

Zapiski s priprav

Podatkovne strukture

ACM priprave na tekmovanja

Različica z dne 12. april 2025

Kazalo

1	Seznami	4
1.1	Vektorji	7
1.2	Skladi	8
1.3	Vrste	11
2	Množice	13
2.1	Slovarji	15
2.2	Neurejene množice in slovarji	17
3	Kopice	18
4	Primerjava	19

Pogosto si moramo za rešitev različnih problemov podatke organizirati drugače, kot da jih le shranimo v običajen seznam. Z uvedbo posebnih *podatkovnih struktur* lahko hitreje odgovarjamo na različna vprašanja o shranjenih informacijah, kar lahko uporabimo kot sestavno komponento v raznih algoritmih. Poglejmo si primer. Recimo, da imamo dan seznam števil dolžine n , nato pa bi radi q -krat ugotovili, če je neko število v tem seznamu.

```
1  #include <stdio.h>
2
3  int arr[1000];
4
5  int main() {
6      int n, q;
7      scanf("%d%d", &n, &q);
8      for (int i = 0; i < n; i++) {
9          scanf("%d", &arr[i]);
10     }
11
12     for (int j = 0; j < q; j++) {
13         int k;
14         scanf("%d", &k);
15         bool najdeno = false;
16         for (int i = 0; i < n; i++) {
17             if (arr[i] == k) {
18                 printf("Število %d se pojavi v seznamu.\n", k);
19                 najdeno = true;
20                 break;
21             }
22         }
23         if (!najdeno) {
24             printf("Števila %d ni v seznamu.\n", k);
25         }
26     }
27     return 0;
28 }
```

Če si podatke le shranimo v seznam, kot zgoraj, potem bomo v najslabšem primeru za vsako od poizvedb morali preiskati celoten seznam, torej je naša rešitev $O(nq)$. Za hitrejšo rešitev bi lahko seznam na začetku uredili, nato pa pojavitve števil k iskali z bisekcijo:

```
1  #include <stdio.h>
2  #include <algorithm>
```

```

3  using namespace std;
4
5  int arr[1000];
6
7  int main() {
8      int n, q;
9      scanf("%d%d", &n, &q);
10     for (int i = 0; i < n; i++) {
11         scanf("%d", &arr[i]);
12     }
13
14     sort(arr, arr+n);
15
16     for (int j = 0; j < q; j++) {
17         int k;
18         scanf("%d", &k);
19         int idx = lower_bound(arr, arr+n, k) - arr;
20         if (arr[idx] == k)
21             printf("Stevilo %d se pojavi v seznamu.\n", k);
22         else
23             printf("Stevila %d ni v seznamu.\n", k);
24     }
25
26     return 0;
27 }

```

Takšna rešitev ima časovno zahtevnost $O((n + q) \log n)$, saj za urejanje porabimo $O(n \log n)$ časa, vsaka izmed q poizvedb pa se izvede v $O(\log n)$ korakih. Povedano drugače smo v zgornjem primeru opazili naslednje dejstvo: Operacija „preveri, ali se dan podatek nahaja v podatkovni strukturi“ ima zahtevnost $O(n)$ v seznamu in $O(\log n)$ v urejenem seznamu. Kaj pa lahko povemo o drugih operacijah, ki bi si jih želeli uporabljati v algoritmih?

V nadaljevanju bomo za različne podatkovne strukture obravnavali časovno zahtevnost naslednjih operacij:

- (O1) dodaj podatek v podatkovno strukturo,
- (O2) odstrani podatek iz podatkovne strukture,
- (O3) preveri, če je podatek v podatkovni strukturi.

Število podatkov, shranjenih v podatkovni strukturi, bomo označili z n .

