

Wydanie IV

Profesjonalne programowanie w Pythonie

Poznaj najlepsze praktyki kodowania
i zaawansowane koncepcje
programowania

Michał Jaworski
Tarek Ziadé



Helion 

Packt 

Tytuł oryginału: Expert Python Programming: Master Python by learning the best coding practices and advanced programming concepts, 4th Edition

Tłumaczenie: Tomasz Walczak

ISBN: 978-83-283-8743-0

Copyright © Packt Publishing 2021. First published in the English language under the title 'Expert Python Programming - 4th Edition - (9781801071109)'.

Polish edition copyright © 2022 by Helion S.A.
All rights reserved.

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 the Publisher.

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)

Pliki z przykładami omawianymi w książce można znaleźć pod adresem:

<https://ftp.helion.pl/przyklady/prprp4.zip>

Drogi Czytelniku!

Jeżeli chcesz ocenić tę książkę, zajrzyj pod adres

<https://helion.pl/user/opinie/prprp4>

Możesz tam wpisać swoje uwagi, spostrzeżenia, recenzję.

Printed in Poland.

- Kup książkę
- Poleć książkę
- Oceń książkę

- Księgarnia internetowa
- Lubię to! » Nasza społeczność

Spis treści

O autorach	9
O recenzencie	10
Przedmowa	11
Rozdział 1. Aktualny stan Pythona	17
Gdzie jesteśmy i dokąd zmierzamy?	18
Co zrobić z kodem w Pythonie 2?	18
Jak być na bieżąco?	20
Dokumenty PEP	21
Aktywne społeczności	22
Inne źródła informacji	25
Podsumowanie	26
Rozdział 2. Nowoczesne środowiska programistyczne Pythona	27
Wymagania techniczne	28
Ekosystem pakietów Pythona	29
Instalowanie pakietów Pythona za pomocą narzędzia pip	29
Izolowanie środowiska uruchomieniowego	31
Izolacja na poziomie aplikacji a izolacja na poziomie systemu	33
Izolacja środowiska na poziomie aplikacji	34
Poetry jako system zarządzania zależnościami	37
Izolacja środowiska na poziomie systemu	41
Konteneryzacja a wirtualizacja	43
Zarządzanie środowiskami wirtualnymi z użyciem Dockera	44
Wirtualne środowiska programistyczne oparte na narzędziu Vagrant	61

Popularne narzędzia do zwiększania produktywności	63
Niestandardowe powłoki Pythona	63
Stosowanie powłoki IPython	65
Stosowanie powłok we własnych skryptach i programach	67
Interaktywne debugery	68
Inne narzędzia do zwiększania produktywności	69
Podsumowanie	71
Rozdział 3. Nowości w Pythonie	72
Wymagania techniczne	73
Niedawne dodatki do języka	73
Operatory scalania i aktualizacji słownika	74
Wyrażenia przypisania	78
Wskazówki dotyczące typów w typach generycznych	81
Parametry czysto pozycyjne	83
Moduł zoneinfo	85
Moduł graphlib	86
Nie tak nowe, ale wciąż błyszczące	90
Funkcja breakpoint()	90
Tryb roboczy	92
Funkcje __getattr__() i __dir__() na poziomie modułu	94
Formatowanie łańcuchów znaków za pomocą obiektów f-string	95
Podkreślenia w literałach liczbowych	96
Moduł secrets	96
Co może się pojawić w przyszłości?	98
Tworzenie sumy typów za pomocą operatora	98
Strukturalne dopasowywanie wzorców	99
Podsumowanie	103
Rozdział 4. Porównanie Pythona z innymi językami	104
Wymagania techniczne	105
Model klas i programowanie obiektowe	105
Dostęp do klas bazowych	106
Wielodziedziczenie i porządek MRO	108
Inicjalizowanie instancji klasy	113
Wzorce dostępu do atrybutów	116
Deskryptory	117
Właściwości	123
Dynamiczny polimorfizm	127
Przeciążanie operatorów	129
Przeciążanie funkcji i metod	135
Klasy danych	138
Programowanie funkcyjne	141
Funkcje lambda	142
Funkcje map(), filter() i reduce()	144
Obiekty i funkcje częściowe	146
Generatory	147
Wyrażenia generatora	148
Dekoratory	149

Wyliczenia	151
Podsumowanie	153
Rozdział 5. Interfejs, wzorce i modułowość	154
Wymagania techniczne	155
Interfejsy	155
Odrobina historii: zope.interface	157
Stosowanie adnotacji funkcji i abstrakcyjnych klas bazowych	164
Tworzenie interfejsów z wykorzystaniem adnotacji określających typ	169
Odwroćenie sterowania i wstrzykiwanie zależności	172
Odwroćenie sterowania w aplikacjach	173
Stosowanie platform do wstrzykiwania zależności	181
Podsumowanie	185
Rozdział 6. Współbieżność	186
Wymagania techniczne	186
Czym jest współbieżność?	187
Wielowątkowość	189
Czym jest wielowątkowość?	189
Obsługa wątków w Pythonie	192
Kiedy należy stosować wielowątkowość?	194
Przykładowa aplikacja wielowątkowa	197
Wieloprocusowość	212
Wbudowany moduł multiprocessing	214
Stosowanie puli procesów	217
Stosowanie modułu multiprocessing.dummy jako interfejsu do obsługi wielowątkowości	219
Programowanie asynchroniczne	220
Kooperytywna wielozadaniowość i asynchroniczne operacje wejścia – wyjścia	221
Słowa kluczowe async i await w Pythonie	222
Praktyczny przykład zastosowania programowania asynchronicznego	225
Dostosowywanie nieasynchronicznego kodu do asynchroniczności za pomocą obiektów future	228
Podsumowanie	231
Rozdział 7. Programowanie sterowane zdarzeniami	233
Wymagania techniczne	234
Czym dokładnie jest programowanie sterowane zdarzeniami?	234
Sterowanie zdarzeniami nie jest tożsame z asynchronicznością	235
Programowanie sterowane zdarzeniami w GUI	236
Komunikacja sterowana zdarzeniami	238
Różne style programowania sterowanego zdarzeniami	240
Styl oparty na wywołaniach zwrotnych	241
Styl oparty na obserwowaniu obiektów	242
Styl oparty na tematach	246
Architektury sterowane zdarzeniami	248
Kolejki zdarzeń i komunikatów	249
Podsumowanie	252

Rozdział 8. Elementy metaprogramowania	253
Wymagania techniczne	254
Czym jest metaprogramowanie?	254
Stosowanie dekoratorów do modyfikowania działania funkcji przed jej użyciem	255
Następny krok: dekoratory klas	256
Przechwytywanie procesu tworzenia instancji klasy	260
Metaklasy	263
Ogólna składnia	264
Stosowanie metaklas	267
Pułapki związane z metaklasami	270
Stosowanie metody <code>__init_subclass__()</code> jako alternatywy dla metaklas	271
Generowanie kodu	273
Funkcje <code>exec</code> , <code>eval</code> i <code>compile</code>	273
Drzewa składni abstrakcyjnej	274
Haczyki importu	276
Ważne przykłady generowania kodu w Pythonie	276
Podsumowanie	279
Rozdział 9. Łączenie Pythona z kodem w C i C++	280
Wymagania techniczne	281
C i C++ jako podstawa rozszerzalności w Pythonie	282
Kompilowanie i wczytywanie w Pythonie rozszerzeń napisanych w C	283
Kiedy należy używać rozszerzeń?	285
Zwiększanie wydajności kluczowych fragmentów kodu	285
Integrowanie istniejącego kodu napisanego w różnych językach	286
Integrowanie zewnętrznych bibliotek dynamicznych	287
Tworzenie wydajnych niestandardowych typów danych	287
Pisanie rozszerzeń	288
Rozszerzenia w czystym C	289
Pisanie rozszerzeń za pomocą Cythona	304
Wady korzystania z rozszerzeń	310
Dodatkowa złożoność	310
Trudniejsze debugowanie	311
Komunikacja z bibliotekami dynamicznymi bez używania rozszerzeń	312
Moduł <code>ctypes</code>	312
CFFI	318
Podsumowanie	320
Rozdział 10. Automatyzacja testów i kontroli jakości	321
Wymagania techniczne	322
Zasady programowania sterowanego testami	322
Pisanie testów z użyciem platformy <code>pytest</code>	325
Parametryzacja testów	331
Konfiguracje testów w platformie <code>pytest</code>	334
Stosowanie „fałszywych” obiektów	342
Atrapy i moduł <code>unittest.mock</code>	345

Automatyzacja kontroli jakości	349
Pokrycie kodu testami	349
Narzędzia do poprawiania stylu i lintery	353
Statyczna analiza typów	356
Testowanie mutacyjne	358
Przydatne narzędzia związane z testami	363
Generowanie realistycznych danych	363
Generowanie dat i czasu	365
Podsumowanie	366
 Rozdział 11. Tworzenie pakietów i udostępnianie kodu w Pythonie	 367
Wymagania techniczne	368
Tworzenie pakietów bibliotek i ich udostępnianie	368
Budowa pakietu Pythona	369
Rodzaje dystrybucji pakietów	377
Rejestrowanie i publikowanie pakietów	381
Wersjonowanie pakietów i zarządzanie zależnościami	383
Instalowanie własnych pakietów	387
Pakiety przestrzeni nazw	388
Skrypty i punkty wejścia w pakietach	390
Tworzenie pakietów aplikacji i usług do użytku w internecie	394
Manifest Twelve-Factor App	394
Korzystanie z Dockera	396
Zarządzanie zmiennymi środowiskowymi	398
Rola zmiennych środowiskowych w platformach do tworzenia aplikacji	402
Tworzenie samodzielnych aplikacji wykonywalnych	406
Kiedy samodzielne aplikacje wykonywalne są przydatne?	407
Popularne narzędzia	407
Bezpieczeństwo kodu Pythona w pakietach wykonywalnych	414
Podsumowanie	415
 Rozdział 12. Monitorowanie pracy i wydajności aplikacji	 417
Wymagania techniczne	418
Rejestrowanie błędów i logów	418
Podstawy rejestrowania logów w Pythonie	419
Zalecane praktyki z obszaru rejestrowania logów	430
Rozproszone rejestrowanie logów	433
Rejestrowanie błędów w celu ich późniejszej analizy	435
Instrumentacja kodu z wykorzystaniem niestandardowych wskaźników	438
Stosowanie aplikacji Prometheus	440
Śledzenie rozproszone aplikacji	448
Śledzenie rozproszone za pomocą Jaegera	451
Podsumowanie	456

Rozdział 13. Optymalizacja kodu	457
Wymagania techniczne	458
Częste przyczyny niskiej wydajności	458
Złożoność kodu	459
Nadmierne wykorzystanie zasobów i ich wyciekanie	462
Nadmierna liczba operacji wejścia – wyjścia i operacji blokujących	463
Profilowanie kodu	464
Profilowanie procesora	465
Profilowanie wykorzystania pamięci	472
Zmniejszanie złożoności przez wybór odpowiednich struktur danych	480
Przeszukiwanie listy	481
Stosowanie zbiorów	482
Stosowanie modułu collections	482
Architektoniczne kompromisy	487
Stosowanie heurystyk i algorytmów aproksymacyjnych	487
Stosowanie kolejek zadań i przetwarzania odroczonego	489
Stosowanie probabilistycznych struktur danych	491
Zapisywanie wyników w pamięci podręcznej	492
Podsumowanie	500
Skorowidz	501

Nowości w Pythonie

Jednym z najistotniejszych etapów w historii Pythona było prawdopodobnie udostępnienie wersji 3.0. Oto najważniejsze zmiany, jakie zostały w niej wprowadzone:

- rozwiązanie wielu problemów dotyczących obsługi tekstu, danych i kodowania Unicode,
- rezygnacja z klas w dawnym stylu,
- rozpoczęcie reorganizacji biblioteki standardowej,
- wprowadzenie adnotacji funkcji,
- wprowadzenie nowej składni obsługi wyjątków.

Z rozdziału 1., „Aktualny stan Pythona”, wiesz już, że Python 3 nie jest zgodny z Pythonem 2. Jest to główny powód, dla którego pełna akceptacja Pythona 3 przez społeczność zajęła tak dużo czasu. Była to bolesna, choć konieczna lekcja dla głównych programistów Pythona i społeczności skupionej wokół tego języka.

Na szczęście problemy związane z przyjmowaniem Pythona 3 nie zatrzymały procesu ewolucji języka. Od 3 grudnia 2008 roku (jest to data oficjalnego udostępnienia Pythona 3.0) konsekwentnie wprowadzane są nowe poważne aktualizacje Pythona. Każda nowa wersja dodaje nowe usprawnienia do języka, biblioteki standardowej i interpretera. Ponadto od wersji 3.9 obowiązuje roczny cykl udostępniania nowych edycji. To oznacza, że co rok otrzymamy nowe funkcje i usprawnienia.

Jeśli chcesz się dowiedzieć czegoś więcej o cyklu udostępniania wersji Pythona, zapoznaj się z dokumentem *PEP 602 — Annual Release Cycle for Python* (<https://www.python.org/dev/peps/pep-0602/>).

W tym rozdziale przyjrzymy się bliżej ewolucji Pythona w ostatnich latach. Omówimy kilka ważnych dodatków z ostatnich wersji. Spekulacyjnie wybiegniemy też w przyszłość i zaprezentujemy kilka funkcji, które zostały zaakceptowane w procesie przygotowywania dokumentów PEP i już niedługo staną się oficjalną częścią Pythona. Oto zagadnienia, którymi się zajmiemy:

- niedawne dodatki do języka,
- nie tak nowe, ale jeszcze błyszczące,
- co może nas czekać w przyszłości.

Jednak przed omówieniem nowych funkcji warto zacząć od analizy wymagań technicznych.

Wymagania techniczne

Oto omawiane w tym rozdziale pakiety Pythona, które możesz pobrać z repozytorium PyPI:

- mypy,
- pyright.

Informacje o instalowaniu pakietów znajdziesz w rozdziale 2., „Nowoczesne środowiska programistyczne Pythona”.

Pliki z kodem do tego rozdziału są dostępne na stronie <https://github.com/PacktPublishing/Expert-Python-Programming-Fourth-Edition/tree/main/Chapter%203> i w witrynie wydawnictwa Helion.

Niedawne dodatki do języka

W każdej wersji Pythona wprowadzane są liczne zmiany różnego rodzaju. W prawie wszystkich edycjach Pythona pojawiają się nowe elementy składni. Jednak większość zmian związana jest z biblioteką standardową Pythona, interpreterem CPython, API Pythona i API C w interpreterze CPython. Z powodu ograniczonego miejsca nie da się ich wszystkich omówić w tej książce. Dlatego skupiamy się tylko na nowych elementach składni i nowościach w bibliotece standardowej.

W dwóch najnowszych wersjach Pythona można wskazać cztery główne zmiany w składni:

- operatory scalania i aktualizacji słownika (dodane w Pythonie 3.9),
- wyrażenia przypisania (dodane w Pythonie 3.8),
- wskazówki dotyczące typów w typach generycznych (dodane w Pythonie 3.9),
- argumenty czysto pozycyjne (dodane w Pythonie 3.9).

Te cztery mechanizmy można zaliczyć do grupy zmian poprawiających jakość życia. Nie wprowadzają one żadnych nowych paradygmatów programowania. Nie zmieniają też radykalnie sposobu pisania kodu. Umożliwiają jedynie stosowanie lepszych wzorców programistycznych lub precyzyjniejsze definiowanie API.

W ostatnich latach główni programiści Pythona koncentrowali się przede wszystkim na usuwaniu martwych lub nadmiarowych modułów z biblioteki standardowej, a nie na dodawaniu czegokolwiek nowego. Jednak od czasu do czasu wprowadzane są jakieś dodatki do biblioteki standardowej. W ostatnich dwóch wersjach zyskałyśmy dwa zupełnie nowe moduły. Są to:

- moduł `zoneinfo` do korzystania z bazy danych ze strefami czasowymi organizacji **Internet Assigned Numbers Authority (IANA)** (dodany w Pythonie 3.9),
- moduł `graphlib` do operowania strukturami grafowymi (dodany w Pythonie 3.8).

