

HASHING

① provides insert, delete, ~~also~~ search operations in $O(1)$ time

② No duplicates are allowed.

③ Hashing is not used

- when we have to find closer values.
 - data is arranged in sorted fashion
 - prefix searching → "ge" in "geeksforgeeks"
- } AVL or Red black tree

Application of Hashing

Hash table is the second most used data structure

- Implementing dictionaries
- Database indexing
- cryptography
- caches
- symbol tables in compilers / interpreters
- Routers
- Getting data from databases

Direct Address table

We use keys as indexes of an array. on insert operation $arr[key] = 1$ on delete operation

$arr[key] = 0$ and on search we just check if $(arr[key] = 1)$

return true;

else

return false;

• If we have a range of $[0 - n]$ we can make an array of $arr[n]$

• If we have range from $[m - n]$ we can use $arr[m - n]$
 $\neq arr[key - m] \leftarrow 1$

0	1	10	100	1000	10000	100000	1000000
---	---	----	-----	------	-------	--------	---------

table 17

- since away has random access

delet(i) {

table Lin = 0;

et $\text{search}(i)$ }

return table[i]

next li) {

$$table[i] = 1$$

3

۴

5

Problems

① large range of keys

⑧ Floating point numbers (as keys are used as index & index can only be positive integer).

⑧ keys can be strings.

Hashing

Traversing
Take large tree universes of keys and by using
traverse function convert them to small values.

- * Every time hash function must produce same value for same key

* should generate value between 0 to $m-1$.

$n \rightarrow$ size of hash table

* Hash function should be fast

should uniformly distribute large keys into Hash table slots.

Examples of Hash Function

⑦ nl range-key) - range-key % m.

→ m is chosen as prime no
• to have less common factors

For strings $\rightarrow$ weighted sum

① $\rightarrow$ sum of acid values of all the

characteristics of strong.
Problem is we have some under-
for acid, acid, acid...

② → str → "abcd"

$$stu[0] \times x^0 + stu[1] \times x^1 + stu[2] \times x^2 + stu[3] \times x^3.$$

→ choose a number x

→ at producing unique keys most of the time

check a
number 2

③ Universal Training

- group of hash function, you pick one hash function randomly.

collusion handling

It is possible that two large values map to the same key. This situation is called as collision.

② If data is known to us in advance, we can use perfect Hedging to guarantee zero collation.

① If we do not know keys in advance, To handle collision, we have two techniques

→ Chaining
→ open Addressing

Linear probing

→ quadratic probing

→ Double Flanking

Chaining

Hash function

Hash Table $\rightarrow$ Array of linked list headers.

0	21	49	15	4
1	50	15	42	6
2	58	17	23	
3				
4	25			
5				
6				
7				

Hash (value) $\rightarrow$ value % 7
 keys $\rightarrow$ {50, 21, 58, 17, 15, 49, 56, 23, 25, 4}.

Basic formula is to link the expected value with the head.

50 % 7 $\rightarrow$ 1
21 % 7 $\rightarrow$ 0
58 % 7 $\rightarrow$ 2
17 % 7 $\rightarrow$ 3
15 % 7 $\rightarrow$ 1
49 % 7 $\rightarrow$ 0
56 % 7 $\rightarrow$ 0

Performance of chaining

$m \rightarrow$ no of slots
 $n \rightarrow$ no of keys to be inserted
 load factor $\alpha = n/m$.

If hash table is small / load factor is big $\rightarrow$ more no of collision.

Average chain length $\rightarrow \alpha$

Expected time for search $\rightarrow$ $O(1 + \alpha)$
 or insert $\rightarrow$ $O(1) \rightarrow$ hash function computation

$O(\alpha) \rightarrow$ traversing in array of length α

Under case where keys will be the same key. map to same index. hence chain length $\rightarrow$ 2 computation for 0th.

