

Algoritmen: procedurele kennis (hoe dingen uit te voeren)

Convex hull = convex omhullende figuur

↳ kan men maken buiten de figuur

↳ algoritme: binnenhoek $< 180^\circ \rightarrow ok$

buitenhoek $< 180^\circ \rightarrow$ probleem - punt overlaan, naar volgend punt

zo met binnentert punten verbinden, dan onduidelijk zelfde herhalen

door alle punten op geschikte manier te overlappen (2x)

↳ noemen in convex hull zijn "equivalent" m.b.t. hun complexiteit

meest efficiënte algoritme $\leftrightarrow$ meest efficiënte convex hull algoritme

$$T_{\text{convex}}(n) = T_{\text{sort}}(n) + c \cdot n + c \cdot n$$

binnentert + onduidelijk doorlopen

$$T_{\text{sort}}(n) = T_{\text{convex}}(n) + c \cdot n$$

onduidelijk convex omhullende uitbreiden $\rightarrow$ geschikte lijst



Analyseren $\rightarrow$ prestatie, vergelijken met andere, benchmark tgd, theorie basis

N-lichamenprobleem

- slecht algoritme: brute force

2 for loops: elk lichaam met elk ander vergelijken

$\rightarrow N(N-1)$ berekeningen $\Rightarrow \sim n^2$

- beter: Barnes-Hut:

clusters bekijken met hun zwaartekracht

$\rightarrow N \log_2 N$

Wetenschappelijke methode voor analyseren

- observeren [natural world]

- hypothesen opstellen [model consistent met observaties]

- voorspellen [gemiddelde hypothesen]

- verifiëren [of nieuwe observ.]

- valideren [herhalen tot het klopt]

↳ experimenten reproduceerbaar

& hypothesen falsifieerbaar

gemeten uitvoertijd experimenten [stopwatch] plotten i/v de

inputgrootte

log-log plot

$T(N)$ i/v N

$\rightarrow \log_2(T(N))$ i/v $\log_2(N)$

stel $T = a \cdot N^b$

$\rightarrow \log_2(T) = b \cdot \log_2(N) + c$, met $a = 2^c$

rechte lijnen fitten en parameters verifiëren

verduubelingsexperiment

[weert inder by power laws]

geeft een macht b

een experimenten met telkens een verdubbelende grootte N

we kunnen nu b verifiëren

aangaven voor hypothesen

in dan in de verhouding

$$T = a \cdot N^b$$

$$\rightarrow T(2N) = a(2N)^b$$

$$\frac{T(2N)}{T(N)} = \frac{a(2N)^b}{a(N)^b} = 2^b$$

$$\Rightarrow b = \log_2(\text{ratio } \frac{T(2N)}{T(N)})$$

eenmaal je zo b hebt gevonden, kan je a uitrekenen

dan voor eenzelfde N het experiment uitvoeren te kunnen

en dan oplossen voor a : $a = \frac{T}{N^b}$

met N bekend, T gemeten & b bekend

$\Rightarrow$ hypothesen $T = a \cdot N^b$ met

www.lmsintl.com

bepaalde a en b

by $T = a N^b$ is

- a afhankelijk v/h spec. systeem, v.b. snelheid computer, OS, software, geheugen
→ systeem afh. effecten
- b onafhankelijk v/h spec. systeem
→ systeem onafh. effecten → v.b. algoritme, input data

mathematische modellen

(kost van elke operatie) $\times$ (frequency v/d operatie) $\Rightarrow$ totale tijd

↳ tegenoverstell. manieren

- simplificaties
- 1) belijven enkel de bewerkingen die herachte werk doen niet de boekhouding e.d.
↳ proxy voor heel het algoritme
→ tellen enkel de body, werk i/h binnenste v/d lus
 - 2) belijven enkel de meest stijgende term, in termen van N, want we zijn enkel geïnteresseerd in gedrag bij zeer grote N
↳ orde van groei