Oba moduły mają dość małe API. Dalej omawiamy kilka obszarów, w których można jest zastosować. Najpierw jednak przyjrzyjmy się zmianom w składni wprowadzonym w Pythonie 3.8 i 3.9.

Operatory scalania i aktualizacji słownika

Python umożliwia stosowanie zestawu operatorów arytmetycznych do operowania wbudowanymi typami kontenerowymi, w tym listami, krotkami, zbiorami i słownikami.

Dla list i krotek można używać operatora dodawania `+`, aby złączyć dwie zmienne, pod warunkiem jednak, że są to zmienne tego samego typu. Dostępny jest też operator `+=`, który umożliwia modyfikowanie istniejących zmiennych. Poniższy fragment kodu zawiera przykłady złączania list i krotek w interaktywnej sesji:

```
>>> [1, 2, 3] + [4, 5, 6]
[1, 2, 3, 4, 5, 6]
>>> (1, 2, 3) + (4, 5, 6)
(1, 2, 3, 4, 5, 6)
>>> value = [1, 2, 3]
>>> value += [4, 5, 6]
>>> value
[1, 2, 3, 4, 5, 6]
>>> value = (1, 2, 3)
>>> value += (4, 5, 6)
>>> value
(1, 2, 3, 4, 5, 6)
```

Dla zbiorów dostępne są cztery operatory dwuargumentowe, które dają nowy zbiór:

- **Operator części wspólnej** — reprezentowany przez symbol `&` (bitowy operator I). Tworzy zbiór z elementami występującymi i w jednym, i w drugim zbiorze jednocześnie.
- **Operator sumy** — reprezentowany przez symbol `|` (bitowy operator LUB). Tworzy zbiór z wszystkimi elementami z obu zbiorów.
- **Operator różnicy** — reprezentowany przez symbol `-` (odejmowanie). Tworzy zbiór z elementami lewego zbioru, które nie występują w prawym zbiorze.
- **Operator XOR** — reprezentowany przez symbol `^` (bitowy operator XOR). Tworzy zbiór z elementami, które występują albo w jednym zbiorze, albo w drugim, ale nie w obu jednocześnie.

Poniższy fragment kodu zawiera przykłady zastosowania operatorów części wspólnej i sumy do zbiorów w interaktywnej sesji:

```
>>> {1, 2, 3} & {1, 4}
{1}
```

```
>>> {1, 2, 3} | {1, 4}
{1, 2, 3, 4}
>>> {1, 2, 3} - {1, 4}
{2, 3}
>>> {1, 2, 3} ^ {1, 4}
{2, 3, 4}
```

W Pythonie przez bardzo długi czas nie istniał specjalny operator dwuargumentowy, który umożliwiałby uzyskanie nowego słownika na podstawie dwóch istniejących słowników. Od wersji 3.9 można używać operatorów `|` (bitowy operator LUB) i `|=` (złożony bitowy operator LUB), które pozwalają scalać i aktualizować słowniki. Należy je stosować jako idiomatyczny sposób tworzenia sumy dwóch słowników. Powody dodania tych operatorów są opisane w dokumencie *PEP 584 — Add Union Operators To Dict*.

Idiom programistyczny jest typowym i zalecanym sposobem wykonywania określonego zadania w danym języku programowania. Pisanie idiomatycznego kodu jest ważną częścią kultury użytkowników Pythona. W dokumencie *Zen of Python* można przeczytać, że „powinien istnieć jeden — i najlepiej tylko jeden — oczywisty sposób na wykonanie danego zadania”.

Więcej idiomów omawiamy w rozdziale 4., „Porównanie Pythona z innymi językami”.

Aby scalać dwa słowniki w nowy, zastosuj następujące wyrażenie:

```
słownik_1 | słownik_2
```

Wynikowy słownik to zupełnie nowy obiekt z wszystkimi kluczami z obu słowników źródłowych. Jeśli w obu słownikach występują pokrywające się klucze, w wynikowym obiekcie zachowane są wartości z prawego obiektu.

Poniżej pokazujemy, jak zastosować tę składnię do dwóch literalów słownikowych. Lewy słownik jest modyfikowany wartościami ze słownika podanego po prawej stronie:

```
>>> {'a': 1} | {'a': 3, 'b': 2}
{'a': 3, 'b': 2}
```

Jeśli chcesz zmodyfikować zmienną reprezentującą słownik kluczami z innego słownika, możesz zastosować następujący operator złożony:

```
istniejący_słownik |= inny_słownik
```

Oto przykład zastosowania tego operatora do rzeczywistej zmiennej:

```
>>> mydict = {'a': 1}
>>> mydict |= {'a': 3, 'b': 2}
>>> mydict
{'a': 3, 'b': 2}
```

W starszych wersjach Pythona najprostszym sposobem na aktualizację istniejącego słownika zawartością innego słownika było użycie metody `update()`, tak jak w poniższym przykładzie:

```
istniejący_słownik.update(inny_słownik)
```

Ta metoda modyfikuje istniejący słownik w miejscu i nie zwraca wartości. To oznacza, że nie jest możliwe proste utworzenie scalonego słownika w ramach wyrażenia; metoda ta zawsze jest stosowana jako instrukcja.

Różnica między słowami „wyrażenie” i „instrukcja” jest wyjaśniona w podrozdziale „Wyrażenia przypisania”.

Inna możliwość: rozpakowywanie słownika

Mało znanym faktem jest to, że Python już przed wersją 3.9 udostępniał dość zwężony sposób scalania dwóch słowników. Służył do tego mechanizm **rozpakowywania słownika**. Obsługę rozpakowywania słowników w literalach typu `dict` dodano w Pythonie 3.5. Technika ta jest opisana w dokumencie *PEP 448 Additional Unpacking Generalizations*. Składnia pozwalająca rozpakować dwa słowniki (lub większą ich liczbę) do nowego obiektu wygląda tak:

```
{**słownik_1, **słownik_2}
```

Oto przykład z rzeczywistymi literalami:

```
>>> a = {'a': 1}; b = {'a': 3, 'b': 2}
>>> {**a, **b}
{'a': 3, 'b': 2}
```

Ten mechanizm, podobnie jak rozpakowywanie list (składnia `*wartość`), może wyglądać znajomo dla osób, które mają doświadczenie w pisaniu tak zwanych **funkcji wariadycznych**, przyjmujących niezdefiniowany zestaw argumentów i argumenty nazwane. Składnia ta jest przydatna przede wszystkim przy pisaniu dekoratorów.

Funkcje wariadyczne i dekoratory omawiamy szczegółowo w rozdziale 4., „Porównanie Pythona z innymi językami”.

Należy zapamiętać, że rozpakowywanie słownika, choć niezwykle popularne w definicjach funkcji, jest wyjątkowo rzadko stosowaną techniką scalania słowników. Może być ona niezrozumiała dla mniej doświadczonych programistów czytających Twój kod. Dlatego w kodzie używającym Pythona 3.9 i nowszych wersji języka należy preferować nowy operator scalania. W starszych wersjach Pythona czasem lepiej jest posłużyć się tymczasowym słownikiem i prostą metodą `update()`.

Inna możliwość: ChainMap z modułu collections

Jeszcze inny sposób tworzenia obiektu, który technicznie jest połączeniem dwóch słowników, polega na użyciu klasy `ChainMap` z modułu `collections`. Jest to klasa nakładkowa, która przyjmuje kilka obiektów odwzorowań (tu są nimi słowniki) i działa tak, jakby był to jeden obiekt odwzorowań.

Składnia scalania dwóch słowników za pomocą klasy `ChainMap` wygląda tak:

```
nowe_odwzorowanie = ChainMap(słownik_2, słownik_1)
```

Zauważ, że kolejność słowników jest tu odwrotna niż w operatorze `|`. To oznacza, że jeśli pobierzesz określony klucz z obiektu `nowe_odzworowanie`, opakowane obiekty będą przeszukiwane w kolejności od lewej do prawej. Przyjrzyj się poniższemu kodowi, który ilustruje operacje z użyciem klasy `ChainMap`:

```
>>> from collections import ChainMap
>>> user_account = {"iban": "GB71BARC20031885581746", "type":
"konto"}
>>> user_profile = {"display_name": "Jan Kowalski", "type": "profil"}
>>> user = ChainMap(user_account, user_profile)
>>> user["iban"]
'GB71BARC20031885581746'
>>> user["display_name"]
'Jan Kowalski'
>>> user["type"]
'konto'
```

W tym przykładzie widać, że wynikowy obiekt `user` w obiekcie typu `ChainMap` zawiera klucze z obu słowników — `user_account` i `user_profile`. Jeśli któreś z kluczy się powtarzają, obiekt typu `ChainMap` zwróci wartość z pierwszego od lewej odwzorowania, które zawiera dany klucz. Operator scalania słowników działa więc na odwrot.

`ChainMap` jest obiektem nakładkowym. To oznacza, że nie kopiuje zawartości przekazanych mu źródłowych odwzorowań, tylko zapisuje je jako referencje. Dlatego jeśli podstawowy obiekt się zmieni, `ChainMap` może zwrócić zmodyfikowane dane. Przyjrzyj się poniższej kontynuacji wcześniejszej sesji interaktywnej:

```
>>> user["display_name"]
'Jan Kowalski'
>>> user_profile["display_name"] = "Adam Mickiewicz"
>>> user["display_name"]
'Adam Mickiewicz'
```

Ponadto obiekty klasy `ChainMap` umożliwiają zapis i zachowują zmiany w powiązonym odwzorowaniu. Należy jednak pamiętać, że operacje zapisu, aktualizacji i usuwania dotyczą tylko lewego odwzorowania. Jeśli nie zachowasz należytej ostrożności, mogą wystąpić zaskakujące sytuacje, tak jak w poniższej kontynuacji poprzedniej sesji:

```
>>> user["display_name"] = "Jan Kowalski"
>>> user["age"] = 33
>>> user["type"] = "rozszerzenie"
>>> user_profile
{'display_name': 'Adam Mickiewicz', 'type': 'profil'}
>>> user_account
{'iban': 'GB71BARC20031885581746', 'type': 'rozszerzenie', 'display_name':
'Jan Kowalski', 'age': 33}
```

W tym przykładzie widać, że klucz `'display_name'` jest zapisywany w słowniku `user_account`, a początkowym słownikiem źródłowym tego klucza był `user_profile`. W wielu kontekstach takie zapisywanie danych w źródłowych słownikach w obiekcie klasy `ChainMap` jest niepożądane.

Dlatego często stosowany idiom związany z używaniem tych obiektów do scalania dwóch słowników obejmuje jawną konwersję na nowy słownik. Oto przykład dotyczący wcześniej zdefiniowanych słowników wejściowych:

```
>>> dict(ChainMap(user_account, user_profile))
{'display_name': 'Jan Kowalski', 'type': 'konto', 'iban':
'GB71BARC20031885581746'}
```

Jeśli chcesz tylko scalać dwa słowniki, wybierz nowy operator scalania zamiast klasy `ChainMap`. Nie oznacza to jednak, że takie obiekty stały się bezużyteczne. Jeżeli potrzebne jest wprowadzanie zmian w obie strony, należy posłużyć się klasą `ChainMap`. Ponadto działa ona dla odwzorowań dowolnego typu. Dlatego jeśli chcesz zapewnić ujednolicony dostęp do wielu obiektów działających tak jak słowniki, `ChainMap` umożliwia utworzenie jednostki działającej jak efekt ich scalenia.

Jeśli masz niestandardową klasę reprezentującą dzienniki, zawsze możesz rozszerzyć ją o specjalną metodę `__or__()`, aby zapewnić zgodność z operatorem `|`, zamiast używać klasy `ChainMap`. Przesłanianie metod specjalnych omawiamy w rozdziale 4., „Porównanie Pythona z innymi językami”. Jednak użycie klasy `ChainMap` jest zwykle łatwiejsze niż pisanie niestandardowej metody `__or__()` i umożliwia pracę z istniejącymi instancjami klas, których nie możesz zmodyfikować.

Zazwyczaj najważniejszym powodem stosowania klasy `ChainMap` zamiast rozpakowywania słowników lub operatora sumy jest zachowanie zgodności wstecz. W wersjach Pythona starszych niż 3.9 nie można stosować nowego operatora scalania słowników. Dlatego jeśli musisz pisać kod za pomocą starszych wersji Pythona, użyj klasy `ChainMap`. W nowych wersjach języka lepiej jest korzystać z operatora scalania.

Inną zmianą składniową, która ma duży wpływ na zgodność wstecz, są wyrażenia przypisania.

Wyrażenia przypisania

Jest to ciekawy mechanizm, ponieważ jego wprowadzenie wpływa na jedną z podstawowych cech składni Pythona — rozróżnienie na wyrażenia i instrukcje. Wyrażenia i instrukcje są podstawowymi cegiełkami prawie każdego języka programowania. Różnica między nimi jest prosta: wyrażenia mają wartość, a instrukcje nie.

Instrukcje możesz traktować jak kolejne operacje, które wykonuje program. Tak więc przypisania wartości, klauzule `if` lub pętle `for` i `while` są instrukcjami. Także definicje funkcji i klas są instrukcjami.

Wyrażenie to cokolwiek, co można umieścić w klauzuli `if`. Typowymi przykładami wyrażeń są literały, wartości zwracane przez operatory (w tym operatory złożone) oraz wyrażenia listowe, słownikowe i zbiorowe. Także wywołania funkcji i metod są wyrażeniami.

Niektóre elementy wielu języków programowania są prawie zawsze instrukcjami. Często są to:

- definicje funkcji i klas,
- pętle,
- klauzule `if...else`,
- przypisania wartości.

W Pythonie udało się wyeliminować opisane rozróżnienie dzięki udostępnieniu składni, która pozwala stosować odpowiedniki takich elementów języka mające postać wyrażeń:

- Wyrażenia `lambda` do tworzenia funkcji anonimowych są odpowiednikiem definicji funkcji:

```
lambda x: x**2
```

- Tworzenie instancji typu `type` jest odpowiednikiem definicji klasy:

```
type("MyClass", (), {})
```

- Różne wyrażenia składania kolekcji (ang. *comprehensions*) są odpowiednikiem pętli:

```
squares_of_2 = [x**2 for x in range(10)]
```

- Wyrażenia złożone są odpowiednikiem instrukcji `if ... else`:

```
"nieparzysta" if number % 2 else "parzysta"
```

Jednak przez wiele lat niedostępna była składnia, która pozwalałaby przypisywać wartość do zmiennej w formie wyrażenia. Była to bez wątpienia świadoma decyzja projektowa twórców Pythona. W językach takich jak C, gdzie przypisanie do zmiennej może mieć postać zarówno wyrażenia, jak i instrukcji, często się zdarza, że operator przypisania jest mylony z operatorem równości. Każdy, kto programował w C, może potwierdzić, że jest to irytujące źródło błędów. Przyjrzyj się przykładowemu kodowi w C:

```
int err = 0;
if (err = 1) {
    printf("Wystąpił błąd");
}
```

Porównaj go z poniższym fragmentem:

```
int err = 0;
if (err == 1) {
    printf("Wystąpił błąd");
}
```

Obie wersje są składniowo poprawne w C, ponieważ `err = 1` jest w tym języku wyrażeniem o wartości 1. Porównaj to z Pythonem, gdzie poniższy kod spowoduje błąd składniowy:

```
err = 0
if err = 1:
    printf("Wystąpił błąd")
```

Jednak w rzadkich sytuacjach wygodne może być używanie operatora przypisania do zmiennej, który zwraca wartość. Na szczęście w Pythonie 3.8 wprowadzono specjalny operator `:=`, który przypisuje wartość do zmiennej, ale działa jak wyrażenie, a nie jak instrukcja. Z powodu wyglądu został on szybko nazwany operatorem **morsa**.