OR worse scenario max key will uniformly distributed

$\alpha \rightarrow$ Trade-off between 1st & 2nd time

Data structures to store chains

Linked list $\rightarrow$ search $O(n)$ delete $O(n)$ insert $O(n)$ } But it is cache unfriendly

Dynamic size $\rightarrow$ search/insert/delete $O(n)$ } But it provides cache friendliness

Self Balancing BST $\rightarrow$ search/insert/delete $O(\log n)$ } But does not provide cache friendliness

Implementation of linked list

Bucket $\rightarrow$ 7

0	NULL	70	56
1	NULL	71	
2	NULL	9	72
3	NULL	NULL	
4	NULL	NULL	
5	NULL	NULL	
6	NULL	NULL	

Basic structure of classes

struct MyNode {
 int BUCKET;

list <int> *table; } use to implement linked list
 MyNode (int b) {
 Bucket = b;

table = new list <list> [b];

void insert (int key) {

use wait prevent size array of linked list using new.

we can use push-back to insert values in linked list

void insert (int key) {

int i = key % bucket;

table[i].push-back (key);

void remove (int key) {

int i = key % bucket;

table[i].remove (key);

bool search (int key) {

int i = key % bucket;

for (auto x : table[i])

if (x == key) return true;

return false;

Open Addressing

No of slots in the hash table $\geq$ no of keys to be inserted

Advantage $\rightarrow$ cache friendly.

Ex $\rightarrow$ key $\rightarrow$ { 50, 51, 49, 16, 56, 15, 19 }

0	49
1	50
2	51
3	16
4	56
5	15
6	19

hash(key) = key % 7

linearly search for next empty slot when there is collision

- If the collision occurs at first index we search for the next slot in linear sequential fashion

0	
1	50
2	51
3	15
4	
5	
6	

insert (50), insert (51), insert (15)
search (15), delete (15), search (15)

① In search operation, we first search in 1st index $i = \text{key} \% 7$

$\rightarrow$ and if arr[i] is not empty & contains some other value

we linearly move in search of x we stop at search when

(i) arr[i] = x or return true

(ii) arr[i] = null $\rightarrow$ return false

(iii) or we traverse back at the same slot

In delete operation

we cannot simply make a slot empty as during searching we even for empty slots and this may modify our results

Ex

2
3
4

delete $\rightarrow$

2
4

let $i = 1$ & $x = 4$
we start at $i = 1$ arr[i] = 2
move linearly
arr[i] = null
 $\rightarrow$ stop
return false

hence we mark the slot as deleted slot (not true empty slot)

Problems with linear probing

• clusters are formed

• suppose we have a collision $\rightarrow$ cluster of size 2

• if another key maps to the index of any one of them, cluster will grow $\rightarrow$ 3

• because of this if k is large $k+1 \rightarrow$ all insert/search/delete operation become costly



We can write linear probing as

$$\text{hash}(\text{key}, i) = (\text{hash}(\text{key}) + i) \% 7$$

→ i is incremented till 6 and if we don't find any empty slot it shows hash table is full.

Quadratic Probing

The problem of clusters are removed to much extent in this, instead of linearly searching for next slot, we search for i^2 , (next 1st slot), 2^2 ($i+4$), 3^2 ($i+9$) ...

* Secondary clusters were formed.

$$\text{hash}(\text{key}, i) = (\text{hash}(\text{key}) + i^2) \% m$$

Disadvantages

There may be a case when quadratic probing will not find an empty slot even if there exist an empty slot

And it is proven that

$$\alpha = 0.5 \pm m \rightarrow \text{prime}$$

then only quadratic probing guarantees that it'll find free slot

Double Hashing

→ two hash function

→ one hash function to find the slot

→ second hash is to find next slot if the calculated by the first is not free.

→ if m is relatively prime, then it always find a free slot if there is one.