$\Rightarrow$ tilde notatie

$$f(x) \sim g(x) : \lim_{x \rightarrow \infty} \frac{f(x)}{g(x)} = 1$$

→ ongeveer gelijk aan → leading term
verwachtigheid

andere notaties

- grote O : $f(x)$ is $O(g(x))$: $\exists c, x_0 : x \geq x_0 : f(x) \leq c g(x)$ ↳ bovengrens
- grote Θ → asymptotische grootheid : f is $\Theta(g)$: f is $O(g)$ en g is $O(f)$
- Ω : Θ en kleiner → ondergrens
- ω : Θ en groter → ondergrens

by tilde: vergelijk de meest stijgende term ook meenemen
kan heel bel! zijn in efficiënte algoritmes met welke grote O
vergelijken, want $\sim \frac{1}{10} N^2$ zal veel efficiënter zijn dan $\sim 50 N^3$

waars in dergelijke analyse

- grote de coeff. lagere orde termen
- niet-dominante binnenste lus
- instruction time niet voor elke instructie gelijk
- systeem pc
- sterke afh. input
- versch. probleemformaat, niet enkel N

waars van - best case - average case vergelijken

amortized analysis: $\frac{\text{totale kost alle operaties}}{\# \text{ operaties}}$

→ gem. kost laag proberen te houden

geheugengebruik ook tell!

geheugen object = som geheugen alle instance var. + geheugengebr. overhead object

int[] ~ 4n bytes, double[] ~ 8n, double[][] ~ 8mn

ref. objectwaars
+ variabelen info
+ objecten omvattende info

managing arrays of elements
based on the help of the items

SORTING: sorteren van data

sorteralgoritme kiezen op basis van data

- hoe makkelijk zijn de objecten te verplaatsen / kopiëren (waarden)
- " " " " is het om objecten te vergelijken (wijzende)

⇒ 2 soorten sorteralgoritmen: - vergelijkingen niet veel gedaan
- manipuleren (verplaatsen / kopiëren) niet veel gedaan

in-place
sorting
+ invariabele
gebruik

Selection sort

minimum in de rij zoeken, dit selecteren en vervangen
dan in resterende rij hetzelfde doen

dit blijven herhalen totdat heel de rij gesorteerd is (gunstige rij men even)

public class SelectionSort {

public static void sort (Comparable[] a)

{ int n = a.length;

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

int min = i;

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

if (less(a[j], a[min]))

min = j;

exchange(a, i, min); }

}

}

2 for lussen: 1 om door de lijst te lopen totdat alles gesorteerd is
1 om door de lijst te lopen om telkens het min te zoeken

steeds kleiner wordend

alle elementen links van i zijn steeds / nu gesorteerd en in dus niet opnieuw bekijken
inclusief i

alle rechts van i hoe meer kleiner zijn dan de elementen links van i
invarianten van algoritme, altijd waar na de lus 1 heen te draaien

[symm algoritme: max zoeken en van achter zetten, van achter ivv voor werken]

private static boolean less (Comparable v, Comparable w) {

return v.compareTo(w) < 0; }

private static void exchange (Comparable[] a, int i, int j) {

Comparable t = a[i];

a[i] = a[j];

a[j] = t; }

Analyse: # vergelijkingen x # verwisselingen tellen
n elementen

op iteratie i: (n-i-1) vergelijkingen, 1 verplaatsing

⇒ # vergelijkingen voor alle iteraties

$$\sum_{i=0}^{n-1} (n-i-1) = (n-1) + (n-2) + \dots + 1 + 0 = \frac{n(n-1)}{2}$$

$$\Rightarrow \sim \frac{n^2}{2}$$

⇒ # verwisselingen

$$1 + 1 + \dots + 1 = n-1 \Rightarrow \sim n$$

Insertion Sort

in-place
sorting
variabel
gedrag

1^e element vld onge sorteerde lijst vergelijken met zijn (al gesorteerde) voorganger
dit herhalen totdat het ongesorteerde element op de juiste plaats in de
reeds gesorteerde lijst kan worden invoegen
herhalen voor alle elementen in de (steeds korter wordende) onge sorteerde lijst
public void InsertionSort {

public static void sort (Comparable[] a) {

int n = a.length;

for (int i = 1; i < n; i++) {

for (int j = i; j > 0 && !less(a[j], a[j-1]); j--)

exchange(a, j, j-1); }

}

element pakken

vergelijken met zijn linker-
buren

indien nodig wisselen
als linker > rechter

↳ 2 for lussen: 1 om de neg te sorteren vj steeds kleiner te maken
1 om het onge sorteerde elementje steeds in de gesorteerde lijst in te voegen

invar: { alle elementen links (inclusief) i zijn gesorteerd
alle elementen rechts van i zijn nog niet beloken

analyse / performance

voor element i: per grote element links van i 1 vergelijking en 1 verwisseling
+ max 1 vergelijking met een kleiner element

best case: alle elementen links van i zijn kleiner
→ rij is al gesorteerd
0 verplaatsingen