1 Seznami

Najbolj osnovna podatkovna struktura je običajen seznam. V računalniškem spominu je predstavljen z enim velikim blokom spomina, ki si ga lahko predstavljamo kot n spremenljivk istega tipa, položenih eno za drugo. Ko deklariramo seznam, moramo nujno podati njegovo dolžino, da lahko prevajalnik rezervira blok primerne velikosti. Kakor smo povedali v uvodnem razdelku poglavja, za operacijo (O3) potrebujemo $O(n)$ časa. Pri določanju časovne zahtevnosti za operaciji (O1) in (O2) pa naletimo na težavo. Seznami namreč ne podpirajo dodajanja ali odstranjevanja podatkov — število elementov seznama je fiksno in določeno v kodi. Lahko pa se pretvarjamo, da seznam podpira ti operaciji tako, da si zapomnimo trenutno število podatkov v seznamu:

```
1 // V spremenljivki arr hranimo elemente seznama (in nekaj praznega
2 // prostora), v spremenljivki velikost pa število teh elementov.
3 // Začnemo s praznim seznamom
4 int arr[1000];
5 int velikost = 0;
6
7 // Funkcija doda element k na i-to mesto našega seznama
8 void dodaj_element(int k, int i) {
9     // Zapisati želimo arr[i] = k, a tega še ne smemo narediti,
10    // ker bomo s tem prepisali element, ki je tam trenutno shranjen.
11    // Zato moramo prvo prekopirati vse elemente na eno mesto desno.
12    for (int j = velikost; j > i; j--) {
13        arr[j] = arr[j-1];
14    }
15    arr[i] = k;
16    velikost++; // dodali smo element, zato je naš seznam za 1 daljši
17 }
18
19 // Funkcija odstrani element na indeksu i iz našega seznama
20 void odstrani_element(int i) {
21     // Če želimo odstraniti i-ti element, moramo poskrbeti, da bodo
22     // vsi kasnejši elementi na pravem mestu, torej da bo prvi "prazen"
23     // prostor na mestu arr[velikost], in ne prej
24     for (int j = i; j < velikost; j++) {
25         arr[j] = arr[j+1];
26     }
27     velikost--;
28 }
```

Iz zgornje kode je jasno razvidno, da imata tako operaciji (O2) kot (O3) časovno zahtevnost $O(n)$, saj moramo pri obeh v najslabšem primeru premakniti vse elemente.

Opazimo pa lahko tudi, da potrebujeta operaciji „dodaj na konec seznama“ in „odstrani iz konca seznama“ le konstantno mnogo časa, saj se v obeh primerih zanka takoj konča. To bo pomembno kasneje.

Druga težava z zgornjo implementacijo pa je ta, da naš seznam lahko shrani največ 1000 elementov. Seveda bi lahko pri deklaraciji `int arr[1000]` le povečali številko tako, da bi bila dovolj velika za naše potrebe, vendar tega včasih ne želimo narediti (na primer če imamo veliko majhnih seznamov in nekaj velikih, a ne vemo v naprej, kateri od teh seznamov bodo veliki). Lahko bi si namesto fiksnega seznama `arr` hranili kazalec na del spomina, kjer je seznam shranjen, ter širino tega dela spomina, ter ga po potrebi povečali, kot spodaj:

```
1  int *arr = nullptr;
2  int kapaciteta = 0; // število elementov, ki jih lahko shranimo v arr
3  int velikost = 0; // število elementov, trenutno shranjenih v seznamu
4
5  void dodaj_element(int k, int i) {
6      if (velikost == kapaciteta) {
7          // Ne moremo dodati novega elementa, saj imamo rezerviranega
8          // premalo spomina.
9          // Na začetku si torej sposodimo več spomina, in v nov spomin
10         // prekopiramo stare podatke
11         int *nov_arr = (int*) malloc((kapaciteta + 1) * sizeof(int));
12         for (int j = 0; j < velikost; j++) {
13             nov_arr[j] = arr[j];
14         }
15         kapaciteta++;
16
17         // Starega bloka spomina ne potrebujemo več, zato ga vrnemo
18         // operacijskemu sistemu
19         if (arr != nullptr) free(arr);
20         arr = nov_arr;
21     }
22     // Nadaljujemo tako kot prej
23     for (int j = velikost; j > i; j--) {
24         arr[j] = arr[j-1];
25     }
26     arr[i] = k;
27     velikost++;
28 }
```