→ distribute keys more uniformly than linear probing quadratic

$$\text{th}(\text{key}, i) = (\text{th}_1(\text{key}) + i \cdot \text{th}_2(\text{key})) \% m$$

No clustering

Also h_2 can never return 0, because if it returns zeroes it will repeatedly perform collision.

$$E \rightarrow 6 - (\text{key} \% 6)$$

give value key % 6

$$6 - \{0-5\} \rightarrow \{1-6\}$$

Keys → {49, 63, 56, 52, 54, 48}

0	49
1	
2	54
3	63
4	56
5	52
6	48

→ key = 54

$$h_1(\text{key}) = 54 \% 7 = 5$$

$$h_2(\text{key}) = 6 - (5 \times 6) = 6 - 30 = 6$$

$$\text{hash}(\text{key}) = (5 + 6 \times 6) \% 7 = 3$$

$$\text{hash}(\text{key}, i) = (\text{th}_1(\text{key}) + i \cdot \text{th}_2(\text{key})) \% m$$

→ key = 49

$$\text{hash}(49) = 0$$

→ key = 63

$$\text{hash}(63) = 0$$

$$h_2(63) = 6 - (6 \times 6) = 6 - 36 = 3$$

For the first time put

$$\text{hash}(63) = (0 + 1 \cdot (3)) \% 7 = 3$$

$$= 3$$

$$(5 + i(6)) \% 7$$

first time $i = 1$

$$(5 + 6) \% 7 \rightarrow 4 \text{ (filled)}$$

second $i = 2$

$$(5 + 12) \% 7$$

$$17 \% 7 = 3 \rightarrow \text{filled}$$

third $i = 3$

$$(5 + 18) \% 7 = 23 \% 7 \rightarrow 2 \text{ (found)}$$

• initially we put $i = 0$, if and calculate way h_1 only.

• h_2 is used only when collision happens initially put $i = 1$ then $i = 2, 3, \dots$

• we increment i each time we see collision and initially we start with $i = 1$ for first collision.

Requirements

• $h_2(\text{key}) - m$ should be relatively prime.

Algorithm for open Addressing

```
void doubleHashing (key) {
    if (table is full)
        return error;
```

```
probe = h1(key), offset = h2(key);
while (table[probe] is occupied)
    probe = (probe + offset) % m;
```

```
table[probe] = key;
```

};

For linear probing offset = 1.

Performance Analysis of search

$\alpha = n/m$ (should be ≤ 1) i.e. no of slots should be greater than keys

let $\alpha = 0.8$

it means 80% of table is occupied &

20% is still empty. (1/5)

To find an empty slot you need 5 iterations i.e. 4 filled slots & then 1 empty slot.

$\alpha = 0.9$

Average no of probes: req $\rightarrow 9$

$$\approx \left[\frac{1}{1-\alpha} \right]$$

Implementation of open Addressing

Assumption:- -1, -2 are not present as keys.

+ is for ~~deleted~~ empty
-2 is for deleted.

Structure in C++

```
struct myHash {
```

```
    int *arr;
```

```
    int cap, size;
```

```
    myHash (int c) {
```

```
        cap = c;
```

```
        size = 0; arr = new int [cap];
```

```
        for (int i=0; i<cap; i++) {
```

```
            arr[i] = -1;
```

```
        } int hash (int key) {
```

```
            return key % cap;
```

```
        };
```

```
        bool search (int key) {
```

```
            bool insert (int key) {
```

```
                bool erase (int key) {
```

```
        };
```

bool search (int key) {

```
    int h = hash (key);
```

```
    int i = h;
```

```
    while (arr[i] != -1) {
```

```
        if (arr[i] == key) {
```

```
            return true;
```

```
            i = (i+1) % cap;
```

```
            if (i == h)
```

```
                return false;
```

```
        }
```

```
    return false;
```

```
};
```

} hash function

we return false when

① we are an empty slot i.e. arr[i] = -1

② we have reached the whole arr (i = h)

we return true when

we find the key
simple if

bool insert (int key) {

if (size == cap)

return false;

int i = trash (key);

while (arr[i] != -1 && arr[i] != -2 && arr[i] != key)

i = (i+1) % cap

}

if (arr[i] == key)

return false;

else {

arr[i] = key;

size++;

return true;