$n-1$ vergelijkingen $\sim n$

worst case: alle elementen links van i zijn groter
→ rij is al omgekeerd gesorteerd

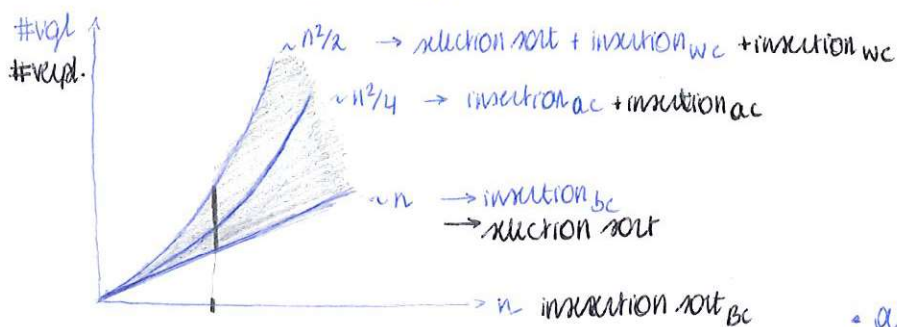
$\sum_{i=0}^{n-1} (n-i-1) = \frac{n(n-1)}{2}$ vergelijkingen & verwisselingen $\sim \frac{n^2}{2}$

average case helft vld elementen links van i zijn groter

→ elk element moet tot het midden vld verplaatst w

⇒ # vergelijkingen & verplaatsingen $\sim \frac{n^2}{4}$

Selection vs Insertion Sort



qua # vgl'en is IS in het
algemeen beter dan SS

qua # verplaatsingen is SS
echter algemeen beter dan IS

↳ verhouding t/m de 2 afwegen

• als het vergelijken makkelijk is, v.s. int
→ selection sort

• als het vergelijken moeilijk is, v.s. strings
→ insertion sort

insertion sort is een heel goed algoritme als je rij bijna gesorteerd is,
v.s. als alle elementen max een vast getal x van hun finale positie staan

dan lineair gedrag (practisch) , m $(\frac{x}{2} + 1)n$ # vgl'en
en $\frac{x}{2}n$ # verpl. } average

bij SS werkt dit niet zo → blijven nog steeds $\sim n^2/2$

bottom-up merge sort

public static void sort (Comparable[] a) {

int N = a.length;

aux = new Comparable[N];

for (int sz = 1; sz < N; sz = sz * 2) {

for (int lo = 0; lo < N - sz; lo = lo + sz) {

merge(a, lo, lo + sz - 1, Math.min(lo + sz + sz - 1, N - 1));

probleemgrootte in 2 delen - elk apart sorteren $\rightarrow 2 \times \sim (N/2)^2/2 = N^2/8 \rightarrow \text{tot} \sim N^2/4$

$\rightarrow$ 2 delen terug samenvoegen (mergen): kleinste van 2 delen verglijken

en kleinste van die 2 in nieuwe rij zetten

$\Rightarrow$ extra geheugen nodig (extra hulpanaly: kopiëren die linker helft naar copy array)

recursieve opwerking + hulpanaly

public class MergeSort {

private static Comparable[] aux;

public static void sort (Comparable[] a) {

aux = new Comparable[a.length];

sort(a, 0, a.length - 1);

aux: array for merges

allocate space just once

top-down merge sort

sort a[lo..hi]

public static void sort (Comparable[] a, int lo, int hi) {

if (hi <= lo) return;

int mid = lo + (hi - lo) / 2;

sort(a, lo, mid);

sort(a, mid + 1, hi);

merge(a, lo, mid, hi);

$\rightarrow$ array of length 1 reached

sort left half

sort right half

merge;

merge a[lo..mid]

with a[mid+1..hi]

public static void merge (Comparable[] a, int lo, int mid, int hi) {

for (int k = lo; k <= hi; k++)

aux[k] = a[k];

copy a[lo..hi] to aux[lo..hi]

int i = lo, j = mid + 1;