Funkcija `odstrani_element` je v tem primeru enaka kot prej. Napisali smo *dinamični seznam*, tj. seznam, ki prilagaja svojo hitrost glede na število shranjenih elementov. S

tem pa smo izgubili zgoraj omenjeno lastnost, da je dodajanje na oz. brisanje iz konca seznama $O(1)$, saj bo tekom zgornjega programa velikost vedno enaka kapaciteti, zato se bo prvi pogojni stavek v funkciji vedno aktiviral. To lastnost pa lahko še vedno v neki obliki obdržimo, če naredimo majhno spremembo:

```
1 void dodaj_element(int k, int i) {
2     if (velikost == kapaciteta) {
3         // Ne moremo dodati novega elementa, saj imamo rezerviranega
4         // premalo spomina.
5         // Na začetku si torej sposodimo več spomina, in v nov spomin
6         // prekopiramo stare podatke
7         int nova_kapaciteta = 2*kapaciteta;
8         if (nova_kapaciteta == 0) nova_kapaciteta = 1;
9         int *nov_arr = malloc(nova_kapaciteta * sizeof(int));
10        for (int j = 0; j < velikost; j++) {
11            nov_arr[j] = arr[j];
12        }
13        kapaciteta = nova_kapaciteta;
14
15        // Starega bloka spomina ne potrebujemo več, zato ga vrnemo
16        // operacijskemu sistemu
17        if (arr != nullptr) free(arr);
18        arr = nov_arr;
19    }
20    // Nadaljujemo tako kot prej
21    for (int j = velikost; j > i; j--) {
22        arr[j] = arr[j-1];
23    }
24    arr[i] = k;
25    velikost++;
26 }
```

Ko `arr` povečamo, torej njegovo kapaciteto podvojimo. Kaj pa smo s tem dejansko dosegli? Ko v seznam dodamo prvi element, ne potrebujemo ničesar kopirati. Ko dodamo drugi element, moramo prekopirati prvega. Ko dodamo tretji element, moramo prekopirati prvega in drugega. Pri četrtem elementu pa nam ni treba ničesar kopirati, saj je `arr` že dovolj velik, da lahko četrti element takoj shranimo vanj. Za peti dodan element bomo morali spet prekopirati prve štiri, za šesti, sedmi in osmi element pa nam ne bo treba ničesar kopirati. Zapišimo še nekaj členov tega zaporedja (koliko elementov moramo prekopirati, ko dodamo nov element v seznam):

0, 1, 2, 0, 4, 0, 0, 0, 8, 0, 0, 0, 0, 0, 0, 16, ...

Z nekaj razmisleka vidimo, da je vsota prvih n členov tega zaporedja enaka $2n$, če je n ravno potenca števila 2, oziroma n v vseh ostalih primerih. V povprečju bomo torej za dodajanje enega elementa na konec seznama porabili le $O(1)$ časa. Pravimo, da je ta operacija *asimptotično konstantna* — v povprečju potrebujemo konstanten čas, da jo izvedemo, vsake toliko pa bo operacija trajala dlje.

1.1 Vektorji

K sreči ne potrebujemo vsakič znova pisati svojih dinamičnih seznamov, saj so že vključeni v standardni knjižnici C++ pod imenom **vector**. Poglejmo si, kako delujejo.

```
1  #include <stdio.h>
2  #include <vector>
3  #include <algorithm> // sort
4  using namespace std; // nujno za vector
5
6  int main() {
7      vector<int> seznam;
8
9      for (int i = 0; i < 1000; i++) {
10         seznam.push_back(i);
11     }
12
13     for (int i = 0; i < seznam.size(); i++) {
14         if (seznam[i] > 500) {
15             seznam[i] = 7;
16         }
17     }
18
19     seznam.insert(seznam.begin() + 852, 17);
20     printf("%d\n", seznam[852]);
21     seznam.erase(seznam.begin() + 852);
22     printf("%d\n", seznam[852]);
23
24     sort(seznam.begin(), seznam.end());
25     printf("%d\n", seznam[852]);
26
27     seznam.clear();
28     return 0;
29 }
```

Ko deklariramo vektor, v trikotnih oklepajih (znaka za manjše oz. večje) dopišemo še tip podatkov, ki jih bomo v vektor shranjevali. V tem primeru uporabimo

kar `int`, lahko pa bi uporabili kakršenkoli drugi tip, na primer `char`, `long long*`, `vector<vector<int>>` itd. Potem nam je na voljo več metod, ki jih uporabimo tako, da napišemo ime metode za imenom spremenljivke in piko:

- metoda `size()` vrne trenutno velikost vektorja,
- metoda `push_back(x)` doda `x` na konec vektorja v $O(1)$,
- metoda `begin()` vrne iterator na začetek vektorja,
- metoda `end()` vrne iterator tik čez konec vektorja,
- metoda `insert(it, x)` doda `x` tik pred mesto `it` v $O(n)$,
- metoda `erase(it)` izbriše element na mestu `it` v $O(n)$,
- metoda `clear()` izbriše vse elemente vektorja.