}

- * we returned false
- (1) trash table is full.
- (2) when key is already present
- * returned true
- + found spot for the key & size++

bool erase (int key) {

if (size == 0)

return false;

int i = trash (key);

while (arr[i] != -1) { → search for this key.

if (arr[i] == key) {

arr[i] = -2; size--;

return true;

}

i = (i+1) % cap;

if (i == -1) {

return false;

}

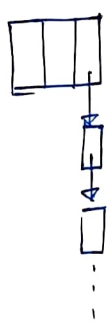
return false;

Found the key

Comparison between two Techniques

Chaining (Better)

Hash table would never fill.



May be the size of chain becomes huge but the we can always add another node on the cost of performance

- ① less sensitive to trash functions
- ② the cache friendly as it is very poor.

③ Extra space is required for links.

④ Performance → $O(1 + \alpha)$

$1 + 0.9 \rightarrow 1.9$

Open Addressing

① Table may become full and requires

resizing after a point

There may be a case when keys are added, for that we need to resize the table if it is full

② Extra case is required for clustering.

③ Cache Friendly

④ No extra space required

⑤ Performance $1/\alpha$

$\frac{1}{1-0.9} \approx \frac{1}{0.1} = 10$

How to handle the case when -1 & -2 are input key? As what libraries do, they store references or pointers. we don't use actual keys we use their pointers as for.

① empty - we can simply use NULL

② deleted - we have a dummy node, whenever we delete something we store pointer or reference to this dummy node.

* This dummy node is not part of class & is shared by all the members

can be used when count of keys is

unordered_set in C++

- It internally was having for storing elements
- No particular order is maintained.
- Output can be any perm. permutation of elements

`unordered_set < int > s;`

① `s.insert(10)` → to insert an element

② `s.find(15)` → check if the element is present
`if (s.find(15) == s.end())`
 If element is present iterator to that element is returned.
 If element is not present iterator to element just after last element is returned (`s.end()`)

③ `s.begin()` → iterator to first element

④ `s.size()` → gives the size of unordered set

⑤ `s.clear()` → it clears the unordered set completely

⑥ `s.count()` → alternative of find function and returns either 1 or 0

when element is present → 1
 when element is not present → 0

⑦ `s.erase()` → `s.erase(15)` → find & remove 15

- we can pass iterator also.

`it = s.begin() + 2;`

`s.erase(it);`

- we can also use `erase` to remove set of elements
- `s.erase(s.begin(), s.end());`

unordered_map in C++ STL

- Based on having
- unordered container with fast insert, delete & search operations
- It stores key-value pair

`unordered_map < string, int > m;`

data type 1

data type 2

`m["ggg"] = 20`

If key is present i.e. if key value pair having "ggg" as key is present, it returns the sequence of the value corresponding to the key and if key is not present, it simply inserts the value and returns the sequence to the key.

`m["ggg"] = 20`

→ inserting key & returning sequence to the value.
 (initialised default as 0)
 with that value as 20

⑧ `m.insert({"courses", 15});`

standard function for inserting a key-value pair

⑨ Traversing

`for (auto x : m)`

`cout << x.first << " " << x.second << endl;`

accessing key

accessing value

③ m.find("ide") → if m.find("ide") = m.end() it means key is not present

we can also get the corresponding value

* auto it = m.find("ide");

if (it != m.end()) {

cout << it->second;

} it is a pointer

③ m.begin() → iterator pointing to 1st element

③ m.end() → iterator pointing beyond last element

③ m.size() → number of key, value pairs

③ m.erase("ide") → remove key value pair

Time comp → $O(1)$

Count Distinct Elements

I/P → { 10, 20, 30, 20, 10 }

O/P → 3

I/P → { 10, 20, 30 }

O/P → 3

Naive Approach

int countDistinct (int arr[], int n) {

int i = 0;

for (int i = 0; i < n; i++) {

bool flag = false;

for (int j = 0; j < i; j++) {

if (arr[i] == arr[j]) {

flag = true;

break;

}

if (flag == false) res++;

return res;

}

Time complexity = $O(n^2)$

Space comp = $O(1)$

Efficient solution

we use an unordered set in STL which does hashing and insert all the elements in it.