$\rightarrow$ aux[j] = a[0]

for (int k = lo; k <= hi; k++) {

merge back to a[lo..hi]

if (i > mid)

a[k] = aux[j++];

if a₁ is fully in a

$\rightarrow$ insert a₂ in a

else if (j > hi)

a[k] = aux[i++];

if a₂ is fully in a

$\rightarrow$ insert a₁ in a

else if (aux[j] < aux[i])

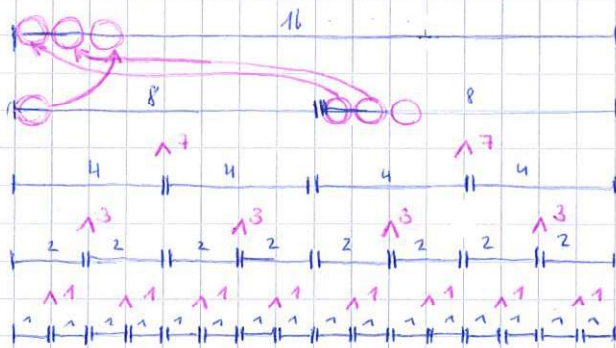
if a₂[0] < a₁[0]

$\rightarrow$ put a₂[0] in a

a[k] = aux[j++];

else a[k] = aux[i++];

if a₁[0] <= a₂[0] $\rightarrow$ put a₁[0] in a



15 verglijken = $1(2 \cdot 8 - 1)$

voor $N = 2^n$
met $n = \#$ niv, hier 4
 $N - 1$

14 verglijken = $2(2 \cdot 4 - 1)$

$2(\frac{N}{2} - 1) = N - 2$

12 verglijken = $4(2 \cdot 2 - 1)$

$4(\frac{N}{4} - 1) = N - 4$

8 verglijken = $8(2 \cdot 1 - 1)$

$8(\frac{N}{8} - 1) = N - 8$

$\sim 2k$ verplaatsingen en $\sim k$ verglijken om 2 arrays van lengte $k/2$ te mergen

$\rightarrow \sim n \log_2 n$ vergelijkingen en verplaatsingen [forum - 2.23]

$\sim n$ extra plaats nodig (geheugen hulpanaly)

worst case voor $N = 2^n$ (met n het # niveaus, N het # elementen)

$$\begin{aligned} N \log_2 N - (1 + 2 + 4 + 8 + \dots) &= N \log_2 N - (2^0 + 2^1 + 2^2 + \dots + 2^{n-1}) \\ &= N \log_2 N - (N-1) \\ &= N \log_2 N - N + 1 \text{ exact} \\ &\sim N \log_2 N \end{aligned}$$

best case $\sim \frac{N}{2} \log_2 N$ (ene deling helemaal geïntegreerd valt groter/kleiner dan andere $\Rightarrow$ maar $N/2$ vgl'en nodig)

stel dat $N \neq 2^n$, dan gebruiken we dat $2^{n-1} < N < 2^n$ en kan we nog steeds onder- en bovengrenzen bepalen

$\hookrightarrow$ worst & average case zijn sneller dan insertion sort (asymptotisch)
best case is insertion sort asymptotisch sneller dan merge sort

verbeteringen merge sort

- overschakelen op insertion sort voor kleine arrays (mergesort heeft te veel overhead voor kleine arrays)
- merge enkel wennr nodig
laatste element linker subarray vergelijken met eerste element rechter subarray
- gebruik steeds dezelfde subarray

Snelheid sorteralgoritmen

- Comparison sorts:

sorteren die 2 elementen te vergelijken

Hinstens $\sim n \log_2 n$ vergelijkingen $\rightarrow$ merge sort is asympt. optimaal

- Non-comparison sorts

$\hookrightarrow$ counting sort, radix sort, bucket sort
in lineair ($\sim n$) tijd

$\sim n \log_2 n$ als ondergrens voor comparison sorts

"B" als beslissingsboom: elke mogelijke input is een blad in d boom

$\Rightarrow n!$ verschillende bladen in de boom

elke knoop heeft 2 vertakkingsopties (binair)

- 2 elementen die vergeleken w

aangenomen elke binaire boom met een hoogte $h \leq 2^h$ bladen heeft

$$\Rightarrow n! \leq \# \text{bladen} \leq 2^h \Rightarrow h \geq \log_2(n!)$$

aangenomen h (de min. hoogte vld beslissingsboom)

het worst case # vgl'en is

kunnen we nu met Stirling's benadering

$$\text{afleiden dat } \log_2(n!) \sim n \log_2 n$$

$\Rightarrow$ geen enkel sorteralgoritme kan worst case beter dan $\sim n \log_2 n$