Iterator je poseben tip podatkov, ki si ga lahko predstavljamo kot kazalec (ker tudi podpira običajne operacije kazalcev). Iterator dobimo iz metod `begin()` in `end()`, moramo pa ga podati kot argument metodama `insert()` in `erase()`. Iteratorjem lahko prištevamo številke, s čimer lahko dostopamo do poljubnega elementa vektorja (ali pa to naredimo tudi z oglatimi oklepaji).

1.2 Skladi

Sklad je podatkovna struktura, ki podpira le dve operaciji: „dodaj element na vrh sklada“ in „odvzemi element z vrha sklada“. Predstavljamo si, da elemente zlagamo na kup, iz katerega lahko jemljemo le vrhnji element, ker bi se kup sicer prevrnil.

```
1  #include <stdio.h>
2  #include <stack>
3  using namespace std;
4
5  int main() {
6      stack<int> sklad;
7
8      // Za dodajanje na sklad uporabimo metodo push()
9      sklad.push(3);
10     sklad.push(7);
11     sklad.push(9);
12     sklad.push(4);
13
14     // Z metodo top() lahko pogledamo vrhnji element sklada
15     printf("%d\n", sklad.top());
16
17     // Za odstranjevanje s sklada uporabimo metodo pop()
```

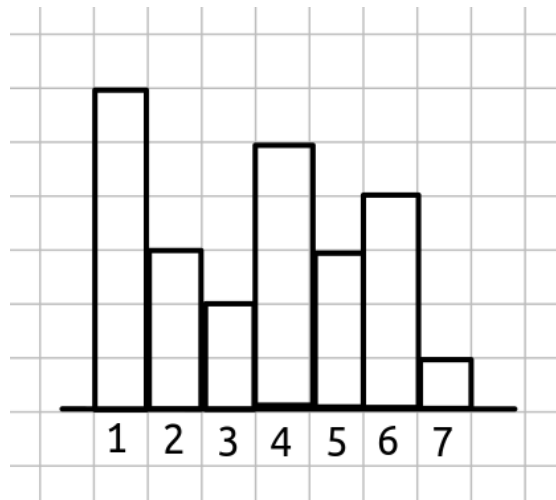
```

18     sklad.pop();
19     sklad.pop();
20
21     printf("%d\n", sklad.top());
22     return 0;
23 }

```

Poleg metod `push()` in `pop()` za dodajanje oz. odstranjevanje elementov imamo na voljo tudi metodi `top()`, ki pogleda, kaj je shranjeno na vrhu sklada, in `size()`, ki vrne število elementov na skladu. Vse metode imajo časovno zahtevnost $O(1)$.

Poglejmo si primer naloge, kjer v rešitvi uporabimo sklad. Zadan nam je naslednji problem: Dan je seznam števil, ki opisuje višino stolpnice, postavljenih v vrsto ena za drugo. Za vsako stolpnico ugotovi, katera je njej najbližja stolpnica na levi, ki je višja od nje. Za tretjo stolpnico na spodnji sliki je to stolpnica številka 2, za četrto stolpnico pa je odgovor enak 1, saj sta stolpnici 2 in 3 nižji od stolpnice 4.



```

1  #include <stdio.h>
2
3  int arr[1000];
4
5  int main() {
6      int n;
7      scanf("%d", &n);
8      for (int i = 0; i < n; i++) {
9          scanf("%d", &arr[i]);
10         bool najdena = false;

```

```

11     for (int j = i-1; j >= 0; j--) {
12         if (arr[j] > arr[i]) {
13             printf("%d\n", j);
14             najdena = true;
15             break;
16         }
17     }
18     if (!najdena) {
19         printf("Ni visje stolpnice na levi.\n");
20     }
21 }
22 return 0;
23 }

```

Zgornji program bo rešil dan problem na očiten način v $O(n^2)$. Obstaja pa tudi linearna rešitev, ki se poslužuje dveh skladov. Osnovna ideja je, da si hranimo padajoče zaporedje stolpnice, ki so kandidati za najbližjo višjo stolpnico na levi. Ko pridemo do stolpnice, ki je nižja od najnižjega kandidata, lahko za njo takoj zatrdimo, da je najnižji kandidat ravno iskani odgovor. Če pa je najnižji kandidat višji od stolpnice, ki jo trenutno obravnavamo, vemo, da ne bo ta kandidat nikoli več najbližja višja stolpnica, saj je trenutna stolpnica hkrati višja od kandidata in bližje preostalim stolpnicam. Torej lahko kandidate, ki so nižji od trenutne stolpnice, kar odstranimo iz sklada.