As we know in unordered set duplicates are not allowed. Hence we can simply check the size of set for distinct elements.

int countDistinct (int arr[], int n) {

unordered_set <int> s;

for (int i = 0; i < n; i++)

s.insert (arr[i]);

return s.size();

}

Time complexity → $O(n)$

Space comp → $O(n)$

Frequencies of Array Element

I/P → { 10, 12, 10, 15, 10, 20, 12, 12 }

O/P → 10-3

12-3

15-1

20-1

* Naive approach → for every element traverse to the right of it and calculate frequency.

You must also check if that element is appeared at left, if it appeared at left, it must have processed, that's

// at any point if
// i, we check from
// 0 to i-1 if element
// is present or
// not.

we can directly pass integer array as unordered_set <int> s(arr, arr+n);

Efficient solution

using an unordered map.

```
int countFreq (int arr[], int n) {  
    unordered_map<int> h;  
    for (int x: arr) {  
        h[x]++;  
    }  
    for (auto e: h) {  
        cout << e.first << " " << e.second;  
    }  
}
```

Time comp $\rightarrow O(n)$
Aux space $\rightarrow O(n)$

Intersection of Two Arrays

IFP $a[] \rightarrow \{10, 15, 20, 15, 30, 30, 5\}$
 $b[] \rightarrow \{30, 5, 30, 80\}$

OP $\rightarrow \{30, 5\}$

Naive solution

```
int intersection (int a[], int b[], int m, int n) {
```

```
    int res = 0;
```

```
    for (int i=0; i<m; i++) {
```

```
        bool flag = false;
```

```
        for (int j=0; j<n; j++) {
```

```
            if (a[i] == b[j]) { flag = true;
```

```
                break;
```

```
        }
```

```
        if (flag == true) { continue; }
```

```
        for (int j=0; j<n; j++)
```

```
            if (a[i] == b[j]) { res++; break; }
```

```
    }
```

```
    return res;
```

```
}
```

Time comp $\rightarrow O(m \times (m+n))$.

Efficient solution $O(m+n)$.

idea is to use unordered-set.

IFP $\rightarrow a[] = \{10, 15, 20, 15, 30, 35, 5\}$
 $b[] = \{30, 5, 30, 80\}$

```
int intersection (int a[], int b[], int m, int n) {
```

```
    unordered_set<int> s;
```

```
    for (int i=0; i<m; i++)
```

```
        s.insert (a[i]);
```

```
    for (int j=0; j<n; j++) {
```

```
        if (s.find (b[j]) != s.end()) {
```

```
            res++;
```

```
        s.erase (b[j]);
```

```
    }
```

```
    }
```

```
    return res;
```

```
}
```

To ensure that
duplicates are not allowed.

Union of two arrays

IFP $\rightarrow \{15, 20, 5, 15\}$, $\{15, 15, 20, 10\}$

OP $\rightarrow 4$ $[15, 20, 5, 10]$

Naive solution

* Make an auxiliary array $c[m+n]$

* Copy contents of A and B into auxiliary array

* And then count distinct elements in that array $c[m+n]$

Time comp $\rightarrow O((m+n) * (m+n))$

Aux space $\rightarrow O(m+n)$.