Trzeba przy tym zauważyć, że zastosowania tego operatora są dość ograniczone. Pomaga on pisać bardziej zwężony kod, a taki kod często jest łatwiejszy do zrozumienia, ponieważ ma lepszy stosunek sygnału do szumu. Najczęściej stosuje się go wtedy, kiedy trzeba wykonać skomplikowane obliczenia wartości, a następnie natychmiast posłużyć się nią w dalszych instrukcjach.

Często podawanym przykładem jest używanie wyrażeń regularnych. Wyobraź sobie prostą aplikację, która wczytuje kod źródłowy napisany w Pythonie i sprawdza go za pomocą wyrażeń regularnych, szukając zaimportowanych modułów.

Bez wyrażeń przypisania kod może wyglądać tak:

```
import os
import re
import sys

import_re = re.compile(
    r"^\s*import\s+\.{0,2}((\w+\.)*(\w+))\s*$"
)
import_from_re = re.compile(
    r"^\s*from\s+\.{0,2}((\w+\.)*(\w+))\s+import\s+(\w+|\*)\s*$"
)

if __name__ == "__main__":
    if len(sys.argv) != 2:
        print(f"użytkowanie: {os.path.basename(__file__)} file-name")
        sys.exit(1)

    with open(sys.argv[1]) as file:
        for line in file:
            match = import_re.search(line)
            if match:
                print(match.groups()[0])

            match = import_from_re.search(line)
            if match:
                print(match.groups()[0])
```

Widać tu, że trzeba dwukrotnie powtórzyć kod, który sprawdza dopasowanie złożonych wyrażeń, a następnie pobiera pogrupowane tokeny. Za pomocą wyrażeń przypisania ten blok kodu można zmodyfikować w następujący sposób:

```
if match := import_re.match(line):
    print(match.groups()[0])

if match := import_from_re.match(line):
    print(match.groups()[0])
```

Widać więc niewielką poprawę czytelności, ale nie jest ona istotna. Tęgo rodzaju zmiany są przydatne w sytuacjach, gdy ten sam wzorec trzeba powtórzyć wielokrotnie. Ciągłe przypisywanie tymczasowych wyników do tej samej zmiennej może sprawić, że kod będzie niepotrzebnie rozbudowany.

Innym zastosowaniem wyrażeń przypisania jest używanie tych samych danych w wielu miejscach w większych wyrażeniach. Przyjrzyj się literalowi słownikowemu, który reprezentuje dane na temat fikcyjnego użytkownika:

```
first_name = "Jan"
last_name = "Kowalski"
height = 168
weight = 70

user = {
    "first_name": first_name,
    "last_name": last_name,
    "display_name": f"{first_name} {last_name}",
    "height": height,
    "weight": weight,
    "bmi": weight / (height / 100) ** 2,
}
```

Załóżmy, że ważne jest zachowanie spójności wszystkich elementów. Dlatego wyświetlane nazwisko zawsze powinno obejmować imię i nazwisko, a wskaźnik BMI należy obliczać na podstawie wagi i wzrostu. Aby zapobiec pomyłkom w trakcie edycji poszczególnych komponentów danych, trzeba zdefiniować je jako odrębne zmienne. Jednak po utworzeniu słownika nie są one potrzebne. Wyrażenia przypisania umożliwiają zapisanie tego słownika w bardziej zwęży sposób:

```
user = {
    "first_name": (first_name := "Jan"),
    "last_name": (last_name := "Kowalski"),
    "display_name": f"{first_name} {last_name}",
    "height": (height := 168),
    "weight": (weight := 70),
    "bmi": weight / (height / 100) ** 2,
}
```

Widać tu, że wyrażenia przypisania zostały umieszczone w nawiasach. Niestety składnia `:=` powoduje konflikt ze znakiem `:` używanym jako operator przypisania wartości do klucza w literalach słownikowych. Nawiasy pozwalają rozwiązać ten problem.

Wyrażenia przypisania są narzędziem do „dopieszczania” kodu i niczym więcej. Zawsze sprawdzaj, czy ich zastosowanie rzeczywiście poprawia czytelność, a nie zmniejsza ją.

Wskazówki dotyczące typów w typach generycznych

Adnotacje określające typ są w pełni opcjonalnym mechanizmem Pythona, ale zyskują coraz większą popularność. Dodaje się je do zmiennych, argumentów i typu wartości zwracanych przez funkcję. Tęgo rodzaju adnotacje pełnią funkcję dokumentacji, ale można też wykorzystać

je do sprawdzania poprawności kodu za pomocą zewnętrznych narzędzi. Wiele środowisk IDE potrafi przetwarzać adnotacje dotyczące typów i wizualnie wskazywać możliwe problemy z typami. Dostępne są też narzędzia do statycznego sprawdzania typów, na przykład **mypy** i **pyright**, które można wykorzystać do zbadania całego kodu bazowego i wskazania wszystkich błędów związanych z typami w jednostkach kodu opatrzonych adnotacjami.

Bardzo ciekawa jest historia projektu **mypy**. Początkowo powstał on w ramach rozprawy doktorskiej Jukka Lehtosalo, ale zaczął nabierać kształtu, gdy Jukka zaczął pracować nad nim razem z Guidem van Rossumem (twórcą Pythona) w firmie Dropbox. Więcej na ten temat dowiesz się z listu pożegnalnego skierowanego do Guido (<https://blog.dropbox.com/topics/company/thank-you--guido>).

W najprostszym podejściu wskazówki dotyczące typów można stosować do podawania oczekiwanych typów dla wbudowanych lub niestandardowych typów, argumentów wejściowych funkcji, wartości zwracanych przez funkcje i zmiennych lokalnych. Przyjrzyj się poniższej funkcji. Umożliwia ona wyszukiwanie kluczy bez uwzględniania wielkości liter w słowniku, którego kluczami są łańcuchy znaków:

```
from typing import Any

def get_ci(d: dict, key: str) -> Any:
    for k, v in d.items():
        if key.lower() == k.lower():
            return v
```

Ten przykład jest oczywiście naiwną implementacją wyszukiwania bez uwzględniania wielkości liter. Aby uzyskać wyższą wydajność, prawdopodobnie trzeba zastosować specjalną klasę. Do tego problemu wracamy w dalszej części książki.

Pierwsza instrukcja importuje typ **Any** z modułu **typing**. Ten typ definiuje, że zmienna lub argument może być dowolnego typu. W sygnaturze funkcji określone jest, że pierwszy argument, **d**, powinien być słownikiem, a drugi, **key**, łańcuchem znaków. Sygnatura kończy się specyfikacją zwracanej wartości, która może być dowolnego typu.

Jeśli używasz narzędzi do sprawdzania typów, takie adnotacje wystarczą do wykrycia wielu pomyłek. Jeżeli na przykład w jednostce wywołującej zostanie przedstawiona kolejność argumentów pozycyjnych, będzie można szybko wykryć błąd, ponieważ adnotacje argumentów **key** i **d** określają inne typy. Jednak te narzędzia nie zgłoszą problemu w sytuacji, gdy użytkownik przekaże słownik używający różnych typów dla kluczy.

Z tego powodu typy generyczne takie jak **tuple**, **list**, **dict**, **set**, **frozenset** i wiele innych można opatrzyć adnotacjami określającymi typ zawartości. Dla słowników takie adnotacje mają następującą postać:

```
dict[TypKlucza, TypWartości]
```

Sygnatura funkcji `get_ci()` z bardziej restrykcyjnymi adnotacjami określającymi typ może wyglądać tak:

```
def get_ci(d: dict[str, Any], key: str) -> Any: ...
```

W starszych wersjach Pythona nie można było tak łatwo dodawać adnotacji określających typ zawartości wbudowanych kolekcji. Moduł `typing` obejmuje specjalne typy, które można stosować w tym celu. Oto niektóre z tych typów:

- `typing.Dict` dla słowników,
- `typing.List` dla list,
- `typing.Tuple` dla krotek,
- `typing.Set` dla zbiorów,
- `typing.FrozenSet` dla zablokowanych zbiorów.

Te typy są przydatne, jeśli kod ma działać w różnych wersjach języka. Jeżeli jednak piszesz kod dostosowany tylko do Pythona 3.9 i nowszych wersji, należy używać wbudowanych typów generycznych. Importowanie ich z modułu `typing` to przestarzałe rozwiązanie, które w przyszłości zostanie wycofane z Pythona.

Adnotacje określające typ omawiamy dokładniej w rozdziale 4., „Porównanie Pythona z innymi językami”.

Parametry czysto pozycyjne

Python daje dużą swobodę, jeśli chodzi o przekazywanie argumentów do funkcji. Są dwa sposoby przekazywania argumentów do funkcji:

- w formie **argumentów pozycyjnych**,
- za pomocą **nazw argumentów**.

W wielu funkcjach to programista jednostki wywołującej decyduje o tym, jak przekazywać argumenty. Jest to korzystne, ponieważ użytkownik funkcji może uznać, że w danej sytuacji dana technika jest czytelniejsza lub wygodniejsza. Przyjrzyj się funkcji, która łączyła łańcuchy znaków z użyciem określonego separatora:

```
def concatenate(first: str, second: str, delim: str):
    return delim.join([first, second])
```

Funkcję tę można wywołać na wiele sposobów:

- za pomocą argumentów pozycyjnych — `concatenate("Jan", "Kowalski", " ")`,
- za pomocą nazw argumentów — `concatenate(first="Jan", second="Kowalski", delim=" ")`,
- łącząc argumenty pozycyjne i nazwy argumentów — `concatenate("Jan", "Kowalski", delim=" ")`.

Jeśli piszesz bibliotekę do wielokrotnego użytku, możliwe, że wiesz, w jaki sposób będzie ona używana. Czasem z doświadczenia wiadomo, że określony wzorzec wywołań zwiększy lub zmniejszy czytelność wynikowego kodu. Być może jeszcze nie masz pewności co do projektu lub chcesz się upewnić, że API biblioteki można będzie w bliskiej przyszłości zmienić bez wpływu na jej użytkowników. W obu sytuacjach warto tworzyć sygnatury funkcji tak, aby zachęcać do zalecanego sposobu użytkowania, a jednocześnie umożliwić wprowadzanie zmian w przyszłości.

Po opublikowaniu biblioteki sygnatura funkcji stanowi kontrakt jej użytkowania. Każda zmiana nazw argumentów lub ich kolejności może spowodować błędy w aplikacjach programistów, którzy korzystają z danej biblioteki.

Jeśli w przyszłości stwierdzisz, że nazwy argumentów `first` i `second` nie wyjaśniają poprawnie ich przeznaczenia, nie będzie można ich zmienić bez naruszania zgodności wstecz. Wynika to z tego, że jakiś programista mógł zastosować następujące wywołanie:

```
concatenate(first="Jan", second="Kowalski", delim=" ")
```

Jeżeli zechcesz przekształcić funkcję na postać, która przyjmuje dowolną liczbę łańcuchów znaków, też nie będzie to wykonalne, ponieważ jakiś programista mógł napisać poniższy kod:

```
concatenate("Jan", "Kowalski", " ")
```

Na szczęście w Pythonie 3.8 dodano możliwość definiowania określonych argumentów jako czysto pozycyjnych. Dzięki temu możesz wskazać, których argumentów nie można przekazywać za pomocą nazw. Pomaga to uniknąć późniejszych problemów z zachowaniem zgodności wstecz. Możesz też oznaczyć, że określone argumenty można wywoływać tylko za pomocą nazw. Staranne ustalenie, które argumenty należy przekazywać tylko jako pozycyjne, a które tylko za pomocą nazw, powoduje, że w przyszłości funkcję łatwiej będzie zmodyfikować. Funkcja `concatenate()` z argumentami czysto pozycyjnymi i podawanymi tylko za pomocą nazw może wyglądać tak:

```
def concatenate(first: str, second: str, /, *, delim: str):
    return delim.join([first, second])
```

Tę definicję należy czytać następująco:

- wszystkie argumenty przed znakiem `/` są argumentami czysto pozycyjnymi;
- wszystkie argumenty po znaku `*` są argumentami podawanymi tylko za pomocą nazw.

Ta definicja gwarantuje, że jedynym poprawnym wywołaniem funkcji `concatenate()` jest następująca postać:

```
concatenate("Jan", "Kowalski", delim=" ")
```

Jeśli spróbujesz wywołać ją w inny sposób, wystąpi błąd `TypeError`:

```
>>> concatenate("Jan", "Kowalski", " ")
Traceback (most recent call last):
  File "<stdin>", line 1, in <module>
TypeError: concatenate() takes 2 positional arguments but 3 were given
```

Załóżmy, że funkcja w tej postaci została opublikowana w bibliotece. Teraz chcesz, aby przyjmowała ona nieograniczoną liczbę argumentów pozycyjnych. Ponieważ istnieje tylko jeden sposób używania tej funkcji, w celu wprowadzenia tej zmiany możesz zastosować rozpakowywanie argumentów:

```
def concatenate(*items, delim: str):
    return delim.join(items)
```

Argument `*items` reprezentuje wszystkie argumenty pozycyjne z krotki `items`. Dzięki takim zmianom użytkownicy będą mogli korzystać z tej funkcji z różną liczbą argumentów pozycyjnych, tak jak w poniższych przykładach:

```
>>> concatenate("Jan", "Kowalski", delim=" ")
'Jan Kowalski'
>>> concatenate("Ronald", "Reuel", "Tolkien", delim=" ")
'Ronald Reuel Tolkien'
>>> concatenate("Jacek", delim=" ")
'Jacek'
>>> concatenate(delim=" ")
''
```

Argumenty czysto pozycyjne i podawane tylko za pomocą nazw to doskonałe narzędzia dla autorów bibliotek, ponieważ umożliwiają późniejsze wprowadzanie w projekcie zmian, które nie dotyczą użytkowników. Są też znakomitym narzędziem w trakcie pisania aplikacji, zwłaszcza jeśli współpracujesz z innymi programistami. Możesz wykorzystać argumenty czysto pozycyjne i podawane tylko za pomocą nazw, aby mieć pewność, że funkcja będzie wywoływana w oczekiwany sposób. Ułatwi to późniejszą refaktoryzację kodu.

Moduł zoneinfo

Obsługa czasu i stref czasowych jest jednym z największych wyzwań w programowaniu. Głównym tego powodem są liczne błędne założenia programistów w tym obszarze. Inną przyczyną są niekończące się zmiany definicji stref czasowych. Te zmiany mogą pojawiać się każdego roku i często są wprowadzane z przyczyn politycznych.

Python od wersji 3.9 zapewnia łatwiejszy niż kiedykolwiek wcześniej dostęp do informacji na temat bieżących i historycznych stref czasowych. Biblioteka standardowa Pythona udostępnia moduł `zoneinfo`, który jest interfejsem dla bazy danych ze strefami czasowymi dostępnej w systemie operacyjnym lub pobieranej jako pakiet standardowy `tzdata` z PyPI.

Pakiety z PyPI są uznawane za pakiety zewnętrzne, natomiast moduły biblioteki standardowej to pakiety standardowe. Pakiet `tzdata` jest wyjątkowy, ponieważ zarządzają nim główni programiści implementacji CPython. Zawartość bazy danych IANA została przeniesiona do odrębnych pakietów w PyPI, aby zagwarantować regularne aktualizacje niezależne od częstotliwości udostępniania implementacji CPython.

Aby użyć tego modułu, należy utworzyć obiekt typu `ZoneInfo` za pomocą następującego wywołania konstruktora:

```
ZoneInfo(klucz_strefy_czasowej)
```

W tym kodzie `klucz_strefy_czasowej` to nazwa pliku z bazy danych stref czasowych organizacji IANA. Te nazwy plików odpowiadają sposobowi, w jaki strefy czasowe są często prezentowane w różnych aplikacjach. Oto kilka przykładów:

- Europe/Warsaw,
- Asia/Tel_Aviv,
- America/Fort_Nelson,
- GMT-0.