V rešitvi imamo dve gnezdeni zanki (**while** znotraj **for**). Zakaj je potem še vedno linearna? Enostaven argument je, da bomo vsako stolpnico enkrat dodali v sklad, in jo zato največ enkrat odstranili iz sklada, torej bomo v najslabšem primeru opravili $2n$ operacij na skladu.

```

1  #include <stdio.h>
2  #include <stack>
3  using namespace std;
4
5  stack<int> sklad, indeksi;
6
7  int main() {
8      int n;
9      scanf("%d", &n);
10     for (int i = 0; i < n; i++) {
11         int t;
12         scanf("%d", &t);
13         while (sklad.size() > 0 && sklad.top() < t) {
14             sklad.pop();
15             indeksi.pop();

```

```

16     }
17     if (sklad.size() == 0) {
18         printf("Ni visje stolpnice na levi.\n");
19     } else {
20         printf("%d\n", indeksi.top());
21     }
22     sklad.push(t);
23     indeksi.push(i);
24 }
25 return 0;
26 }

```

1.3 Vrste

Skladom sorodne podatkovne strukture so *vrste*. Vrsta podpira le dodajanje elementov na konec in odvzemanje elementov z začetka. Tako v nasprotju v skladom pri vrstah elemente odstranjujemo v istem vrstnem redu, v katerem smo jih dodali.

V C++-u imamo na voljo podatkovni tip `queue`, ki podaja implementacijo vrste. Za dodajanje elementov v vrsto uporabljamo metodo `push()`, za odvzemanje metodo `pop()`, za dostop do prvega elementa pa metodo `front()`. Vse te operacije imajo časovno zahtevnost $O(1)$.

```

1  #include <stdio.h>
2  #include <queue>
3  using namespace std;
4
5  int main() {
6      queue<int> vrsta;
7
8      // Za dodajanje v vrsto uporabimo metodo push()
9      vrsta.push(3);
10     vrsta.push(7);
11     vrsta.push(9);
12     vrsta.push(4);
13
14     // Z metodo front() lahko pogledamo prvi element vrste
15     printf("%d\n", vrsta.front()); // izpiše 3
16
17     // Za odstranjevanje iz vrste uporabimo metodo pop()
18     vrsta.pop();
19     vrsta.pop();

```

```

20
21     printf("%d\n", vrsta.top()); // izpiše 9
22     return 0;
23 }

```

Razmislimo, kako bi napisali svojo implementacijo vrste. Lahko bi uporabili razširljiv seznam iz začetka poglavja, kjer bi za dodajanje uporabili funkcijo `dodaj_element` in ji vedno podali $i = \text{velikost} - 1$, za odstranjevanje elementov pa bi uporabili funkcijo `odstrani_element`, ki bi ji vedno podali argument $i = 0$. Problem je, da bi potem za odstranjevanje elementov vedno potrebovali $O(n)$ časa, saj bi vedno morali prekopirati celoten seznam.

V spodnji kodi se poslužimo naslednje ideje: Namesto, da se vrsta vedno začne na indeksu 0 tabele `arr`, bomo pustili, da se tudi indeks začetka spreminja. Potem lahko element vrste „odstranimo“ preprosto tako, da povečamo ta indeks; podatek bo sicer še vedno shranjen v spominu, a bomo mi vedeli, da ga moramo ignorirati.

```

1  int *arr = nullptr;
2  int velikost = 0;
3  int kapaciteta = 0;
4  int zacetek = 0;
5
6  int prvi_element() {
7      return arr[zacetek];
8  }
9
10 // V vrsto vedno dodajamo elemente na konec
11 void dodaj_element(int k) {
12     if (velikost == kapaciteta) {
13         int nova_kapaciteta = 2*kapaciteta;
14         if (nova_kapaciteta == 0) nova_kapaciteta = 1;
15         int *nov_arr = malloc(nova_kapaciteta * sizeof(int));
16
17         // Prekopiramo seznam na novo lokacijo, pri čemer lahko
18         // nastavimo nov začetek na 0
19         for (int j = 0; j < velikost; j++) {
20             nov_arr[j] = arr[(zacetek+j)%kapaciteta];
21         }
22         kapaciteta = nova_kapaciteta;
23         zacetek = 0;
24
25         if (arr != nullptr) free(arr);
26         arr = nov_arr;