Efficient solution

we can insert all the elements in an unordered set and can simply return its size

int findUnsorted (int a[], int b[], int m, int n) {

unordered_set <int> s;

for (int i = 0; i < m; i++) {

s.insert(a[i]);

for (int i = 0; i < n; i++) {

s.insert(b[i]);

return s.size();

}

Pair with a given sum in unsorted array

I/P → {3, 8, 5} - 8

sum → 17

O/P → Yes

Naive Approach

for (i = 0; i < n; i++)

for (j = i + 1; j < n; j++)

if (arr[i] + arr[j] == sum)

return true; }

}

return false;

Efficient solution

By using set

At any point before inserting an element into the set look for (sum - arr[i]) inside the set

int findUnsorted (int a[], int m, int n) {

unordered_set <int> s;

for (int i = 0; i < m; i++) {

if (s.find(sum - arr[i]) != s.end())

return true;

else {

s.insert(arr[i]);

}

}

return false;

}

Time complexity → O(m)

Aux space → O(m)

Subarray with zero sum

contiguous elements

I/P → {1, 4, 13, -3, -10, 5}

O/P → Yes

I/P → {4, -3, 2, 1}

O/P → Yes

Naive solution

we start with current element as first element and then run loop for remaining elements to check if sum gets 0.

bool isSubarray (int arr[], int n) {

for (int i = 0; i < n; i++) {

int curr = arr[i]; int j = i + 1;

while (curr != 0 && j < n) {

curr = curr + arr[j];

j++;

if (curr == 0)

return true;

}

return false; }

Time comp - O(n²)

Effective solution $O(n)$

Prefix sum $\rightarrow$

$$\text{sum} = 0$$

$$a_1 a_2 a_3 a_4 a_5 a_6 \dots a_{i-1} a_i a_{i+1} a_{i+2} \dots a_{n-2} a_{n-1}$$

Prefix sum (A) $\uparrow$ Prefix sum (B) $\uparrow$

$$\text{Prefix sum (A)} = \text{Prefix sum (B)}$$

We will calculate prefix sum and if there is a ~~term~~ subarray having 0 sum then the prefix sum of two elements will be equal (To check this we use looping)

We will return true in two cases.

- ① When the prefix sum we have calculated is already present in ~~the~~ set
- ② When the calculated prefix sum comes to be 0 i.e. first i elements sum up to 0.

Summary with given sum

- This is the extension of previous problem
- Instead of comparing prefix sum $= 0$ we compare prefix-sum $= \text{sum}$
- And we search for (prefix-sum - sum) in the unordered set

```
bool isSubarray (int arr[], int n, int sum) {
    unordered_set<int> s;
    int pre-sum = 0;
    for (int i = 0; i < n; i++) {
        pre-sum += arr[i];
        if (s.find(pre-sum) != s.end())
            return true;
        if (pre-sum == 0)
            return true;
        s.insert(pre-sum);
    }
    return false;
}
```

Longest subarray with given sum

I/P : arr[] = {5, 8, -4, 1, 9, -2, 2} sum = 0

O/P : 3. (8, -4, -4)

I/P : arr[] = {3, 1, 0, 1, 8, 2, 3, 6} sum = 5.

O/P : 5

I/P : arr[] = {8, 3, 7} sum = 15

O/P : 0 (No subarray is present).

Naive solution $O(n^2)$

- ① We will find the length of subarray using (j-i+1)

```
int maxLen (int arr[], int n, int sum) {
    int res = 0;
    for (int i = 0; i < n; i++) {
        int curr-sum = 0;
        for (int j = i; j < n; j++) {
            curr-sum += arr[j];
            if (curr-sum == sum)
                res = max(res, j-i+1);
        }
    }
    return res;
}
```

Time comp $\rightarrow O(n^2)$

Effective solution

We will use unordered-map for this problem. We will store prefix sum as well as index of first occurrence of that prefix sum.

arr[] $\rightarrow$ {8, 3, 1, 5, -6, 6, 2, 2}

pre-sum $\rightarrow$ {8, 11, 12, 17, 11, 17, 19, 21}

We will store occurrence because we want to find the longest subarray.

$$\text{length} = j - i + 1$$

$$0 \ 1 \ 2 \ 3 \ 4$$

$$\underline{1}$$

$$2 - 0 + 1 = 3$$


```
int maxSub (int arr [], int n, int m)
```

```
unordered_map<int, int> m;
int pre_sum = 0, res = 0;
for (int i = 0; i < n; i++) {
    pre_sum += arr[i];
    if (pre_sum == sum) {
        res = i + 1;
    }
    if (m.find(pre_sum) == m.end()) {
        m.insert({pre_sum, i});
    }
    if (m.find(pre_sum - sum) != m.end()) {
        res = max(res, i - m[pre_sum - sum]);
    }
}
return res;
```

Longest subarray with equal 0's and 1's.