Instancje klasy `ZoneInfo` mogą być używane jako parametr `tzinfo` w konstruktorze obiektów `datetime`:

```
from datetime import datetime
from zoneinfo import ZoneInfo

dt = datetime(2020, 11, 28, tzinfo=ZoneInfo("Europe/Warsaw"))
```

W ten sposób można tworzyć obiekty `datetime` uwzględniające strefę czasową. Takie obiekty są niezbędne do poprawnego obliczania różnic między określonymi strefami czasowymi, ponieważ potrafią uwzględniać takie kwestie jak czas letni i zimowy, a także historyczne zmiany wprowadzane w bazie danych stref czasowych organizacji IANA.

Pełną listę wszystkich stref czasowych dostępnych w systemie możesz pobrać za pomocą funkcji `zoneinfo.available_timezones()`.

Moduł graphlib

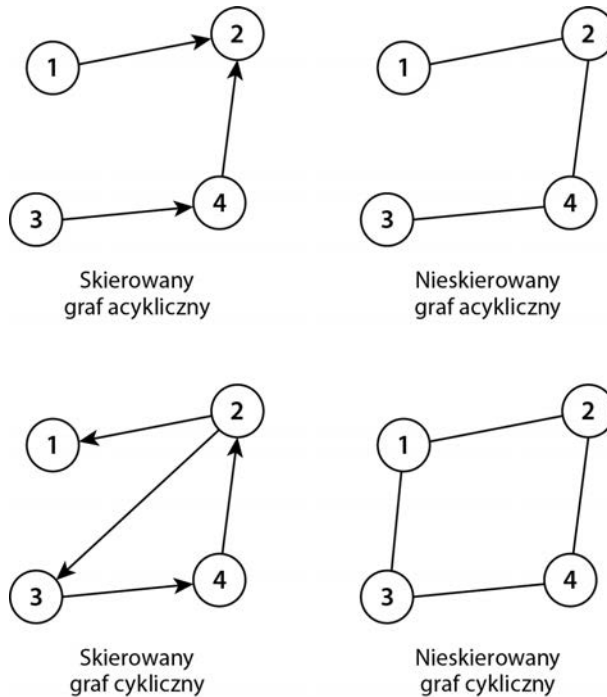
Innym ciekawym dodatkiem do biblioteki standardowej Pythona jest moduł `graphlib`, dodany w Pythonie 3.9. Ten moduł zapewnia narzędzia do pracy ze strukturami grafowymi.

Graf jest strukturą danych składającą się z węzłów połączonych krawędziami. Grafy są koncepcją z dziedziny matematyki nazywanej **teorią grafów**. W zależności od typu krawędzi można wyróżnić dwa podstawowe rodzaje grafów:

- **Graf nieskierowany** to graf, w którym wszystkie krawędzie są nieskierowane. Jeśli graf reprezentuje system miast połączonych drogami, krawędzie w grafie nieskierowanym są drogami dwukierunkowymi, którymi można się poruszać w obie strony. Dlatego jeżeli dwa węzły, *A* i *B*, są połączone krawędzią *E* w grafie nieskierowanym, można przejść z *A* do *B* i z *B* do *A*, używając tej samej krawędzi *E*.
- **Graf skierowany** cechuje się tym, że każda krawędź jest w nim skierowana. Jeśli graf reprezentuje system miast połączonych drogami, krawędzie w grafie skierowanym odpowiadają drogom jednokierunkowym, którymi można podróżować

tylko z jednego punktu źródłowego. Jeżeli dwa węzły, A i B , są połączone jedną krawędzią E , która wychodzi z węzła A , ta krawędź umożliwia przejście z A do B , ale już nie z B do A .

Ponadto grafy mogą być cykliczne lub acykliczne. **Graf cykliczny** ma przynajmniej jeden cykl — zamkniętą ścieżkę, która rozpoczyna się i kończy w tym samym węźle. **Graf acykliczny** nie obejmuje żadnych cykli. Na rysunku 3.1 pokazana jest przykładowa reprezentacja grafów skierowanych i nieskierowanych.



Rysunek 3.1. Wizualna reprezentacja różnych rodzajów grafów

Teoria grafów analizuje wiele problemów matematycznych, których modele można przedstawić za pomocą grafu. W programowaniu grafy są używane do rozwiązywania wielu problemów algorytmicznych. W informatyce grafy można wykorzystać do reprezentowania przepływu danych lub relacji między obiektami i mają wiele praktycznych zastosowań. Oto niektóre z nich:

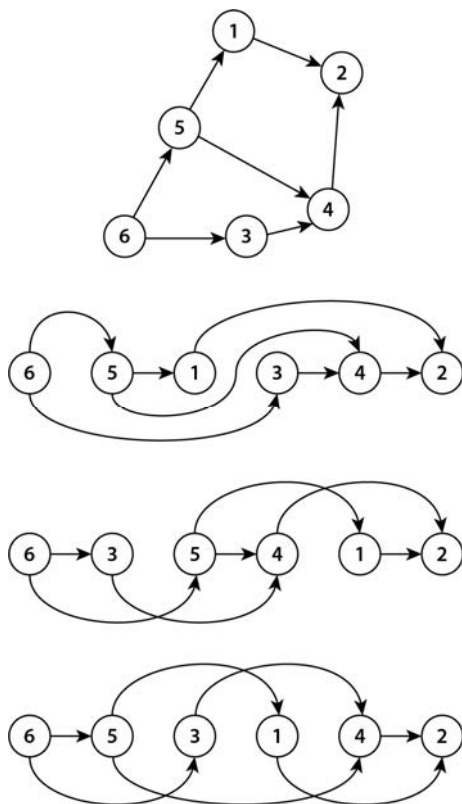
- modelowanie drzew zależności,
- reprezentowanie wiedzy w formacie czytelnym dla maszyn,
- wizualizowanie informacji,
- modelowanie systemów transportu.

Moduł `graphlib` ma ułatwiać programistom Pythona pracę z grafami. Jest to nowy moduł, dlatego obecnie obejmuje jedną klasę narzędziową, `TopologicalSorter`. Zgodnie z nazwą ta klasa przeprowadza sortowanie topologiczne skierowanych grafów acyklicznych.

Sortowanie topologiczne jest operacją porządkowania węzłów **skierowanych grafów acyklicznych** w określony sposób. Wynikiem sortowania topologicznego jest lista wszystkich węzłów, na której każdy węzeł pojawia się przed wszystkimi węzłami, do jakich można z niego przejść. Oznacza to, że:

- jako pierwszy zapisany jest węzeł, do którego nie można przejść z żadnego innego węzła;
- z każdego następnego węzła nie można przejść do poprzednich węzłów;
- z ostatniego węzła nie można przejść do żadnego innego.

W niektórych grafach występuje wiele uporządkowań zgodnych z warunkami sortowania topologicznego. Rysunek 3.2 przedstawia przykładowy skierowany graf acykliczny z trzema możliwymi porządkami topologicznymi.



Rysunek 3.2. Różne sposoby topologicznego sortowania tego samego grafu

Aby lepiej zrozumieć zastosowania sortowania topologicznego, rozważ następujący problem. Istnieje złożona operacja, która składa się z wielu zależnych od siebie zadań. Tą operacją może być na przykład przenoszenie wielu tabel bazy danych między dwoma różnymi systemami bazodanowymi. Jest to często występujący problem i istnieje już wiele narzędzi umożliwiających migrację danych do różnych systemów zarządzania bazami. Jednak w celach edukacyjnych założmy, że takie rozwiązanie nie istnieje i trzeba zbudować je od podstaw.

W relacyjnym systemie bazodanowym wiersze w tabelach często zawierają referencje do wierszy z innych tabel. Integralność takich referencji jest chroniona za pomocą **ograniczeń klucza obcego**. Jeśli chcesz zagwarantować, że w danym momencie docelowa baza danych zachowuje integralność referencyjną, musisz przenosić wszystkie tabele w określonej kolejności. Załóżmy, że używane są następujące tabele:

- tabela `customers` zawierająca informacje o klientach;
- tabela `accounts` przechowująca informacje o kontach użytkowników, w tym o stanie tych kont; jeden użytkownik może mieć wiele kont (na przykład osobiste i firmowe), ale jedno konto nie może być przypisane do kilku użytkowników;
- tabela `products` z informacjami na temat produktów dostępnych w sprzedaży w systemie;
- tabela `orders` obejmująca pojedyncze zamówienia wielu produktów złożone z jednego konta przez jednego użytkownika;
- tabela `order_products` z danymi o liczbie poszczególnych produktów w jednym zamówieniu.

Python nie ma specjalnego typu danych reprezentującego grafy. Dostępny jest jednak typ słownikowy, który świetnie nadaje się do zapisywania relacji między kluczami i wartościami. Zdefiniuj referencje między wymienionymi fikcyjnymi tabelami:

```
table_references = {
    "customers": set(),
    "accounts": {"customers"},
    "products": set(),
    "orders": {"accounts", "customers"},
    "order_products": {"orders", "products"},
}
```

Jeśli graf referencji nie ma cykli, można go posortować topologicznie. Wynikiem tego sortowania jest kolejność przenoszenia tabel. Konstruktor klasy `graphlib.TopologicalSorter` przyjmuje jako dane wyjściowe jeden słownik, w którym kluczami są węzły źródłowe, a wartościami — zbiory węzłów docelowych. To oznacza, że możesz przekazać zmienną `table_references` bezpośrednio do konstruktora klasy `TopologicalSorter()`. W celu przeprowadzenia sortowania topologicznego możesz użyć wywołania `static_order()`, tak jak w poniższym zapisie sesji interaktywnej:

```
>>> from graphlib import TopologicalSorter
>>> table_references = {
...     "customers": set(),
...     "accounts": {"customers"},
...     "products": set(),
...     "orders": {"accounts", "customers"},
...     "order_products": {"orders", "products"},
... }
>>> sorter = TopologicalSorter(table_references)
>>> list(sorter.static_order())
['customers', 'products', 'accounts', 'orders', 'order_products']
```

Sortowanie topologiczne można wykonywać tylko na skierowanych grafach acyklicznych. Klasa `TopologicalSorter` w trakcie inicjalizacji nie sprawdza występowania cykli, choć wykrywa cykle na etapie sortowania. Po znalezieniu cyklu metoda `static_order()` zgłasza wyjątek `graphlib.CycleError`.

Ten przykład jest oczywiście prosty i łatwo można wykonać zadanie ręcznie. Jednak prawdziwe bazy danych często zawierają dziesiątki, a nawet setki tabel. Ręczne przygotowanie planu dla tak dużych baz to bardzo żmudne zadanie, w którym łatwo o pomyłki.

Mechanizmy omówione do tego miejsca są stosunkowo nowe, dlatego musi minąć trochę czasu, zanim staną się powszechnie stosowane w Pythonie. Wynika to z tego, że nie są zgodne wstecz, a wiele osób zajmujących się konserwacją bibliotek nadal zapewnia obsługę starszych wersji Pythona.

W następnym podrozdziale omawiamy ważne elementy Pythona wprowadzone w wersjach 3.6 i 3.7, a więc dostępne w większej liczbie wersji. Jednak nie wszystkie te nowe elementy są popularne, dlatego mamy nadzieję, że nauczysz się czegoś nowego.

Nie tak nowe, ale wciąż błyszczące

Każda wersja Pythona zawiera jakieś nowości. Niektóre zmiany są prawdziwym przełomem — znacznie ułatwiają programowanie i są prawie natychmiast przyjmowane przez społeczność. Jednak zalety innych modyfikacji nie zawsze są od razu oczywiste, dlatego ich spopularyzowanie zajmuje więcej czasu.

Dotyczy to na przykład adnotacji funkcji, które są dostępne w Pythonie już od wersji 3.0. Musiały minąć lata, aby powstał ekosystem narzędzi korzystających z takich adnotacji. Obecnie adnotacje są stosowane w prawie wszystkich nowoczesnych aplikacjach Pythona.

Główni programiści Pythona bardzo konserwatywnie podchodzą do dodawania nowych modułów do biblioteki standardowej, dlatego rzadko pojawiają się w niej nowości. Możliwe, że zapomnisz o modułach `graphlib` i `zoneinfo`, jeśli nie będziesz mieć okazji do pracy nad problemami, które wymagają operowania grafami lub starannego uwzględniania stref czasowych. Możliwe, że już nie pamiętasz o innych przydatnych dodatkach do Pythona, które zostały wprowadzone w kilku ostatnich latach. Dlatego przedstawiamy krótki przegląd ważnych zmian wprowadzonych do wersji 3.7. Są to albo małe, ale interesujące dodatki, które łatwo jest przeoczyć, albo rozwiązania, do których trzeba się przyzwyczaić.

Funkcja `breakpoint()`

Debugery omówiliśmy w rozdziale 2., „Nowoczesne środowiska programistyczne Pythona”. Wspomnieliśmy tam o funkcji `breakpoint()` jako o idiomatycznym sposobie pracy z debugerem Pythona.

Funkcja ta została wprowadzona w Pythonie 3.7, dlatego jest dostępna już od jakiegoś czasu. Jest to jednak jedna ze zmian, do których przyzwyczajenie się wymaga trochę wysiłku. Przez wiele lat uczono nas, że najprostszym sposobem wywoływania debugera w kodzie Pythona jest następujący fragment:

```
import pdb; pdb.set_trace()
```

Ten kod nie jest ani zgrabny, ani prosty, ale jeśli ktoś — podobnie jak wielu innych programistów — stosował go codziennie przez wiele lat, posługuje się nim automatycznie. Wystąpił problem? Wystarczy przejść do kodu, wpisać kilka znaków, aby wywołać debugger `pdb`, a potem ponownie uruchomić program. Znajdziesz się wtedy w powłoce interpretera w miejscu wystąpienia błędu. Gotowe? Wystarczy wrócić do kodu, usunąć fragment `import pdb; pdb.set_trace()`, a następnie zacząć pracę nad poprawkami.

Dlaczego więc masz się tym przejmować? Czy to nie kwestia indywidualnych preferencji? Czy punkty przerwania są czymś, co kiedykolwiek trafia do kodu produkcyjnego?

Prawda jest taka, że debugowanie często jest samotniczym i bardzo osobistym zadaniem. Często przez wiele godzin zmagamy się z błędami, szukamy wskazówek i w kółko czytamy ten sam kod, desperacko starając się zlokalizować drobny błąd, przez który aplikacja nie działa. Gdy koncentrujesz się na znalezieniu przyczyny problemu, zdecydowanie warto używać najwygodniejszego dla Ciebie narzędzia. Niektórzy programiści preferują debugery zintegrowane ze środowiskiem IDE. Inni w ogóle nie korzystają z debuggerów i wolą rozbudowane wywołania `print()` w różnych miejscach kodu. Zawsze wybieraj to, co jest dla Ciebie najwygodniejsze.

Jeśli jednak masz w zwyczaju używać staromodnego debugera działającego w powłoce, funkcja `breakpoint()` może ułatwić Ci pracę. Główną zaletą tej funkcji jest to, że nie jest powiązana z jednym debugerem. Domyślnie uruchamia ona sesję debugera `pdb`, ale można to zmienić, używając zmiennej środowiskowej `PYTHONBREAKPOINT`. Jeśli wolisz korzystać z innego debugera (na przykład `ipdb`, który opisaliśmy w rozdziale 2., „Nowoczesne środowiska programistyczne Pythona”), możesz przypisać do tej zmiennej wartość informującą Pythona, jaką funkcję ma wywoływać.

Standardową praktyką jest konfigurowanie preferowanego debugera w skrypcie profilu powłoki, dzięki czemu nie trzeba modyfikować zmiennej w każdej sesji powłoki. Na przykład jeśli używasz Basha i zawsze chcesz korzystać z `ipdb` zamiast z `pdb`, możesz wstawić poniższą instrukcję w pliku `.bash_profile`:

```
PYTHONBREAKPOINT=ipdb.set_trace()
```

To podejście dobrze sprawdza się także w zespołach. Na przykład jeśli ktoś chce, abyś pomógł mu w debugowaniu, możesz poprosić, aby wstawił instrukcje `breakpoint()` w miejscach, które mogą powodować problem. Dzięki temu gdy uruchomisz kod na własnym komputerze, użyty zostanie wybrany przez Ciebie debugger.