```

```

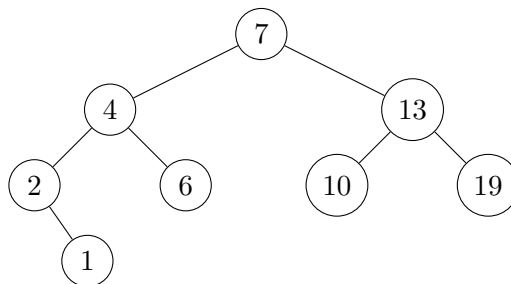
27     }
28
29     // Na kateri indeks dodamo element?
30     int i = (zacetek + velikost) % kapaciteta;
31     arr[i] = k;
32     velikost++;
33 }
34
35 // Iz vrste vedno odstranimo prvi element
36 void odstrani_element() {
37     zacetek++;
38     velikost--;
39 }

```

2 Množice

Vrnimo se k osrednjemu vprašanju, hitrosti operacij (O1), (O2) in (O3). V seznamu, vrsti in skladu smo podatke imeli organizirano zaporedoma, eden za drugim (kakor so v seznamu shranjeni v spominu), in s tem uspeli doseči konstantno časovno zahtevnost za dodajanje in odstranjevanje elementov, operacija (O3) pa še vedno traja $O(n)$, ker moramo pogledati vsak element posebej.

Če si podatke organiziramo v drevesno strukturo, lahko dosežemo hitrost $O(\log n)$ na vseh treh operacijah hkrati. Na primer, če želimo shraniti številke 1, 2, 4, 6, 7, 10, 13, 19, jih lahko organiziramo kakor v spodnji sliki, tako da vsaka črta na levo pomeni manjšo številko, črta na desno pa pomeni večjo številko:



Potem lahko preverimo, če je število k vsebovano v naši strukturi tako, da začnemo na vrhu (v sedmici), in se prestavimo levo, če je $k < 7$, oziroma desno, če je $k > 7$. Potem postopek ponavljamo, dokler lahko, oziroma dokler ne najdemo števila k . V najslabšem primeru preverimo toliko števil, kolikor je vrstic v drevesu (na zgornji sliki 4). Če števila pametno razporedimo, lahko poskrbimo, da bo višina drevesa približno $\log n$.

Tudi operaciji dodajanja in odvzemanja števil iz take strukture lahko implementiramo

tako, da bosta imeli časovno zahtevnost $O(\log n)$, vendar sta bolj komplicirani, zato se vanju ne bomo podrobno spuščali.

Takšno drevesno podatkovno strukturo v jeziku C++ uporabljajo *množice*. Vsi elementi v množici so različni — tudi če isti element dodamo večkrat, bo shranjena le ena ponovitev. Množice definiramo tako kot vektorje, na voljo pa so nam naslednje metode:

- `size()` vrne število elementov v množici,
- `clear()` izbriše vse elemente iz množice,
- `insert(x)` vstavi `x` v množico,
- `count(x)` vrne 1, če je `x` v množici, in 0 sicer,
- `erase(x)` izbriše `x` iz množice.

Poglejmo si primer uporabe.

```
1  #include <stdio.h>
2  #include <set>
3  using namespace std;
4
5  int main() {
6      set<int> mnozica;
7
8      mnozica.insert(1);
9      mnozica.insert(2);
10     mnozica.insert(2);
11     mnozica.insert(4);
12     mnozica.insert(6);
13     mnozica.insert(6);
14     mnozica.insert(6);
15     mnozica.insert(7);
16     mnozica.insert(10);
17     mnozica.insert(10);
18     mnozica.insert(13);
19     mnozica.insert(19);
20
21     printf("V mnozici je %d elementov.\n", mnozica.size());
22
23     if (mnozica.count(10)) {
24         printf("10 je v mnozici.\n");
25     } else {
26         printf("10 ni v mnozici.\n");
27     }
28 }
```

```

29     mnozica.erase(10);
30
31     if (mnozica.count(10)) {
32         printf("10 je v mnozici.\n");
33     } else {
34         printf("10 ni v mnozici.\n");
35     }
36
37     return 0;
38 }