$\rightarrow$ merge sort asympt. optimaal [maar heeft wel extra geheugen nodig]

bewijs $\log_2(n!) = \log_2 e \log_e(n!)$

$$\ln(n \cdot (n-1) \cdot (n-2) \dots 2 \cdot 1)$$

$$= \ln(n) + \ln(n-1) + \dots + \ln(1) \quad \text{integraalbenadering}$$

$$\approx \int_1^n \ln x \, dx$$

$$= [x \ln x - x]_1^n = n \ln(n) - n + 1$$

$$\rightarrow \log_2(n!) \approx \log_2 e (n \ln(n) - n + 1)$$

$$= n \log_2(n) - 1.443n + 1.443$$

andere benadering
zou verwilt voor
extra lagere orde
elementen zouge

echt ~ compleet
sort

knop ~ comparison
betw art/kat

echte
ondergrens
is $\lceil \log_2(n!) \rceil$

partitioning

- doen is niet mogelijk nu onderling verder gesorterd is
doen we door bovenste strategie te herhalen, namelijk verder sorteren
de pivot blijft ALTIJD op dezelfde plaats staan
[positionering is echt werk]

```
public static void sort (Comparable[] a) {
```

```
private static void sort (Comparable[] a, int lo, int hi) {
```

```
int j = partition(a, lo, hi); // partitionierung, a[j] = pivot element
sort(a, lo, j-1);           // sort. linke Teil: a[lo .. j-1]
sort(a, j+1, hi);           // sort. rechte Teil: a[j+1 .. hi]
```

- $a[i_0 \dots j-1]$ is nooit groter dan $a[j]$
- $a[j+1 \dots hi]$ is nooit kleiner dan $a[j]$
- $a[j]$ is in zijn juiste plaats in de rij

1. Simple partitioning

inbreuk 2 lijsten: $a[i] < a[j] \rightarrow \text{left}$, $a[i] > a[j] \rightarrow \text{right}$, $a[i] = a[j] \rightarrow \text{vind 2 kiezen}$
 $\rightarrow \text{concatenatie } [\text{left}] + [\text{pivot}] + [\text{right}]$

Leiden: extra gekregen modifier voor 2 nieuwe arrays

+ Extra groentes voor kopieren & concareneren (lin.)

i = next element rechts
 j = kleinste element
 k = pivot

2 mensen: i naar rechts doen lopen

alle elementen kleiner dan pivot op pos i

bleven staan & i loopt verder tot dat we

op pos i een element quotes hebben - Stop

alle elementen gelijk dan pivot op pos; blijven staan x; loopt verder tot kleine-stop

→ elementen $a[i]$ en $a[j]$ van matrix wisselen en opnieuw verder doen

→ STOP wenn i en j "beiden" → Wahlte kleinere Element werden mehr pivot

- ↳ for men on steroids with active sex life

voor

v

↳ hi

tijden

v	$\leq v$	i	j	$\geq v$
---	----------	---	---	----------

na

$\leq v$	v	$\geq v$
----------	---	----------

lo

j

hi

2. Hoare

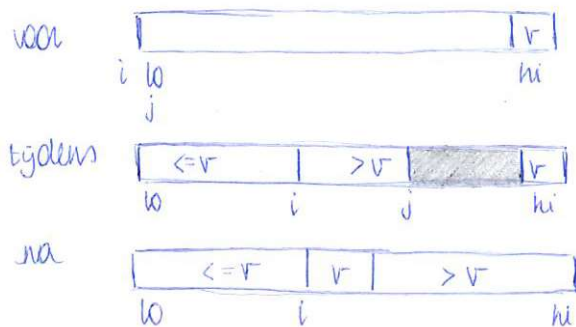
```
private static int partition (Comparable[] a, int lo, int hi) {
    int i = lo, j = hi + 1; // left and right scan indices
    Comparable v = a[lo] → pivot, partitioning value - 1st element
    while (true) {
        while (less (a[++i], v)) { // scan right
            if (i == hi) break;
        }
        while (less (v, a[--j])) { // scan left
            if (j == lo) break;
        }
        if (i >= j) break; // check for scan complete
        exchange (a, i, j); // exchange
    }
    exchange (a, lo, j); // put pivot v in position v = a[j]
    return j; // j is position of pivot, with a[lo..j-1] < a[j] < a[j+1..hi].
}
```

