanej w niezbyt sformalizowany sposób. Opisy, formuły matematyczne i algorytmy, pokazane w postaci czytelnego pseudokodu, cechują się przejrzystością i precyzją. Zaprezentowano tu wszystkie ważne idee i koncepcje sztucznej inteligencji, zgodnie z najnowszymi trendami i osiągnięciami.

W tomie pierwszym między innymi:

- koncepcje sztucznej inteligencji
- różne podejścia do rozwiązywania problemów z wykorzystaniem sztucznej inteligencji
- reprezentacja wiedzy i modelowanie, a także wyszukiwanie i planowanie
- wnioskowanie w warunkach niepewności
- podejmowanie złożonych decyzji, również w środowisku wieloagentowym

Stuart Russell jest profesorem na Uniwersytecie Kalifornijskim w Berkeley, dyrektorem Center for Human-Compatible AI i profesorem inżynierii z ramienia fundacji Smitha-Zadeha. Laureat wielu prestiżowych nagród. Autor kilkuset publikacji dotyczących sztucznej inteligencji.

Peter Norvig jest dyrektorem do spraw badań w Google i członkiem kilku amerykańskich stowarzyszeń akademickich. Był szefem Wydziału Nauk Obliczeniowych NASA Ames Research Center. Napisał kilka cenionych książek dotyczących praktycznych aspektów sztucznej inteligencji.

Helion	KOD KORZYŚCI Sięgnij po więcej ►	
 helion.pl		
 HELION SA ul. Kościuszki 1c 44-100 Gliwice tel.: 32 230 98 63 helion@helion.pl		
ISBN 978-83-283-7608-3		
		
9 788328 376083		
Cena: 169,00 zł		