```

Če želimo pregledati vse elemente množice, tega žal ne moremo narediti z zanko kot pri seznamih, saj ne moremo dostopati do posamičnih elementov — vprašanje „Kaj je *i*-ti element?“ sploh nima smisla pri množicah, saj elementi niso postavljeni zaporedno. Še vedno pa lahko pregledamo vse elemente v množici z zanko naslednje oblike:

```

1 for (auto it : mnozica) {
2     printf("%d\n", *it);
3 }

```

V tem primeru je `it` iterator na elemente množice in ima tip `set<int>::iterator`. Da nebi toliko pisali, lahko uporabimo tudi `auto`.

2.1 Slovarji

Včasih želimo podatke obravnavati v parih (ključ, vrednost), in se hitro sklicati na vrednost pri danem ključu. To lahko dosežemo tako, da ustvarimo množico oblike

```

1 set<pair<tip_kljuca, tip_vrednosti>> mnozica;

```

kjer podamo tip za ključ in vrednosti. Potem do ključa v paru dostopamo s `par.first`, do vrednosti pa s `par.second`.

V C++ pa nam je na voljo podatkovna struktura, ki to naredi namesto nas: *slovar*. Ko ga ustvarimo, moramo podati tip za ključ in tip za vrednosti v slovarju, potem pa lahko do shranjenih vrednosti dostopamo tako, da zapišemo ključ v oglatih oklepajih.

```

1 #include <stdio.h>
2 #include <map>
3 using namespace std;
4

```

```

5  int main() {
6
7      map<char, int> slovar;
8
9      // Dodamo nekaj vrednosti
10     slovar['a'] = 3;
11     slovar['b'] = 6;
12     slovar['g'] = 7;
13     slovar['8'] = 8;
14
15     // Vse shranjene podatke lahko pregledamo s for-each zanko
16     for (auto p : slovar) {
17         // tip spremenljivke p je v tem primeru pair<char, int>
18         printf("%c: %d\n", p.first, p.second);
19     }
20
21     // Do podatkov lahko dostopamo z oglatimi oklepaji
22     printf("Vrednost pri 'g': %d\n", slovar['g']);
23
24     // Preverimo, če ima nek ključ shranjeno vrednost
25     if (slovar.count('c')) {
26         printf("'c' je v slovarju.\n");
27     } else {
28         printf("'c' ni v slovarju.\n");
29     }
30
31     // Podatke lahko odstranjujemo po ključih
32     slovar.erase('b');
33
34     // Na voljo sta nam tudi metodi lower_bound in upper_bound, ki
35     // poiščeta prvi ključ v slovarju, večji ali enak (oziroma večji)
36     // danemu iskanemu ključu.
37     // Funkciji vrneta iterator, ki kaže na par (ključ, vrednost),
38     // oziroma na slovar.end(), če ne najdeta ničesar.
39     map<char, int>::iterator it = slovar.lower_bound('c');
40     printf(
41         "Najmanjsi ključ, večji od 'c': %c. Vrednost = %d\n",
42         (*it).first,
43         (*it).second);
44
45     return 0;
46 }

```

Slovarje lahko uporabimo tudi kot razpršene tabele, torej kot sezname, kjer so indeksi lahko poljubno veliki ali majhni, in kjer nekateri indeksi manjkajo. Običajnega seznama tu ne moremo uporabiti, saj bi lahko potrebovali zelo veliko tabelo — če so indeksi med -10^9 in 10^9 , moramo imeti tabelo z več kot dvema milijardama elementov, ki bi zavzela več gigabajtov spomina. Velikost slovarja pa bo odvisna le od števila elementov, ki jih dejansko shranimo.

```
1  #include <stdio.h>
2  #include <map>
3  using namespace std;
4
5  int main() {
6      map<int, int> arr;
7
8      // Preberemo podatke
9      int n;
10     scanf("%d", &n);
11     for (int i = 0; i < n; i++) {
12         int x, y;
13         scanf("%d%d", &x, &y);
14         arr[x] = y; // O(log n) namesto O(1) pri seznamu
15     }
16
17     // Sedaj imamo podatke shranjene, in lahko z njimi naredimo,
18     // kar algoritem zahteva.
19     // Za demonstracijo jih samo izpišemo, kjer vidimo tudi nekaj
20     // nove sintakse.
21     for (auto[x, y] : arr) {
22         printf("%d: %d\n", x, y);
23     }
24
25     return 0;
26 }
```