$n+1$ vergelijkingen

efficiënt
weinig
verplaatsingen

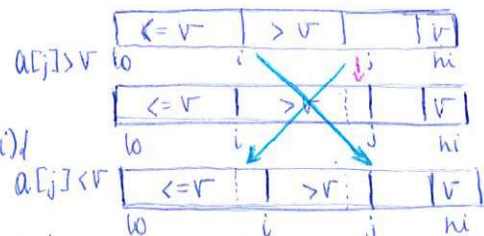
3. Partitionering van Lomuto

het laatste element als pivot
index j = 1st element dat we
moeien behouden
index i = laatste element groep
kleiner dan pivot
2 opties: 1) als $a[j] > v$: element
mag blijven staan en
 j schuift 1 naar rechts



2) als $a[j] < v$: element moet op positie $i+1$ komen te staan → $a[j]$ exch. $a[i+1]$
en $i = i+1$ en $j = j+1$

herhalen totdat $a[j] = v$
⇒ pivot met $a[i+1]$ wisselen



```
private static int partition (Comp[] a, int lo, int hi) {
```

```
    int i = lo - 1; // scan index
```

```
    Comparable v = a[hi]; // pivot, partitioning value - laatste element
```

```
    for (int j = lo; j <= hi - 1; j++) { // scan alle elementen
```

```
        if (less (a[j], v)) { // scan equal elementen
```

```
            i++;
```

```
        exchange (a, i, j); }
```

```
    exchange (a, i+1, hi); // pivot op pos i+1 zetten, v = a[i+1]
```

```
    return i+1; // positie v retourneren, met a[lo..i] < a[i+1] < a[i+2..hi]
```

$n+1$ vergelijkingen

2 elementen worden nooit meer
dan 1 keer met elkaar vergeleken

eligaar

Performantie quicksort

best case: pivot steeds mooi in het midden, gebalanceerd opgesplitst
stel dat de partitie n vergelijkingen uitvoert

$$T(n) = 2T\left(\frac{n}{2}\right) + \text{partitioning}(n)$$

$$= 2T\left(\frac{n}{2}\right) + n \pm 1$$

$$= n \pm 1 + 2\left(\frac{n}{2} \pm 1 + 2T\left(\frac{n}{4}\right)\right)$$

$$= n \pm 1 + n \pm 2 + 4T\left(\frac{n}{4}\right)$$

$$= n \pm 1 + n \pm 2 + n \pm 4 + 8T\left(\frac{n}{8}\right)$$

$$= \dots$$

$$\sim n \log_2(n)$$

→ niet goed als merge sort, maar zonder extra geheugen

tegel om n elementen te sorteren
= tegel partitionering

+ tegel linkse sorteren

+ tegel rechte sorteren

veronderstelling

linkse en rechte sorteren gelijk

$$\left[\log_2 \frac{n}{2} \text{ en } \log_2 \frac{n}{2}\right]$$

+ → Hoare
- → Lomuto

Worst case: 0 elementen in de ene subarray en $n-1$ in de andere
 ↳ pivot steeds kleinste (of grootste) element - is al geordend
 → zelfde gedrag als selection sort

$$T(n) = T(n-1) + T(0) + \text{partitionering}(n) = T(n-1) + \text{partitionering}(n) = \dots$$

$$\sim n^2/2$$

↳ kan dat we in worst case zitten:

2 elementen op n
 2 elementen op $n-1$
 ...
 2 elementen op 2

$$\frac{2 \cdot 2 \cdot 2 \cdot \dots \cdot 2}{n!} = \frac{2^{n-1}}{n!} = P(\text{Worst case})$$

→ $n!$ stijgt sneller dan 2^{n-1}
 v.b. $n=10$, $P=0,014$

(klopt het moet je niet breken, want dan mag maar 1 elementje over)

Average case: behouding de split is in proporties van n

$$\sim c n \log_2 n$$

$$v.b. \text{ split @ } 9/10 \text{ geeft } T(n) = T(9n/10) + T(n/10) + \text{partitionering}(n) \sim c n \log_2 n \quad \text{hier } c = \log_{10/9} 2$$

n comparsen op level 2, $n/10 + 9n/10 = n$ comparsen op level 2

$n/100 + 9/100n + 9/100n + 81/100n = n$ comparsen level 3 ...