Jeśli nie wiesz, gdzie umieścić punkt przerwania, ale aplikacja kończy pracę w wyniku wystąpienia nieobsługiwanego wyjątku, możesz wykorzystać mechanizm debugowania poawaryjnego z pdb. Poniższe polecenie uruchamia skrypt Pythona w sesji diagnostycznej, w której działanie kodu jest wstrzymywane w momencie zgłoszenia wyjątku:

```
python3 -m pdb -c continue script.py
```

Tryb roboczy

Od wersji 3.7 interpreter Pythona można uruchomić w specjalnym trybie roboczym, w którym w czasie wykonywania kodu stosowane są dodatkowe analizy. Pomaga to w diagnozowaniu ewentualnych problemów, jakie mogą wystąpić w trakcie działania programu. W poprawnie działającym kodzie takie analizy generują niepotrzebne koszty, dlatego domyślnie są nieaktywne.

Tryb roboczy można aktywować na dwa sposoby:

- Za pomocą opcji `-X dev` przy wywołaniu interpretera Pythona w wierszu poleceń, na przykład:

```
python -X dev my_application.py
```

- Za pomocą zmiennej środowiskowej `PYTHONDEVMODE`, na przykład:

```
PYTHONDEVMODE=1 my_application
```

Oto najważniejsze skutki zastosowania tego trybu:

- haczyki związane z przydziałem pamięci, co pozwala wykrywać przepełnienie i niedopełnienie bufora, naruszenia API alokatorów pamięci i przypadki niebezpiecznego użycia blokad **Global Interpreter Lock (GIL)**;
- ostrzeżenia związane z możliwymi błędami przy imporcie modułów;
- ostrzeżenia dotyczące nieprawidłowego zarządzania zasobami, na przykład niezamykania otwartych plików;
- ostrzeżenia związane z elementami biblioteki standardowej, które są uznawane za przestarzałe i w przyszłości zostaną usunięte;
- włączona obsługa błędów, co powoduje wyświetlenie śladu stosu aplikacji po otrzymaniu przez nią sygnału systemowego `SIGSEGV`, `SIGFPE`, `SIGABRT`, `SIGBUS` lub `SIGILL`.

Ostrzeżenia generowane w trybie roboczym oznaczają, że coś nie działa w oczekiwany sposób. Może to pomóc w wykryciu problemów, które nie przyjmują formy błędów w trakcie standardowej pracy kodu, ale mogą prowadzić do poważnych usterek, gdy aplikacja będzie działać przez długi czas.

Nieprawidłowe zamykanie otwartych plików może w pewnym momencie doprowadzić do wyczerpania się zasobów w środowisku, w którym aplikacja działa. Zasobami są na przykład deskryptory plików, podobnie jak pamięć RAM i przestrzeń dyskowa. Każdy system operacyjny ma ograniczoną liczbę plików, które mogą być jednocześnie otwarte. Jeśli aplikacja otwiera nowe pliki, ale ich nie zamyka, w pewnym momencie nie będzie mogła utworzyć kolejnych.

Tryb roboczy pozwala wcześniej wykryć takie problemy. To dlatego zaleca się korzystanie z niego na etapie testowania aplikacji. Z powodu dodatkowych kosztów analiz aktywnych w trybie roboczym nie zaleca się stosowania go w środowiskach produkcyjnych.

Czasem tryb roboczy można wykorzystać także do diagnozowania istniejących problemów. Przykładem wysoce problematycznej sytuacji jest błąd segmentacji w aplikacji.

Gdy taki błąd występuje w Pythonie, zwykle nie są wyświetlane żadne szczegółowe informacje na jego temat. Zobaczysz tylko krótki komunikat wyświetlony w powłoce:

```
Segmentation fault: 11
```

W momencie wystąpienia błędu segmentacji proces Pythona otrzymuje sygnał systemowy SIGSEGV i natychmiast kończy pracę. W niektórych systemach operacyjnych wyświetlany jest zrzut pamięci, czyli zapis stanu pamięci procesu zarejestrowany w momencie awarii. Można go wykorzystać przy debugowaniu aplikacji. Niestety w interpreterze CPython wyświetlany jest zapis pamięci procesu interpretera, dlatego debugowanie odbywa się na poziomie kodu w języku C.

Tryb roboczy dodaje kod do obsługi błędów, który wyświetla ślad stosu Pythona po otrzymaniu sygnału błędu. Dzięki temu otrzymujesz dodatkowe informacje na temat tego, które fragmenty kodu mogą powodować problem. Oto przykładowy kod, o którym wiadomo, że powoduje błąd segmentacji w Pythonie 3.9:

```
import sys

sys.setrecursionlimit(1 << 30)

def crasher():
    return crasher()

crasher()
```

Jeśli wykonasz ten kod w interpreterze Pythona z włączoną opcją `-X dev`, otrzymasz dane wyjściowe podobne do poniższych:

```
Fatal Python error: Segmentation fault

Current thread 0x000000010b04edc0 (most recent call first):
  File "/Users/user/dev/crashers/crasher.py", line 6 in crasher
  File "/Users/user/dev/crashers/crasher.py", line 6 in crasher
  File "/Users/user/dev/crashers/crasher.py", line 6 in crasher
  File "/Users/user/dev/crashers/crasher.py", line 6 in crasher
  File "/Users/user/dev/crashers/crasher.py", line 6 in crasher
  ...
```

Ten mechanizm obsługi błędów można też aktywować poza trybem roboczym. W tym celu użyj opcji `-X faulthandler` w wierszu poleceń lub przypisz do zmiennej środowiskowej `PYTHONFAULTHANDLER` wartość 1.

Niełatwo jest wywołać błędy segmentacji w Pythonie. Często występują one dla rozszerzeń Pythona napisanych w języku C lub C++ oraz dla funkcji wywoływanych z poziomu bibliotek współdzielonych (plików DLL, *.dylib* lub *.so*). Jednak istnieje kilka znanych i dobrze udokumentowanych scenariuszy, w których ten błąd może wystąpić w kodzie w czystym Pythonie. W repozytorium implementacji CPython znajduje się kolekcja tego typu znanych „prowokatorów awarii”. Znajdziesz ją na stronie <https://github.com/python/cpython/tree/master/Lib/test/crashers>.

Funkcje `__getattr__()` i `__dir__()` na poziomie modułu

W każdej klasie Pythona można zdefiniować niestandardowe funkcje `__getattr__()` i `__dir__()`, aby zmodyfikować dynamiczny dostęp do atrybutów obiektów. Funkcja `__getattr__()` jest wywoływana, gdy nie można znaleźć podanej nazwy atrybutu — pozwala to obsłużyć wyszukiwanie atrybutu i ewentualnie wygenerować jego wartość „w locie”. Metoda `__dir__()` jest uruchamiana po przekazaniu obiektu do funkcji `dir()`. Powinna ona zwracać listę nazw atrybutów obiektu.

Od Pythona 3.7 funkcje `__getattr__()` i `__dir__()` można definiować na poziomie modułu. Działają one podobnie jak metody obiektów. Funkcja `__getattr__()` z poziomu modułu jest wywoływana po nieudanym wyszukiwaniu składowej modułu. Funkcja `__dir__()` jest uruchamiana po przekazaniu obiektu modułu do funkcji `dir()`.

Omawiane możliwości mogą być przydatne dla osób odpowiedzialnych za konserwację biblioteki, gdy funkcje lub klasy z modułu zostają uznane za przestarzałe. Założmy, że funkcja `get_ci()` z podrzdziału „Wskazówki dotyczące typów w typach generycznych” jest udostępniana w otwartej bibliotece *dict_helpers.py*. Jeśli chcesz zmienić nazwę funkcji na `lookup_ci()`, ale umożliwić importowanie jej po dawną nazwą, możesz zastosować następujący wzorzec wycofywania elementów kodu:

```
from typing import Any
from warnings import warn

def ci_lookup(d: dict[str, Any], key: str) -> Any:
    ...

def __getattr__(name: str):
    if name == "get_ci":
        warn(f"{name} jest przestarzała", DeprecationWarning)
        return ci_lookup

    raise AttributeError(f"Moduł {__name__} nie zawiera atrybutu {name}")
```

Ten wzorzec generuje ostrzeżenie `DeprecationWarning` niezależnie od tego, czy funkcja `get_ci()` została zaimportowana bezpośrednio z modułu (na przykład za pomocą instrukcji `from dict_helpers import get_ci`), czy użyta jako atrybut `dict_helpers.get_ci`.

Ostrzeżenia o przestarzałych elementach nie są domyślnie wyświetlane. Możesz włączyć ich wyświetlanie w trybie roboczym.

Formatowanie łańcuchów znaków za pomocą obiektów f-string

Obiekty f-string, nazywane też **formatowanymi literalami tekstowymi**, są jedną z najbardziej lubianych funkcji Pythona, które pojawiły się w wersji 3.6. Zostały opisane w dokumencie PEP 498 i zapewniają nowy sposób formatowania łańcuchów znaków. W wersjach starszych niż 3.6 dostępne były już dwie różne metody formatowania łańcuchów znaków. Tak więc obecnie łańcuch znaków można sformatować za pomocą trzech technik. Oto one:

- Formatowanie z użyciem symbolu %. Jest to najstarsza technika. Używany jest w niej wzorzec podstawiania przypominający składnię funkcji `printf()` z biblioteki standardowej języka C:

```
>>> import math
>>> "Przybliżona wartość liczby  $\pi$ : %f" % math.pi
'Przybliżona wartość liczby  $\pi$ : 3.141593'
```

- Formatowanie za pomocą metody `str.format()`. Ta technika jest wygodniejsza i mniej narażona na błędy niż formatowanie z użyciem symbolu %, ale wymaga więcej kodu. Umożliwia stosowanie nazwanych zastępowanych pól, a także wielokrotne wykorzystanie tej samej wartości:

```
>>> import math
>>> "Przybliżona wartość liczby  $\pi$ : {:f}".format(pi=math.pi)
'Przybliżona wartość liczby  $\pi$ : 3.141593'
```

- Formatowanie za pomocą obiektów **f-string**. Jest to najbardziej zwięzła, wszechstronna i wygodna technika formatowania łańcuchów znaków. Automatycznie zastępuje elementy literalów na podstawie wartości zmiennych i wyrażeń z lokalnej przestrzeni nazw:

```
>>> import math
>>> f"Przybliżona wartość liczby  $\pi$ : {math.pi:f}"
'Przybliżona wartość liczby  $\pi$ : 3.141593'
```

Obiekty f-string są oznaczone za pomocą przedrostka `f`, a ich składnia przypomina składnię metody `str.format()`, ponieważ stosuje się w niej podobne techniki do podawania zastępowanych pól w formatowanym tekście. W metodzie `str.format()` wstawiany tekst jest określany za pomocą argumentów pozycyjnych i nazw argumentów. Obiekty f-string są wyjątkowe dlatego, że jako zastępowane pola można podawać dowolne wyrażenia Pythona przetwarzane w czasie wykonywania programu. W takich miejscach można używać każdej zmiennej, która jest dostępna w przestrzeni nazw obejmującej formatowany literal.

Możliwość stosowania wyrażeń jako zastępowanych pól sprawia, że kod definiujący formatowanie jest prostszy i krótszy. Dostępne są też te same specyfikatory formatu zastępowanych pól (określające dopełnienie, wyrównanie, wyświetlanie znaków itd.) co w metodzie `str.format()`. Oto używana składnia:

```
f"{wyrażenie_dla_zastępowanego_pola:specyfikator_formatowania}"
```

Poniżej przedstawiamy prosty przykład wykonania w interaktywnej sesji kodu, który wyświetla pierwszych 10 potęg liczby 10 za pomocą obiektów f-string i wyrównuje wyniki z wykorzystaniem dopełnienia:

```
>>> for x in range(10):
...     print(f"10^{x} == {10**x:10d}")
...
10^0 ==      1
10^1 ==     10
10^2 ==    100
10^3 ==   1000
10^4 ==  10000
10^5 == 100000
10^6 == 1000000
10^7 == 10000000
10^8 == 100000000
10^9 == 1000000000
```

Kompletna specyfikacja formatowania łańcuchów znaków Pythona tworzy odrębny minijęzyk. Najlepszym źródłem informacji o nim jest oficjalna dokumentacja — <https://docs.python.org/3/library/string.html>. Inną przydatną stroną poświęconą temu zagadnieniu jest <https://pyformat.info/> — znajdziesz tam omówienie najważniejszych elementów tej specyfikacji razem z praktycznymi przykładami.

Podkreślenia w literałach liczbowych

Podkreślenia w literałach liczbowych są bardzo łatwe do wprowadzenia, ale nadal nie stały się tak popularne, jak mogłyby być. Od Pythona 3.6 można używać znaku podkreślenia (`_`) do oddzielania cyfr w literałach liczbowych. Poprawia to czytelność dużych liczb. Przyjrzyj się następującemu przypisaniu wartości:

```
account_balance = 100000000
```

Przy tak dużej liczbie zer trudno jest szybko stwierdzić, czy chodzi tu o setki milionów, czy o miliardy. Możesz więc użyć znaku podkreślenia do rozdzielania tysięcy, milionów, miliardów itd.:

```
account_balance = 100_000_000
```

Teraz łatwiej jest od razu stwierdzić, że zmienna `account_balance` ma wartość sto milionów; nie wymaga to starannego liczenia zer.

Moduł secrets

Jednym błędów w obszarze zabezpieczeń często popełnianych przez wielu programistów jest zakładanie losowości modułu `random`. Losowość liczb generowanych przez ten moduł wystarcza dla celów statystycznych. W tym module używany jest generator liczb pseudolosowych Mersenne Twister. Zapewnia on znany równomierny rozkład i ma wystarczającą długość okresu, dlatego można go stosować w symulacjach, modelowaniu lub całkowaniu.

Jednak algorytm Mersenne Twister jest w pełni deterministyczny, podobnie jak moduł `random`. To oznacza, że gdy znane są początkowe warunki (**ziarno**), można wygenerować te same liczby pseudolosowe. Ponadto kiedy znana jest wystarczająca liczba kolejnych wartości z generatora pseudolosowego, zwykle można ustalić ziarno i przewidzieć kolejne wyniki. Dotyczy to także algorytmu Mersenne Twister.

Jeśli chcesz się dowiedzieć, jak przewidzieć liczby generowane przez algorytm Mersenne Twister, zapoznaj się z następującym projektem z serwisu GitHub: <https://github.com/kmyk/mersenne-twister-predictor>.

Ta cecha generatorów liczb pseudolosowych oznacza, że nigdy nie należy ich stosować do generowania losowych wartości na potrzeby zabezpieczeń. Na przykład jeśli chcesz wygenerować losową sekwencję, którą można wykorzystać jako hasło użytkownika lub token, znajdź inne źródło losowości.

Moduł `secrets` służy właśnie do tego. Wykorzystuje on najlepsze źródło losowości w poszczególnych systemach operacyjnych. W Uniksie i systemach uniksowych jest nim urządzenie `/dev/urandom`, a w systemie Windows — generator `CryptGenRandom`.

Oto trzy najważniejsze funkcje z tego modułu:

- `secrets.token_bytes(nbytes=None)` — zwraca `nbytes` losowych bajtów. Ta funkcja jest używana wewnętrznie przez funkcje `secrets.token_hex()` i `secrets.token_urlsafe()`. Jeśli programista nie poda wartości `nbytes`, zwracana jest domyślna liczba bajtów, opisana w dokumentacji jako „rozsądna”.
- `secrets.token_hex(nbytes=None)` — zwraca `nbytes` losowych bajtów w postaci szesnastkowego łańcucha znaków (nie jest to obiekt taki, jaki zwraca funkcja `bytes()`). Ponieważ jeden bajt jest kodowany za pomocą dwóch cyfr szesnastkowych, wynikowy łańcuch znaków zawiera `nbytes * 2` znaków. Jeśli wartość `nbytes` nie jest podana, zwracana jest ta sama domyślna liczba bajtów co w funkcji `secrets.token_bytes(nbytes=None)`.
- `secrets.token_urlsafe(nbytes=None)` — zwraca `nbytes` losowych bajtów w postaci dostosowanego do adresów URL łańcucha znaków w formacie `base64`. Ponieważ jeden bajt wymaga mniej więcej 1,3 znaku w kodowaniu `base64`, wynikowy łańcuch znaków zawiera `nbytes * 1,3` znaków. Jeśli wartość `nbytes` nie jest podana, zwracana jest ta sama domyślna liczba bajtów co w funkcji `secrets.token_bytes(nbytes=None)`.