I/P arr → {1, 0, 1, 1, 0, 0, 1}

O/P → {1, 1, 1, 1}

I/P → {0, 0, 1, 1, 1, 1, 0, 1}

O/P → {0, 0, 1, 1, 1, 1}

Naive solution

We maintain two pointers and calculate the number of 0's and 1's in the range.

arr[] → {1, 0, 1, 1, 1, 0, 0, 1}

at certain point

0 → 2

1 → 3

```
int longestSum (bool arr [], int n) {
```

```
int res = 0;
for (int i = 0; i < n; i++) {
    int c0 = 0, c1 = 0;
    for (int j = i; j < n; j++) {
        if (arr[j] == 0) {
            c0++;
        }
        else {
            c1++;
        }
        if (c0 == c1) {
            res = max(res, j - i + 1);
        }
    }
}
return res;
```

O(n²) quadratic approach

Effective solution

Trick is to replace every 0 with -1 and calculate prefix sum and then find longest subarray with 0 sum.

arr[] → {1, 0, 1, 1, 1, 0, 0, 1}

int findMaxSubarray (int arr [], int n)

```
for (int i = 0; i < n; i++) {
    int pref_sum = 0;
    unordered_map<int, int> m;
    m.insert({0, -1});
    for (int j = i; j < n; j++) {
        pref_sum += arr[j];
        if (m.find(pref_sum) != m.end()) {
            int len = j - m[pref_sum];
            res = max(res, len);
        }
        m.insert({pref_sum, j});
    }
}
return res;
```

Longest common span with same sum in

Binary Array

arr1[] $\rightarrow$ {0,1,0,0,0,0}

arr2[] $\rightarrow$ {1,0,1,0,0,1}

O/P $\rightarrow$ 4

arr1[] $\rightarrow$ {0,1,0,1,1,1}

arr2[] $\rightarrow$ {1,1,1,1,0,1}

O/P $\rightarrow$ 5

arr1[] $\rightarrow$ {0,0,0,0,1}

arr2[] $\rightarrow$ {1,1,1,1}

O/P $\rightarrow$ 0.

starting and ending indices should be same in subarray & sum of elements of subarray must be equal.

Naive Approach

Run two pointers i and j on both arrays. and then check sum of both. If sum is equal then we get common subarray, to get maximum, we will store the result of each time and then return max of it.

in maxCommon (bool arr1[], bool arr2[], int n)?

int res = 0

for (int i = 0; i < n; i++) {

int sum1 = 0, sum2 = 0;

for (int j = i; j < n; j++) {

sum1 += arr1[j];

sum2 += arr2[j];

if (sum1 == sum2)

res = max(res, j - i + 1)

return res;

Effective solution

- We subtract one array (index to index) from the another and then calculate the subarray with zero sum

there will be 4 cases

- ① first array contains 2 second also $\rightarrow$ 0 {same subarray}
- ② 1st array contains 0 2 second contains 0 $\rightarrow$ 0 {same subarray}
- ③ 1st array contains 1 2 2nd contains 0 $\rightarrow$ 1 {adding to sum 2}
- ④ 1st array contains 0 2 2nd contains 1 $\rightarrow$ -1 {adding to sum 2}

arr1[] $\rightarrow$ {0,1,0,0,0,0}

arr2[] $\rightarrow$ {1,0,1,0,0,1}

output $\rightarrow$ {1,1,1,0,0,-1}

Req. Ans.