Quicksort kan terugkomen van 1 slechte split - 1 extra kost $\sim n$ (taget oude term)

⇒ $n \log_2 n$ gedrag behouden $\sim n \log_2 n + (\# \text{ slechte stappen}) \cdot n$

"exacte" wiskundige benadering

aanpak 1: obv. is vgl. nodig voor elke partitionering

aannames: elk element evenveel kans om pivot te w. ($P = 1/n$)

x partitionering voor n elementen kost $n+1$ vgl. (Hoare)

vgl'en nodig om de n elementen te sorteren:

$$C(n) = (n+1) + \frac{1}{n} [C(0) + C(n-1)] + \frac{1}{n} [C(1) + C(n-2)] + \dots + \frac{1}{n} [C(n-1) + C(0)]$$

partitionering

$\frac{1}{n}$ kans dat we del. vgl. links en $n-1$ vgl. rechts partitie hebben

$\frac{1}{n}$ kans voor 1 vgl. links en $n-2$ vgl. rechts

$\frac{1}{n}$ kans voor elke mogelijke uitkomst vgl. uitkomsten vgl. partitie

sym. kost

$$\Leftrightarrow C(n) = (n+1) + \frac{2}{n} [C(0) + C(1) + C(2) + \dots + C(n-1)]$$

$$\Leftrightarrow n C(n) = n(n+1) + 2 [C(0) + C(1) + \dots + C(n-1)]$$

$$\text{voor } n-1: (n-1) C(n-1) = (n-1)n + 2 [C(0) + C(1) + \dots + C(n-2)]$$

$$\Rightarrow n C(n) - (n-1) C(n-1) = n(n+1) - (n-1)n + 2 C(n-1) = 2n + 2 C(n-1)$$

$$\Leftrightarrow n C(n) = 2n + (n+1) C(n-1)$$

$$\Leftrightarrow \frac{C(n)}{n+1} = \frac{2}{n+1} + \frac{C(n-1)}{n}$$

→ heel symm. recurrente uitdrukking

$$\Rightarrow \text{expansie: } \frac{C(n)}{n+1} = \frac{2}{n+1} + \frac{2}{n} + \frac{2}{n-1} + \dots + \frac{2}{3} \quad \left[\text{kost } 0, 1, 2 \text{ elementen sorteren in } C(0)=C(1)=0, C(2)=1 \right]$$

$$\Leftrightarrow C(n) = 2(n+1) \left(\frac{1}{n+1} + \frac{1}{n} + \frac{1}{n-1} + \dots + \frac{1}{3} \right)$$



$$\approx \int_2^{n+1} \frac{1}{x} dx = [\ln x]_2^{n+1}$$

willen $\log n$ uitdr. Enkel $\log(n)$

$$\approx \int_2^n \frac{1}{x} dx + \frac{1}{n+1}$$

$$= [\ln x]_2^n + \frac{1}{n+1} = \ln(n) - \ln(2) + \frac{1}{n+1}$$

$$\Rightarrow C(n) = 2(n+1) (\ln n - \ln 2 + \frac{1}{n+1})$$

tilde notatie

$$C(n) \sim 2n \ln(n)$$

$$\sim 1,39 n \log_2 n$$

$$\ln(n) = \log_2(n) = \frac{\log_2 2}{2,303} \log_2 n$$

stel dat we lomuto gekozen ipv Hoare

$$\rightarrow C(n) = (n-1) + \frac{2}{n} [C(0) + C(1) + \dots + C(n-1)]$$

$$\Rightarrow n C(n) = n(n-1) + 2 [C(0) + C(1) + \dots + C(n-1)]$$

$$(n-1) C(n-1) = (n-1)(n-2) + 2 [C(0) + C(1) + \dots + C(n-2)]$$

$$\Rightarrow n C(n) - (n-1) C(n-1) = n(n-1) - (n-1)(n-2) + 2 C(n-1)$$

$$= (n-1)(n - n + 2) + 2 C(n-1) = 2(n-1) + 2 C(n-1)$$

$$\Rightarrow n C(n) = 2(n-1) + (n+1) C(n-1)$$

$$\Rightarrow \frac{C(n)}{n+1} = \frac{2(n-1)}{n(n+1)} + \frac{C(n-1)}{n}$$

$$= \frac{2}{n+1} + \frac{C(n-1)}{n} - \frac{2}{n(n+1)}$$

$$C(n) = 2(n-1) \left(\frac{1}{n} + \frac{1}{n-1} + \dots + \frac{1}{2} \right)$$