Inną ważną, choć często pomijaną funkcją jest `secrets.compare_digest(a, b)`. Porównuje ona dwa łańcuchy znaków lub obiekty bajtowe w sposób, który nie umożliwia napastnikowi oceny na podstawie czasu porównywania, czy wartości przynajmniej częściowo do siebie pasują. Porównywanie dwóch sekretów poprzez standardowe porównywanie łańcuchów znaków (za pomocą operatora `==`) jest podatne na atak przez pomiar czasu. W takim scenariuszu napastnik może sprawdzać wiele sekretów i na podstawie analiz statystycznych stopniowo odgadywać kolejne znaki oryginalnej wartości.

Co może się pojawić w przyszłości?

W czasie gdy powstaje ta książka, Python 3.9 nadal ma tylko kilka miesięcy. Możliwe jednak, że gdy Ty będziesz ją czytać, Python 3.10 będzie prawie gotowy lub nawet już dostępny.

Ponieważ proces rozwoju Pythona jest otwarty i transparentny, mamy wgląd (w dokumentach PEP) w to, co zostało zaakceptowane, a także w to, co zostało już zaimplementowane w wersjach alfa i beta. Dzięki temu możemy omówić funkcje, które pojawią się w Pythonie 3.10. Dalej znajdziesz krótkie omówienie najważniejszych zmian, których można oczekiwać w najbliższej przyszłości.

Tworzenie sumy typów za pomocą operatora |

W Pythonie 3.10 wprowadzone zostanie następne uproszczenie składni wskazówek dotyczących typów. Dzięki nowej składni łatwiej będzie tworzyć adnotacje określające sumę typów.

Python jest językiem z typowaniem dynamicznym i bez obsługi polimorfizmu. Dlatego funkcje mogą przyjmować ten sam argument, który w różnych wywołaniach może być innego typu, i poprawnie go przetwarzać, jeśli używane typy mają ten sam interfejs. Aby lepiej to zrozumieć, wróć do sygnatury funkcji, która umożliwia wyszukiwanie wartości ze słownika z kluczami tekstowymi bez uwzględniania wielkości liter:

```
def get_ci(d: dict[str, Any], key: str) -> Any: ...
```

Wewnątrz używana jest metoda `upper()` kluczy pobranych ze słownika. Jest to główny powód, dla którego typ argumentu `d` został zdefiniowany jako `dict[str, Any]`, a typ argumentu `key` jako `str`.

Jednak typ `str` nie jest jedynym typem wbudowanym udostępniającym metodę `upper()`. Innym typem z tą metodą jest `bytes`. Jeśli chcesz umożliwić stosowanie w funkcji `get_ci()` słowników z kluczami w postaci łańcuchów znaków i kluczami w postaci bajtów, musisz podać sumę dozwolonych typów.

Obecnie jedynym sposobem na podawanie sumy typów jest użycie wskazówki `typing.Union`. Ta wskazówka umożliwia zapisanie sumy typów `bytes` i `str` jako `typing.Union[bytes, str]`. Kompletna sygnatura funkcji `get_ci()` wyglądałaby więc tak:

```
def get_ci(
    d: dict[Union[str, bytes], Any],
    key: Union[str, bytes]
) -> Any:
    ...
```

Już ten zapis jest długi, a w bardziej skomplikowanych funkcjach sytuacja jest jeszcze gorsza. Dlatego w Pythonie 3.10 możliwe będzie podawanie sumy typów za pomocą operatora `|`. W przyszłości wystarczy prosty zapis:

```
def get_ci(d: dict[str | bytes, Any], key: str | bytes) -> Any: ...
```

W odróżnieniu od wskazówek dotyczących typów w typach generycznych wprowadzenie operatora sumy typów nie sprawi, że wskazówka `typing.Union` stanie się przestarzała. Dlatego obu konwencji będzie można używać wymiennie.

Strukturalne dopasowywanie wzorców

Strukturalne dopasowywanie wzorców jest najbardziej kontrowersyjną nową cechą Pythona w ostatniej dekadzie — i z pewnością także najbardziej skomplikowaną.

Akceptacja tego mechanizmu była poprzedzona licznymi zaciętymi dyskusjami i niezliczonymi szkicami projektu. Złożoność tego zagadnienia jest dobrze widoczna, jeśli przyjrzeć się wszystkim dokumentom PEP, które go dotyczą. Oto tabela wszystkich dokumentów PEP związanych ze strukturalnym dopasowywaniem wzorców (status z marca 2021 roku).

Data	PEP	Tytuł	Typ	Status
23 czerwca 2020	622	<i>Structural Pattern Matching</i>	Standardy	Zastąpiony przez PEP 634
12 września 2020	634	<i>Structural Pattern Matching: Specification</i>	Standardy	Zaakceptowany
12 września 2020	635	<i>Structural Pattern Matching: Motivation and Rationale</i>	Informacyjny	Wersja finalna
12 września 2020	636	<i>Structural Pattern Matching: Tutorial</i>	Informacyjny	Wersja finalna
26 września 2020	642	<i>Explicit Pattern Syntax for Structural Pattern Matching</i>	Standardy	Szkic
2 lutego 2021	653	<i>Precise Semantics for Pattern Matching</i>	Standardy	Szkic

To sporo dokumentów, a żaden z nich nie jest krótki. Czym więc jest strukturalne dopasowywanie wzorców i do czego może być przydatne?

W strukturalnym dopasowywaniu wzorców wprowadzono **instrukcję dopasowywania** i dwa nowe kontekstowe słowa kluczowe — `match` i `case`. Zgodnie z nazwami można ich używać do dopasowywania (`match`) określonej wartości do listy podanych przypadków (`case`) i wykonywać operacje na podstawie wyniku dopasowania.

Kontekstowe słowo kluczowe to słowo kluczowe, które nie w każdym kontekście jest zarezerwowane. Poza kontekstem instrukcji dopasowywania słowa `match` i `case` mogą być stosowane jak zwykłe nazwy zmiennych lub funkcji.

Niektórym programistom składnia instrukcji dopasowywania przypomina składnię instrukcji switch z języków takich jak C, C++, Pascal, Java i Go. Rzeczywiście, można ją stosować do zaimplementowania tego samego wzorca, ale daje ona zdecydowanie większe możliwości.

Ogólna (i uproszczona) składnia instrukcji dopasowywania wygląda tak:

```
match wyrażenie:
    case wzorzec:
        ...
```

Jako wyrażenie można zastosować dowolne poprawne wyrażenie Pythona; wzorzec reprezentuje dopasowywany wzorzec i jest nowością w Pythonie. W bloku case można umieścić wiele instrukcji. Złożoność instrukcji dopasowywania wynika głównie z wprowadzenia **dopasowywanych wzorców**, które początkowo mogą być trudne do zrozumienia. Wzorce łatwo też pomylić z wyrażeniami, jednak wzorce nie są przetwarzane w taki sam sposób jak zwykle wyrażenia.

Zanim szczegółowo omówimy dopasowywane wzorce, warto przyjrzeć się przykładowej instrukcji dopasowywania, która działa tak jak instrukcje switch z innych języków programowania:

```
import sys

match sys.platform:
    case "windows":
        print("Uruchomiono w systemie Windows")
    case "darwin":
        print("Uruchomiono w systemie macOS")
    case "linux":
        print("Uruchomiono w systemie Linux")
    case _:
        raise NotImplementedError(
            f"{sys.platform} nie jest obsługiwany!")
)
```

Jest to oczywiście bardzo prosty przykład, ale już uwidacznia kilka ważnych aspektów. Po pierwsze, jako wzorce mogą być stosowane literały. Po drugie, używany jest specjalny wzorzec wieloznaczny w postaci podkreślenia (_). Wzorce wieloznaczne i inne wzorce, które z powodu składni zawsze pasują do szukanych wartości, tworzą **domyślny blok case**. Można go umieścić tylko jako ostatni blok instrukcji dopasowywania.

Ten przykład można oczywiście zapisać za pomocą prostej sekwencji instrukcji if, elif i else. Typowe zadanie na rozmowach rekrutacyjnych na niższe stanowiska polega na napisaniu programu FizzBuzz.

Program FizzBuzz ma odliczać od 0 do dowolnej wartości i w zależności od sprawdzanej liczby wykonywać jedną z czterech operacji:

- wyświetlać Fizz, gdy wartość jest podzielna przez 3;
- wyświetlać Buzz, gdy wartość jest podzielna przez 5;
- wyświetlać FizzBuzz, gdy wartość jest podzielna jednocześnie przez 3 i 5;
- wyświetlać sprawdzaną liczbę we wszystkich pozostałych scenariuszach.

Jest to prosty problem, ale zaskakujące jest, jak wiele osób ma trudności nawet z najłatwiejszymi zadaniami pod presją stresu w trakcie rozmowy rekrutacyjnej. Rozwiązanie można oczywiście opracować za pomocą kilku instrukcji `if`, jednak instrukcja dopasowywania pozwala napisać kod cechujący się naturalną elegancją:

```
for i in range(100):
    match (i % 3, i % 5):
        case (0, 0): print("FizzBuzz")
        case (0, _): print("Fizz")
        case (_, 0): print("Buzz")
        case _: print(i)
```

W tym przykładzie wartość jest w każdej iteracji dopasowywana do warunku `(i % 3, i % 5)`. Potrzebne są obie operacje modulo, ponieważ efekt każdej iteracji zależy od obu wyników. Kod kończy sprawdzanie wzorców po znalezieniu pasującego bloku i wykonuje tylko jeden blok kodu.

Widoczną różnicą w porównaniu z poprzednim przykładem jest to, że teraz zastosowaliśmy wzorce w postaci sekwencji zamiast wzorców w postaci literalów:

- Wzorec `(0, 0)` pasuje do dwuelementowej sekwencji, w której oba elementy są równe 0.
- Wzorec `(0, _)` pasuje do dwuelementowej sekwencji, w której pierwszy element jest równy 0. Drugi element może być dowolną wartością dowolnego typu.
- Wzorec `(_, 0)` pasuje do dwuelementowej sekwencji, w której drugi element jest równy 0. Pierwszy element może być dowolną wartością dowolnego typu.
- Wzorec wieloznaczny `_` pasuje do wszystkich wartości.

W wyrażeniach dopasowywania można stosować nie tylko proste literały i ich sekwencje. Możesz też dopasowywać wartości do określonych klas. I to właśnie wzorce w postaci klas pozwalają tworzyć fantastyczny kod. Jest to zdecydowanie najbardziej skomplikowany aspekt omawianej funkcji.

W czasie gdy powstaje ta książka, Python 3.10 nie został jeszcze udostępniony. Dlatego trudno jest przedstawić typowe i praktyczne zastosowanie dopasowywania wzorców w postaci klas. Zamiast tego przyjrzymy się przykładowi z oficjalnego samouczka. Oto zmodyfikowany przykład z dokumentu PEP 636 obejmujący prostą funkcję `where_is()`, która potrafi dopasować wartość do struktury przekazanej instancji klasy `Point`:

```
class Point:
    x: int
    y: int

    def __init__(self, x, y):
        self.x = x
        self.y = y

def where_is(point):
    match point:
        case Point(x=0, y=0):
            print("Początek układu")
        case Point(x=0, y=y):
```

```

        print(f"Y={y}")
    case Point(x=x, y=0):
        print(f"X={x}")
    case Point():
        print("Inne miejsce")
    case _:
        print("To nie jest punkt")

```

W tym przykładzie dzieje się wiele rzeczy, dlatego warto omówić wszystkie występujące tu wzorce:

- `Point(x=0, y=0)` zostaje dopasowany, jeśli `point` jest instancją klasy `Point`, a atrybuty `x` i `y` mają wartość `0`.
- `Point(x=0, y=y)` zostaje dopasowany, jeśli `point` jest instancją klasy `Point`, a jej atrybut `x` ma wartość `0`. Atrybut `y` jest zapisywany w zmiennej `y`, którą można wykorzystać w bloku `case`.
- `Point(x=x, y=0)` zostaje dopasowany, jeśli `point` jest instancją klasy `Point`, a jej atrybut `y` ma wartość `0`. Atrybut `x` jest zapisywany w zmiennej `x`, którą można wykorzystać w bloku `case`.
- `Point()` zostaje dopasowany, jeśli `point` jest instancją klasy `Point`.
- `_` zostaje dopasowany zawsze.

Widać tu, że dopasowywanie wzorców pozwala sprawdzać atrybuty obiektu. Choć wzorec `Point(x=0, y=0)` wygląda jak wywołanie konstruktora, Python nie wywołuje konstruktorów w trakcie sprawdzania wzorców. Nie sprawdza też argumentów pozycyjnych ani nazwanych metody `__init__()`, dlatego w dopasowywanym wzorcu można uzyskać dostęp do wartości dowolnych atrybutów.

Przy dopasowywaniu wzorców można też posługiwać się składnią pozycyjną dla atrybutów, jednak wymaga to więcej pracy. Trzeba udostępnić w klasie dodatkowy atrybut `__match_args__`, który określa naturalną kolejność atrybutów instancji klasy. Oto przykład:

```

class Point:
    x: int
    y: int

    __match_args__ = ["x", "y"]

    def __init__(self, x, y):
        self.x = x
        self.y = y

def where_is(point):
    match point:
        case Point(0, 0):
            print("Początek układu")
        case Point(0, y):
            print(f"Y={y}")
        case Point(x, 0):
            print(f"X={x}")

```

```
case Point():  
    print("Inne miejsce")  
case _:  
    print("To nie jest punkt")
```

A to tylko mała próbka możliwości. Instrukcje dopasowywania są dużo bardziej skomplikowane, niż można to pokazać w tym krótkim podrozdziale. Jeśli mielibyśmy uwzględnić wszystkie zastosowania, warianty składni i przypadki brzegowe, musielibyśmy poświęcić im cały rozdział. Jeżeli chcesz dowiedzieć się więcej na temat instrukcji dopasowywania, koniecznie przeczytaj trzy „kanoniczne” dokumenty PEP: 634, 635 i 636.

Podsumowanie

W tym rozdziale omówiliśmy najważniejsze zmiany w składni języka i bibliotece standardowej wprowadzone w czterech ostatnich wersjach Pythona. Jeśli nie śledzisz aktywnie informacji o wersjach lub nie przeszedłeś jeszcze na Pythona 3.9, powinno to zapewnić Ci wystarczającą wiedzę o bieżącym stanie języka.

W tym rozdziale wprowadziliśmy też koncepcję „idiomów programistycznych”. Będziemy do niej wielokrotnie wracać w tej książce. W następnym rozdziale dokładnie przyjrzymy się wielu idiomom z Pythona, porównując wybrane mechanizmy z Pythona i z innych języków programowania. Jeśli jesteś doświadczonym programistą, ale dopiero od niedawna korzystasz z Pythona, będzie to doskonała okazja, aby się dowiedzieć, „jak się to robi w Pythonie”. Zobaczysz też, gdzie Python się wyróżnia, a gdzie wciąż pozostaje w tyle za konkurencją.