2.2 Neurejene množice in slovarji

Množice in slovarji, kot smo jih spoznali zgoraj, si elemente na pameten način shranijo urejene po vrsti, s čimer dosežejo časovno zahtevnost $O(\log n)$ za operacije (O1), (O2) in (O3). Obstajata pa tudi neurejeni različici teh struktur, ki podpirata iste operacije (z izjemo `lower_bound` in `upper_bound`), a imata časovno zahtevnost $O(1)$ za naštetje

operacije. V delovanje teh struktur se ne bomo spuščali, lahko pa seveda uporabljamo C++ implementaciji teh struktur.

```
1  #include <stdio.h>
2  #include <unordered_set>
3  #include <unordered_map>
4  using namespace std;
5
6  int main() {
7      unordered_set<int> mnozica;
8      unordered_map<char, int> slovar;
9
10     for (int i = 0; i < 20; i++) {
11         mnozica.insert(2*i+7);
12         slovar[i] = 2*i+7;
13     }
14
15     slovar.erase(0);
16     mnozica.erase(7);
17
18     for (int i = 0; i < 20; i++) {
19         if (mnozica.count(i)) {
20             printf("%d\n", slovar[i]);
21         }
22     }
23
24     return 0;
25 }
```

3 Kopice

Kopica ali *prioritetna vrsta* deluje podobno kot običajna vrsta, le da namesto prvo dodanega elementa vedno odstranimo največji element iz kopice. Na voljo so nam podobne operacije kot pri vrstah: `push(x)` doda `x` na kopico, `pop()` odstrani največji element iz kopice, `top()` pa vrne največji element v kopici. Prvi dve naštetim metodam imata časovno zahtevnost $O(\log n)$, zadnja pa $O(1)$. Da to dosežemo, tudi pri kopici elemente hranimo v drevesni strukturi, ki pa je malo drugačna kot struktura pri množicah.

```
1  #include <stdio.h>
2  #include <queue> // priority_queue se najde v <queue>
3  using namespace std;
```

```

4
5 int main() {
6     priority_queue<int> vrsta;
7
8     vrsta.push(3);
9     vrsta.push(7);
10    vrsta.push(5);
11
12    while (vrsta.size() > 0) {
13        // v prioritetni vrsti uporabimo .top() namesto .front()
14        printf("%d\n", vrsta.top());
15        vrsta.pop();
16    }
17
18    return 0;
19 }

```

Če želimo, da namesto največjega elementa delamo z najmanjšim elementom na kopici, ji moramo ob definiciji dodati še nekaj posebnih argumentov:

```

1 priority_queue<int, vector<int>, greater<int>> vrsta;

```

4 Primerjava

Spoznane strukture lahko primerjamo glede na to, katere operacije podpirajo, in kako hitre so te operacije. V spodnji tabeli je prikazana časovna zahtevnost pogostih operacij na različnih strukturah. Operacija „dodaj“ pomeni, da v strukturo dodamo nek element, „poišči“ pa da preverimo, ali se dani element nahaja v naši strukturi. Ostaneta še operaciji „odstrani“ in „odstrani x “. Prva od njih iz strukture odstrani *nek* element — kateri element bo odstranjen je odvisno od podatkovne strukture (npr. pri skladu bo to vrhni element, pri vrsti pa prvi element). Operacija „odstrani x “ pa poišče dan element x v strukturi in ga odstrani, če je prisoten. Če podatkovna struktura ne podpira neke operacije, je to v tabeli označeno z X.

Operacija	Seznam	Vektor	Sklad	Vrsta	Množica	N. množica	Kopica
dodaj	X	$O(1)$	$O(1)$	$O(1)$	$O(\log n)$	$O(1)$	$O(\log n)$
odstrani	X	X	$O(1)$	$O(1)$	X	X	$O(\log n)$
odstrani x	X	$O(n)$	X	X	$O(\log n)$	$O(1)$	X
poišči	$O(n)$	$O(n)$	X	X	$O(\log n)$	$O(1)$	X