O(n)

Longest consecutive subsequence

I/P - arr1[] $\rightarrow$ {1,9,3,4,2,20}

I/P $\rightarrow$ {20,30,40}

O/P $\rightarrow$ 4

O/P $\rightarrow$ 2

I/P - arr1[] $\rightarrow$ {8,20,7,30}

O/P $\rightarrow$ 2

Method-1

Sort the array and update update curr-res when you see a consecutive value.

1,9,3,4,2,20 $\rightarrow$ 1,2,3,4,20
4 1 1 max $\rightarrow$ 4

Time comp $\rightarrow$ O(nlogn)

Method-2

- create a hash table of size n and insert all elements into the hash table. and then traverse the array, for every element check if it is the starting point of sequence — if yes $\rightarrow$ look for elements after it — if no $\rightarrow$ ignore.

Algorithm

- Insert all elements of array into the
- for (int i = 0; i < n; i++)
- if (arr[i] is not present in arr)

curr = 1

while (arr contains arr[i] + curr)

curr++;

res = max(res, curr);

return res;

Count Distinct Elements in every window of size k

I/P → arr[] = {10, 20, 20, 10, 30, 40, 10}

Effective approach → use map

- create a frequency map of first k elements
- and then get the size of it
- on moving k further
- decrement the frequency of freq[arr[i]]
- if freq become 0 remove element

~~vector<int> countDistinct (int arr[], int n, int k)~~

unordered_map<int, int> freq;

for (int i = 0; i < k; i++) {

freq[arr[i]]++;

}

vector<int> countDistinct (int arr[], int n, int k) {

unordered_map<int, int> freq; vector<int> res;

for (int i = 0; i < k; i++)

freq[arr[i]]++;

for (int i = k; i < n; i++) {

freq[arr[i-k]]--;

if (freq[arr[i-k]] == 0)

freq.erase(arr[i-k]);

freq[arr[i]]++;

res.push_back(freq.size());

}

}

return res;

More than (n/k) occurrences

I/P → arr[] = {30, 10, 20, 20, 10, 20, 30, 30}

k = 4

O/P = {20, 30}

Method-1

Sort the array and then check the count of every element if n/k < count we print that element

Time complexity → O(n log n)

Method-2

Build and frequency map and then return those element whose frequency > n/k.

Time complexity → O(n)

Efficient solution

- we will use the concept of moore's voting algorithm.

fact is $\text{res_count} \leq k-1$

let's suppose every element occurs n/k times

$$k * (n/k + 1) \leq n$$

~~$k * n + k \leq n$~~ (contradicting)

* Algorithm

- create an empty map m
- for ($i=0$; $i < n$; $i++$)
- (a) if (m contains $\text{arr}[i]$)

$m[\text{arr}[i]]++$;

- (b) if ($m.size()$ is less than $(k-1)$)
- $m.put(\text{arr}[i], 1)$.

- (c) else decrease all values in m by one and if value becomes 0 remove it.

For all elements in m , print the elements that actually appears more than n/k times.

$O(nk)$

STRING

- sequence of characters
- small set
- contiguous integer value 'a' to 'z' and 'A' to 'Z' in both ASCII and UTF-16

C/C++

Java:

char: ASCII
8 Bit

UTF-16
16 Bit

Also supports unicode

char x = 'a';
cout << (int)x; // 97

Print the frequencies of character in sorted order

string str = "geeksforgeeks";
int count[26] = {0};

for (int i=0; i < str.length(); i++)

 count[str[i] - 'a']++;

for (int i=0; i < 26; i++) {

 if (count[i] > 0) {

 cout << (char)(i + 'a') << " ";

 cout << count[i];

 }

o/p
e f g k o r s
4 1 1 2
1 1 1 2

Strings in C++ (char arrays)

char str[] = "gfg"

example

g f g \0

if we give definite size

str[6] = "gfg"

g f g \0 \0 \0

strlen
• strlen → gives length of string