Skorowidz

A

- AABB, Axis-Aligned Bounding Box, 158
- ABI, application binary interface, 283
- abstrakcyjne klasy bazowe, 164, 168
- adnotacje, 82
 - funkcji, 90, 164
 - typów, 137, 164, 169
- aktualizacja słownika, 74
- algorytm
 - AABB, 158
 - HLL, 492
 - LRU, 495
 - Mersenne Twister, 97
 - mrówkowy, 488
- algorytmy
 - aproksymacyjne, 487, 488
 - genetyczne, 488
- alokowanie zasobów, 462
- analiza typów, 356
- API Python/C, 292
- API, Application Programming Interface, 156
- aplikacja Prometheus, 440
- aplikacje
 - internetowe, 394
 - sieciowe, 240
 - wielodostępne, 194, 195
 - wielowątkowe, 197
 - wykonywalne, 406
- architektura
 - rozproszona aplikacji, 449
 - sterowana zdarzeniami, 248
 - typu pull, 440
 - typu push, 439
 - usługowa, 280

- archiwum
 - Mailman 2, 23
 - Mailman 3, 23
- AsyncAPI, 251
- asynchroniczne
 - dostarczanie e-maili, 490
 - operacje wejścia – wyjścia, 221
- asynchroniczność, 220, 235
 - dostosowywanie nieasynchronicznego kodu, 228
- atak DoS, 209
- atrapy, mock, 345
 - obiektów, 347
- atrybuty, 116
 - klasy, 113
- aukcje RTB, 286
- automatyzacja kontroli jakości, 349

B

- backend, 178
 - RedisBackend, 338
- baza
 - PostgreSQL, 235
 - Redis, 339
- bezpieczeństwo, 398, 414
- biblioteka, 156, 368
 - blinker, 246, 249
 - ctypes.CDLL, 312
 - ctypes.OleDLL, 313
 - ctypes.PyDLL, 313
 - ctypes.WinDLL, 313
 - dict_helpers.py, 94

biblioteka

- injector, 181
- OpenGL, 121
- PyOpenGL, 122
- requests, 226
- Tk, 234, 237
- tkinter, 239
- unittest.mock, 345

biblioteki

- dynamiczne, 287, 312
- rejestrowania logów, 422
- wczytywanie, 312

błędy, 418, 435

- deklaracji, 163
- strukturalne, 163
- w wątkach, 206

C**C**

- debugowanie pamięci, 479

C++ , 132

CalVer, 386

cechy, traits, 157

Celery, 491

CFFI, 318

ChainMap, 76

ciąg Fibonacciego, 292

Clang, 281

cx_Freeze, 411

Cython, 304

- instalowanie, 305
- jako język, 307

D

dane uwierzytelniające, 383

data i czas, 365

dealokacja obiektu, 303

debuger pdb, 68

debugowanie, 91, 311

- poawaryjne, 68

dekorator, 149, 173

- @app.route(route), 175
- @autorepr, 259
- @dataclass, 256, 257
- @login_required, 150
- @lru_cache, 256
- @pytest.fixture, 340
- @pytest.mark.parametrize, 332
- @runtime_checkable, 170, 256, 257

@wraps, 151

dataclass, 140

functools singledispatch(), 136

report.register(), 136

dekoratory klas, 256

dekorowanie nazw, 116

delegowanie zadań, 194, 196

deskryptor

- danych, 117

- niedanych, 117

diagram referencji

- cyklicznych, 476, 479

- przechowywanych przez obiekty, 475

- wstecznych, 475

Docker, 28, 44, 396

- instrukcje, 45

- kontenery, 396

- przenośność aplikacji, 397

- tworzenie obrazów, 397

Docker Compose, 50, 51, 54, 445

- komunikacja, 54, 56

dokumenty PEP, 21, 385

domknięcia, 173

DoS, Denial of Service, 209

dostęp do atrybutów, 116

drzewa składni abstrakcyjnej, 274

DSN, Data Source Name, 436

duck typing, 157

dynamiczne generowanie kodu, 274

dystrybucje

- bdist, 379

- sdist, 378

- wheel, 379

dziedziczenie, 106, 110

E

eksporter wskaźników, 443

F

FaaS, Function as a Service, 249

fabryka klas, 485

Falcon, 277

FIFO, First In First Out, 201, 484

filtrowanie danych, 435

filtry, 421

FizzBuzz, 100

Flask, 46

fora dyskusyjne, 24

format HTML, 352

formatery, 421
 formatowanie
 komunikatów, 426
 łańcuchów znaków, 95
 metoda `str.format()`, 95
 obiekty `f-string`, 95
 symbol `%`, 95
 funkcja, *Patrz także* metoda
 `basicConfig()`, 432
 `breakpoint()`, 68, 69, 90
 `CFUNCTYPE()`, 316, 317
 `compile()`, 273
 `ctypes.util.find_library()`, 314
 `dictConfig()`, 432
 `echo()`, 47
 `eval()`, 273
 `exec()`, 273
 `fibonacci()`, 148
 `filter()`, 144
 `find_library()`, 314
 `get_ci()`, 83
 `getattr()`, 257
 `glob()`, 244
 `instance_repr()`, 258
 `Lock()`, 192
 `map()`, 144
 `next()`, 148
 `partial()`, 180
 `pause()`, 365
 `pdb.pm()`, 68
 `print()`, 200
 `Py_BuildValue()`, 299
 `qsort()`, 316
 `reduce()`, 144
 `repr_instance()`, 257
 `secrets.token_bytes()`, 97
 `secrets.token_hex()`, 97
 `secrets.token_urlsafe()`, 97
 `send()`, 345, 346, 347, 348
 `setattr()`, 121
 `time.time()`, 365
 `WINFUNCTYPE()`, 317
 `worker()`, 202
 funkcje
 częściowe, 146
 czyste, 142
 haszujące, 500
 języka C, 314
 `lambda`, 142
 obsługi żądań, 175
 pierwszoklasowe, 142
 przeciążanie, 133, 135

przekazywanie argumentów, 83, 173
 widoku, 47

G

GCC Linux, 281
 generator, 147
 pseudolosowy, 97
 generowanie
 dat i czasu, 365
 kodu, 273, 274, 276
 realistycznych danych, 363
 getter, 117, 123
 GIL, Global Interpreter Lock, 92, 193, 300
 GNU Debugger, 69
 graf
 acykliczny, 87
 cykliczny, 87
 nieskierowany, 86
 referencji, 89
 skierowany, 86
 graficzny interfejs użytkownika, GUI, 234, 236
 Graphviz, 458
 Guppy-PE, 473

H

haczyki ścieżki importu, 276
 hermetyzacja, 123
 heurystyki, 487, 488
 hierarchia
 klas, 108, 109
 loggerów, 424
 typów, 357
 HLL, HyperLogLog, 491
 Hy, 278

I

IDE, Intergrated Development
 Environment, 254
 identyfikator PID, 213
 idiom, 104, 154
 programistyczny, 75
 inicjalizacja modułu, 294
 instalowanie
 pakietów, 29, 387
 pakietów w trybie edycji, 388
 instancja klasy, 113
 instrukcja `RUN`, 54

- instrukcje Dockera, 45
- integrowanie
 - kodu, 286
 - zewnętrznych bibliotek, 287
- interfejs, 155
 - CLI, 67
 - do obsługi wielowątkowości, 219
- interfejsy
 - funkcji obcych, 281
 - niejawne, 167
 - responsywne, 194
 - tworzenie, 169
- IPython, 65
- izolacja
 - na poziomie aplikacji, 33, 34
 - na poziomie systemu, 33
 - środowiska na poziomie systemu, 41

J

- Jaeger, 451, 452
- język
 - C, 282
 - C++, 132, 282
 - DOT, 475
 - Hy, 278
 - Lisp, 278
 - opisu grafów, 475
 - YAML, 251
- Jupyter, 65

K

- kanały, 246
- katalog pakietów
 - dystrybucji, 31
 - użytkownika, 31
- klasa, 113
 - Callable, 169
 - ChainMap, 76
 - Container, 168
 - Executor, 229
 - FlaskTracing, 453
 - Future, 229
 - Grepper, 243
 - Hashable, 169
 - InstanceCountingClass, 261
 - Iterable, 168
 - logger, 419, 420
 - multiprocessing.Queue, 215
 - object, 111

- ObserverABC, 242
- Point, 101
- Process, 214
- Queue, 201
- queue.Queue, 215
- RedisModule, 182
- SelfWatch, 246
- Sized, 169
- SubjectABC, 242
- Throttle, 211
- Vector, 138
- ZoneInfo, 86
- klasy
 - abstrakcyjne, 164, 168
 - bazowe, 106
 - danych, 138, 141
 - mieszane, 260
 - pochodne, 106
 - złożoności, 459
- klasyfikatory Trove, 376
- klauszula if isinstance(...), 136
- kolejka
 - FIFO, 201, 484
 - LIFO, 484
- kolejki
 - dwukierunkowe, 204
 - komunikatów, 249, 490
 - komunikatów z brokerem, 250
 - zdarzeń, 249
- kolejkowanie zadań, 489
- kolekcja __dict__, 119
- kompilator
 - Clang, 281, 282
 - Cython, 304
 - GCC, 281, 282
- komponenty Prometheusa, 440
- komunikacja
 - międzyprocesowa, 214, 248
 - sterowana zdarzeniami, 238
 - z bibliotekami dynamicznymi, 312
- komunikaty, 234
- konfiguracje testów, 334
 - lokalne, 335
 - monkeypatch, 348
 - parametryzacja funkcji, 340
 - współdzielone, 335
 - wstrzykiwanie zależności, 339
 - z opcją autouse, 336
 - z wtyczek, 335
 - zasięgi, 336

- konfigurowanie
 - czasu wykonywania, 402
 - rejestrowania logów, 427, 429
 - złożonych środowisk, 50
- konstruktor, 140
- kontener Dockera, 56, 396
- kontenery, 48
 - wstrzykiwania zależności, 182
 - zmniejszanie wielkości, 52
- konteneryzacja, 43
- kontrola
 - jakości, 349
 - typów, 157, 169
- konwencja wywołań, 296, 298
 - METH_KEYWORDS, 296
 - METH_NOARGS, 296
 - METH_O, 296
 - METH_VARARGS, 296
- korutyny, 223–225
- krotki, 485
- Kubernetes, 397

L

- lambda, 142
- leniwe wartościowanie, 120
- liczby pseudolosowe, 97
- LIFO, Last In, First Out, 484
- lintery, 353, 354
- listy, 110, 481
 - dwukierunkowe wiązane, 483
 - mailingowe, 24
- literały liczbowe, 96
 - podkreślenia, 96
- loggery, 419, 420
 - tworzenie, 428
- logi, 418
- LRU, Least Recently Used, 256, 495
- luźne powiązanie, 248, 250

Ł

- łańcuch znaków, 95

M

- makra preprocesora, 300
- makro
 - Py_DECREF(), 302
 - Py_INCREF(), 302
 - Py_XDECREF(), 302
 - Py_XINCREF(), 302

- makroprofilowanie, 465
- manifest Twelve-Factor App, 394, 396
- maszyna wirtualna Javy, 105
- mechanizm
 - GIL, 193, 221, 300
 - odśmieciania pamięci, 478
 - wykrywania usług, 440
- Memcached, 458
- memoizacja, 493, 494, 495
- menedżer pamięci, 302
- metadane pakietu, 374
- metahaczyki, 276
- metaheurystyki, 488
- metaklasy, 263, 266
 - pułapki, 270
 - stosowanie, 267
- metaprogramowanie, 253, 273
 - oparte na introspekcji, 254
 - oparte na kodzie, 255
- metoda
 - __add__(), 131
 - __call__(), 130, 266
 - __contains__(), 130
 - __del__(), 130, 477
 - __delitem__(), 107
 - __dir__(), 94
 - __get__(), 130
 - __getattr__(), 94
 - __getattribute__(), 117
 - __getitem__(), 107, 130
 - __init__(), 261, 265
 - __init_subclass__(), 271
 - __iter__(), 130
 - __len__(), 130
 - __mul__(), 135
 - __new__(), 261, 262, 265
 - __prepare__(), 265
 - __repr__(), 257
 - __set__(), 130
 - __setitem__(), 107
 - __sub__(), 132
- bind(), 238
- critical(), 420
- debug(), 420
- error(), 420
- exception(), 420
- Grepper.grep(), 243
- info(), 420
- join(), 190, 214
- move_to(), 365
- os.environ.get(), 400

metoda

- setFormatter(), 425
- signal(), 247
- signal.connect(), 247
- signal.send(), 247
- start(), 214
- stats(), 179
- submit(), 230
- tick(), 365
- track(), 179
- warning(), 420

metody specjalne, 129

- migotanie, thrashing, 463

- mikrooptymalizacje pamięci, 486

- mikroprofilowanie, 465, 469

moduł

- backends.py, 343
- collections, 76, 482
- collections.abc, 168
- concurrent.futures, 229
- ctypes, 312, 314
- environ-config, 401
- graphlib, 74, 86, 87
- hashlib, 499
- itertools, 147
- logging, 419, 437
- mailer, 346
- multiprocessing, 213, 218
- multiprocessing.dummy, 219
- objgraph, 473, 476
- pdb, 68
- pickle, 215
- secrets, 96
- settings.py, 402
- threading, 200
- tkinter, 237
- typing, 82
- unittest.mock, 345
- zoneinfo, 74, 85

moduły

- inicjalizacja, 294
- struktura, 403
- wady, 403
- z ustawieniami, 404

- modyfikowanie działania funkcji, 255

- monitorowanie wskaźników, 439

- monkey patching, 119, 259, 347, 357

- MRO, 108, 111, 258

- MyInstaller, 409

N

narzędzia

- do poprawiania stylu, 353, 354
- do testowania, 363
- do zwiększania produktywności, 63, 69

narzędzie

- ack-grep, 71
- autojump, 70
- black, 355
- Celery, 491
- cloc, 71
- coverage, 350, 352, 353
- curl, 70
- cx_Freeze, 407, 411
- cythonize, 305
- Docker Compose, 50, 445
- Docker, 28, 44
- FizzBuzz, 100
- Flask, 46
- GNU Debugger, 69
- GNU paralel, 71
- Graphviz, 458
- Guppy-PE, 473
- HTTPIe, 70
- jq, 70
- Memcached, 458, 498
- memory_profiler, 473
- Memprof, 473
- mutmut, 361
- MyInstaller, 409
- objgraph, 473
- pip, 29, 305, 384
- py2app, 408, 413
- py2exe, 408, 413
- PyInstaller, 407, 408
- Pympler, 473
- RQ, 491
- Sentry, 437
- setuptools, 380
- Twine, 383
- Vagrant, 28, 44, 61
- Valgrind, 311, 480
- VirtualBox, 28
- wait-for-it, 58
- watchmedo, 59
- Werkzeug, 443
- yq, 70

notacja

- dużego O, 460
- węzowa, 268

- numery wersji, 383

O

- obiekt, 113
 - DatagramHandler, 423
 - FileHandler, 422
 - future, 228
 - HTTPHandler, 423
 - NullHandler, 422
 - RotatingFileHandler, 422
 - SMTPHandler, 423
 - SocketHandler, 423
 - StreamHandler, 422
 - SysLogHandler, 422
 - TimedRotatingFileHandler, 422
 - tracer, 452
- obiekty
 - częściowe, 146
 - falsywe, 342
 - f-string, 95
 - future, 229
 - iterowalne, 147
 - obsługi logów, 421, 426, 427
- obliczenia ewolucyjne, 488
- obrazy
 - Dockera, 53
 - kontenera, 49
- obserwatory, 242, 245
- obserwowanie obiektów, 242
- obsługa
 - asynchronicznych zadań, 491
 - błędów, 93
 - dealokacji, 477
 - kontenerów Dockera, 396
 - logów, 422, 426
 - pamięci, 180, 337, 340
 - połączeń, 342
 - sygnałów, 246
 - wątków, 192
 - wielowątkowości, 219
 - wstrzykiwania, 182
 - wyjątków, 298
- odporność na awarie, 248
- odśmiecianie pamięci, 478
 - ze śledzeniem, 301
- odwrócenie sterowania, 172, 173, 179
- OpenAPI, 251
- OpenCensus, 451
- OpenGL Shading Language (GLSL), 122
- OpenTelemetry, 451
- OpenTracing, 451

- operacje
 - blokujące, 190, 463
 - wejścia – wyjścia, 221, 463
- operator
 - |, 75, 98
 - | =, 75
 - +, 132
 - + =, 191
 - części wspólnej, 74
 - morsa, 80
 - różnicy, 74
 - sumy, 74
 - XOR, 74
- operatory
 - przeciążanie, 129
- ostrzeżenia, 92

P

- pakiet
 - acme_sdk, 343
 - aiohttp, 186
 - blinker, 234
 - coverage, 322
 - cx_Freeze, 368
 - Cython, 305
 - egg, 380
 - environ-config, 401
 - faker, 322
 - falcon, 254
 - flask, 28, 155, 234
 - flask-injector, 155, 181
 - Flask-OpenTracing, 418
 - freezegun, 322, 418, 424
 - gprof2dot, 458
 - Graphviz, 474
 - hy, 254
 - inflection, 254
 - injector, 155
 - ipdb, 28
 - ipython, 28
 - jaeger-client, 418, 452
 - macropy3, 254
 - mutmut, 322
 - mypy, 73, 155, 322
 - objgraph, 458
 - opentracing-python, 452
 - poetry, 28
 - prometheus-client, 418, 444
 - py2exe, 368
 - pyinstaller, 368

- pakiet
 - pymemcache, 458
 - pymemcache, 498
 - pyright, 73
 - pytest, 322
 - redis, 155, 322
 - redis_opentracing, 418, 453
 - requests, 186
 - sentry-sdk, 418
 - setuptools, 393
 - twine, 368
 - tzdata, 85
 - wait-for-it, 28
 - watchdog, 28, 59
 - wheel, 368, 380
 - zope.interface, 155
- pakiety, 29
 - aplikacji, 394
 - bibliotek, 368
 - budowa, 369
 - instalowanie, 29, 387
 - metadane, 374
 - przestrzeni nazw, 388
 - publikowanie, 381
 - punkty wejścia, 390
 - rejestrwanie, 381
 - rodzaje dystrybucji, 377
 - standardowe, 85
 - tryb edycji, 388
 - tworzenie, 367
 - usług, 394
 - wersjonowanie, 383
 - wheel binarne, 381
 - własne, 387
 - zewnętrzne, 85
- pamięć
 - LRU, 256
 - podręczna, 492
 - narzędzie Memcached, 498
 - równoważnika obciążenia, 498
 - serwera reverse proxy, 498
 - sieci CDN, 498
 - wyniki deterministyczne, 493
 - wyniki niedeterministyczne, 496
 - zarządzanie, 497
- parametry czysto pozycyjne, 83
- parametryzacja testów, 331
- parsowanie argumentów, 297
- PEP, Python Enhancement Proposal, 21
- pętla zdarzeń, 230
- piksele śledzące, 174
- pisanie
 - rozszerzeń, 288
 - testów, 325
- pixel-tracking, 452
- platforma
 - Falcon, 277
 - Flask, 46
 - pytest, 325, 327, 334, 348
 - zakupowa, 286
 - zope.interface, 157
- plik
 - .bash_profile, 91
 - .pypirc, 382, 383
 - application.log, 425
 - docker-compose.yml, 51, 54, 442, 445
 - CHANGELOG.md, 371
 - Dockerfile, 45, 46, 442, 445
 - fibonacci.c, 305
 - fibonacci.pyx, 308
 - LICENSE, 371
 - MANIFEST.in, 371, 373
 - poetry.lock, 41
 - pyproject.toml, 39
 - Python.h, 281
 - README.md, 370
 - settings.py, 402
 - setup.cfg, 350, 371, 373
 - setup.py, 371, 382
 - tracking.py, 441
 - Vagrantfile, 61
- pliki blokady zależności, 40
- podatność na awarię, 251
- podkatalogi, 369
- podział
 - czasu procesora, 192
 - programu, 212
 - przetwarzania, 188
- Poetry, 37
- pokrycie kodu testami, 349, 352
- polecenia, 241
 - dodatkowe, 372
 - standardowe, 372
- polecenie
 - assert, 331
 - black, 355
 - coverage, 350
 - coverage html, 352
 - docker images, 49
 - docker-compose, 446
 - docker-compose up, 57
 - install, 387

- mypy, 357
- pip install, 59
- pytest, 329
- polimorfizm, 173
 - ad hoc, 133
 - dynamiczny, 127
- porządek MRO, 108, 111, 167
- PostgreSQL, 235
- powłoka IPython, 65
- powłoki niestandardowe, 63
- preprocesor makra, 300
- probabilistyczne struktury danych, 491
- problemy NP-zupełne, 487
- proces tworzenia instancji klasy, 260
- procesy, 213
 - stosowanie puli, 217
- produktywność, 63, 69
- profiler
 - deterministyczny, 465
 - statystyczny, 465
- profilowanie
 - kodu, 464
 - procesora, 465
 - wykorzystania pamięci, 472
- programowanie
 - asynchroniczne, 189, 220, 225
 - funkcyjne, 141
 - efekty uboczne, 141
 - funkcje czyste, 142
 - funkcje pierwszoklasowe, 142
 - przejrzystość referencyjna, 142
 - obiektywne, 105, 113
 - oparte na tematach, 246
 - sterowane sygnałami, 246
 - sterowane testami, 322
 - sterowane zdarzeniami, 195, 234
 - obserwowanie obiektów, 242
 - style, 240
 - wywołania zwrotne, 241
- Prometheus, 440
 - komponenty, 440
- protokół
 - deskryptorów, 130
 - iterowania, 130
 - kontenera, 130
 - OpenRTB, 286
 - OpenTelemetry, 451
 - sekwencji, 130
 - wywoływania, 130
- przeciążanie
 - funkcji, 133, 135
 - operatorów, 129

- przejrzystość referencyjna, 142
- przekazywanie
 - argumentów, 83, 173
 - własności, 302
- przestrzenie nazw, 388
- przeszukiwanie tabu, 488
- przetwarzanie
 - odroczone, 489
 - w tle, 188, 194, 196
- przypisanie, 78
- pula
 - procesów, 217
 - tokenów, 209
 - wątków, 201
- punkty wejścia, 390
- py2app, 413
- py2exe, 413
- PyInstaller, 408
- PyPI, Python Package Index, 27
- Python Software Foundation (PSF), 22

R

- raport o pokryciu kodu testami, 352
- Redis, 178, 339
- referencje, 302, 475
 - pożyczone, 302
 - skradzione, 303
- referer, 174
- rejestrwanie
 - błędów, 435
 - logów, 419, 430, 435
 - komponenty, 420
 - konfigurowanie, 427
 - platek śniegu, 434
 - rozproszone, 433
 - scentralizowane, 434
 - sieć, 434
 - ujednolicone, 434
 - zalecane praktyki, 430
 - zdarzeń, 419
- REPL, read-eval-print loop, 67
- repozytorium
 - PyPI, 27, 29
- responsywność aplikacji, 188, 194
- responsywny interfejs, 193
- rozszerzenia, 310
 - języka C, 283, 285, 288, 289
 - koszty, 310
- rozszerzenie .py, 392
- równoważnik obciążenia, 498
- RQ, 491

S

SaaS, Software as a Service, 394
 scalanie, 74
 słowników, 75
 SDK, 343
 SemVer, 385
 Sentry, 436, 437
 serwer
 Prometheusa, 445
 reverse proxy, 498
 SMTP, 346
 VPS, 43
 WWW Prometheusa, 447
 setter, 117, 123
 shebang, 406
 singleton, 182
 skalowalność, 248
 składnia metaklas, 264
 skrypty, 390
 nakładkowe, 392
 rozruchowe, 64
 słowa kluczowe kontekstowe, 99
 słowniki, 107
 rozpakowywanie, 76
 scalanie, 75
 słowo kluczowe
 async, 222
 await, 222
 case, 99
 match, 99
 sortowanie, 318
 topologiczne, 88
 specyfikatory wersji, 383
 sprawdzanie typów, 135
 Stackless Python, 283
 standard
 OpenTelemetry, 451
 OpenTracing, 451
 statyczna analiza typów, 356
 sterowanie zdarzeniami, 235
 strategie próbkowania zdarzeń, 452
 struktura, 294
 aplikacji, 402
 danych, 480
 lista, 481
 moduł collections, 482
 probabilistyczna, 491
 zbiór, 482
 modułów, 403
 plików, 369, 410

strukturalne dopasowywanie wzorców, 99
 suma typów, 98
 sygnały, 246
 anonimowe, 247
 symulowane wyżarzanie, 488
 synchroniczne dostarczanie e-maili, 489
 system rejestrowania logów, 421
 systemy
 operacyjne, 156
 rozproszone, 448
 wielojęzyczne, 280
 sytuacja wyścigu, 191
 szyfr Cezara, 237

Ś

ślady, 456
 śledzenie
 rozproszone, distributed tracing, 448, 450
 rozproszone za pomocą Jaegera, 451
 środowisko
 produkcyjne, 42
 REPL, 67
 uruchomieniowe, 31
 wirtualne, 33, 44
 tworzenie, 37
 złożone, 50

T

tablica
 struktur, 294
 z haszowaniem, 343
 TDD, test-driven development, 322
 tematy, 246
 testowanie mutacyjne, 358, 363
 testy, 323, 325
 integracyjne, 342
 jednostkowe, 342
 konfiguracja, 334, 336
 parametryzacja, 331
 sprawdzenia jakości, 358
 struktura, 327
 throttling, 209
 tokeny, 209
 traser platformy Falcon, 277
 tryb
 edycji, 388
 roboczy, 92
 ostrzeżenia, 92

tworzenie
 aplikacji wykonywalnych, 406
 interfejsów, 169
 kontenerów, 50
 loggera, 428
 loggera bazowego, 424
 modułów, 181
 obrazów Dockera, 397
 obrazu, 49
 pakietów aplikacji, 394
 pakietów bibliotek, 368
 sumy typów, 98
 środowisk wirtualnych, 37
 typów pochodnych, 169
 wzorców projektowych, 155

typ
 defaultdict, 482, 484
 deque, 482, 483
 namedtuple, 482, 485, 487
 typing.Protocol, 169

typy danych, 314
 generyczne, 81
 niestandardowe, 287

U

usługa
 echo-server, 55, 58
 Memcached, 498
 prometheus, 446
 Redis, 178

usługi, 394
 komunikacja, 54
 sieciowe, 156
 uruchamianie, 57

utrwalanie danych, 177, 250
 używanie rozszerzeń, 285

V

Vagrant, 28, 44, 61
 Valgrind, 311, 480
 VirtualBox, 28
 Visual Studio 2019, 281
 VPS, Virtual Private Server, 43

W

wątki, 192, 195
 błędy, 206
 stosowanie puli, 201

Werkzeug, 444
 wersjonowanie
 kalendarzowe, 385, 386
 semantyczne, 385

wiązania, 296
 widzety, 236, 241
 wielodziedziczenie, 106, 108
 wieloprosowość, 188, 212
 wielowątkowość, 188, 189
 moduł multiprocessing.dummy, 219
 wielozadaniowość kooperatywna, 221

wirtualizacja
 lekka, 43
 maszyn, 43

wirtualne środowiska programistyczne, 61
 wizualizacja rozproszonego śladu, 455
 właściwości, 123

WSGI, Web Server Gateway Interface, 196

wskaźnik
 typu Gauge, 443
 typu Info, 443
 typu Summary, 443

wskaźniki
 biznesowe, 439
 monitorowanie, 439
 obciążenia, 438
 wydajności, 439
 zużycie zasobów, 438

współbieżność, 187
 naturalna, 250

wstrzykiwanie zależności, 173, 180, 185, 339

wyciek
 pamięci, 472
 w C, 479
 zasobów, 463

wydajność, 280, 285, 457
 algorytmy aproksymacyjne, 487
 heurystyki, 487
 kolejkowanie zadań, 489
 moduł objgraph, 474
 nadmierne wykorzystanie zasobów, 462
 nadrezerwacja pamięci, 463
 operacje blokujące, 463
 operacje wejścia – wyjścia, 463
 probabilistyczne struktury danych, 491
 profilowanie kodu, 464
 profilowanie procesora, 465
 przetwarzanie odroczone, 489
 wyciek zasobów, 463
 złożoność kodu, 459
 zmniejszanie złożoności, 480

- wyjątek, 293, 298, 437
 - ValueError, 300
- wykonawca, 229
- wykrywanie usług, 440
- wyliczenia, 151
 - nazwa, 152
 - składowa, 152
 - symboliczne, 153
 - wartość, 152
- wyrażenia
 - generatora, 148
 - przypisania, 78
- wysyłanie, shipping, 394
- wyszukiwanie ścieżek, 276
- wywołania zwrotne, 241
 - w C, 315
- wyzwalacze, 236
- wzorce, 102
 - dopasowywanie strukturalne, 99
 - dostępu do atrybutów, 116
 - projektowe, 154
 - wieloznaczne, 100
- wzorzec
 - command [...], 58, 60
 - projektowy obserwator, 242
 - projektowy singleton, 182
- zasada Hollywood, 173
- zbiory, 482
- zdarzenia, 234, 419
 - strategie próbkowania, 452
- zliczanie referencji, 302
- złożoność, 310, 460
 - cyklomatyczna, 459
 - McCabe'a, 459
- zmienna środowiskowa, 398, 402
 - BIND_HOST, 401
 - BIND_PORT, 400, 401
 - DATABASE_URI, 401
 - PYTHONFAULTHANDLER, 93
 - PYTHONSTARTUP, 63
 - SCHEDULE_INTERVAL_SECONDS, 400
- zope.interface, 157

Ż

źródła informacji, 25

Ż

żądania HTTP, 239

Z

- zakres, spa, 451
- zależności
 - przechodnie, 40
 - zewnętrzne, 41
- zapis wewnątrzwierszowy, 175
- zarządzanie
 - pamięcią podręczną, 497
 - pulami wątków, 202
 - środowiskami wirtualnymi, 44
 - zależnościami, 383
 - zmiennymi środowiskowymi, 398

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 

Koduj wszystko w Pythonie. Obiektywnie, strukturalnie i funkcjonalnie!

Python cechuje się dużą prostotą, a przy tym jest wszechstronny. Ma bardzo szeroki zakres zastosowania, przez co coraz więcej osób podejmuje naukę programowania w tym języku. Python należy do języków najczęściej używanych przez programistów, którzy tworzą w nim gry i aplikacje webowe. Świetnie sprawdza się ponadto w pracy z wykorzystaniem sztucznej inteligencji i uczenia maszynowego. Tym, co programiści doceniają w Pythonie, jest też obiektywność. Ucząc się, przyswajamy bowiem również zasady programowania obiektywnego, a więc koncepcji dla wielu innych języków.

Oto książka, którą docenią i osoby rozpoczynające przygodę z programowaniem, i programiści znający już inne języki. Znajdziesz tu zarówno podstawowe informacje o Pythonie, jak i wskazówki dotyczące pisania rozszerzeń, dzięki którym będziesz w stanie korzystać z atutów kilku języków. Przydatnym uzupełnieniem są liczne przykłady, pokazujące, jak rozwiązywać częste problemy. To już czwarte wydanie tego praktycznego podręcznika — docenianego za to, że pozwala dobrze poznać Pythona i uczy, jak pisać wydajny i czytelny kod.

W książce między innymi:

- jakie są najważniejsze usprawnienia w Pythonie
- jak przeprowadzić izolację środowiska
- jak używać najnowszych funkcji w Pythonie
- czym Python się różni od innych języków
- co to jest współbieżność i wielowątkowość
- na czym polega programowanie sterowane zdarzeniami
- jakie są elementy metaprogramowania
- jak przeprowadzić automatyzację kontroli jakości
- jak optymalizować kod

Helion 	<i>Sprawdź nasze szkolenia!</i>	KOD KORZYŚCI <i>Sięgnij po więcej!</i> 	
 helion.pl	SZKOLENIA 	ISBN 978-83-283-8743-0	
 HELION SA ul. Kościuszki 1c 44-100 Gliwice tel.: 32 230 98 63 helion@helion.pl	AKADEMIA IT & BUSINESS HELIONSZKOLENIA.PL	 9 788328 387430	
INFORMATYKA W NAJLEPSZYM WYDANIU		Cena: 109,00 zł	

Packt