$$\approx \int_1^n \frac{1}{x} dx = \ln(n)$$

$$\approx 2(n-1) \ln n \approx 2n \ln(n)$$

extra logische orde kun

rekening met houden in rechte expansie

is uiteindelijk genoegend die ~

$$\Rightarrow \sim 1,39 n \log_2 n$$

(bottom-up) aangepak 2

kan dat 2 elementen in de rij ooit met elkaar vergeleken zullen worden
optellen van deze kansen $\rightarrow$ probabilistisch verwacht # vgl

aanname lomuto partitionering (nooit meer dan 1 keer 2 elementen met elkaar vergelijken)

stel dat de elementen in de te sorteren rij reeds in volgorde gesorteerd zijn

$x_1, x_2, x_3, \dots, x_n$

[veronderstel geen dubbele waarden]

Wat is nu de kans dat x_i ooit met x_j zal $\bar{w}$ vergeleken?

$\rightarrow$ bekijken groep $\{x_i, \dots, x_j\}$

- solange pivot $v < x_i$ en $v > x_j$ zal de groep in 2ën partitie samen blijven

- als $v > x_i$ en $v < x_j$ dan zal de groep in 2 gedeeltes $\bar{w}$ en tot verschillende partities behoren
 $\rightarrow x_i$ zal nooit met x_j vergeleken $\bar{w}$

- als $v = x_i$ of $v = x_j$ op een moment dat ze nog samen in een groep zitten
 $\rightarrow$ enkel dan x_i met x_j vergelijken

$\Rightarrow$ - als v buiten groep: blijven 2 partitie en gaan recursief verder

- als v in groep: x_i en x_j niet vergeleken $\Rightarrow$ 0 vgl

- als $v = x_i$ of $v = x_j \Rightarrow$ 1 vgl, met de kans dat x_i of x_j de pivot is en x_i dus vgl $\bar{w}$

elementen verder van elkaar

$\rightarrow$ kleinere kans om ooit te $\bar{w}$ vergelijken

& vice versa

$$P_{ij} = \frac{1}{j-i+1} + \frac{1}{j-i+1} = \frac{2}{j-i+1}$$

$\rightarrow$ kansen
 $\rightarrow$ #elementen id groep

$$\Rightarrow \text{totale wort: } \sum_{i=1}^{n-1} \sum_{j=i+1}^n \frac{2}{j-i+1}$$

$$\Rightarrow P_{ij} = \# \text{ verwachte vgl voor } x_i \text{ en } x_j$$

$$= \sum_{i=1}^{n-1} \sum_{k=2}^{n-i+1} \frac{2}{k}$$

$$= 2 \sum_{i=1}^{n-1} \left(\sum_{k=2}^{n-i+1} \frac{1}{k} \right)$$

$$\sim 2n \ln(n)$$

$$\sim 1,39 n \log_2 n$$

) haarm. som $H_N = 1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{N}$
 $\sim \ln(n) + 1$

$\rightarrow$ quick sort duurt ongeveer 3,9% langer dan merge sort (wel geen extra geheugen nodig)

Optimalisering Quick sort

* voor kleinere n ($n < 10$ meestal) insertion sort beter (lagere orde kunnen spelen dan een red)

* de median van 3 willekeurige elementen uit de rij als pivot gebruiken

$\rightarrow$ meer kans om in het midden te delen (best case)

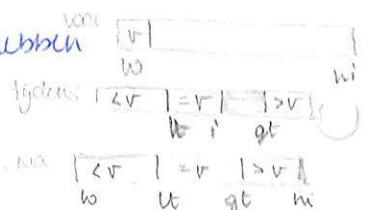
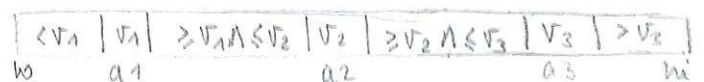
* 3 way partitioning - wim veel elementen gelijkaardige waarden hebben

$\rightarrow$ less, equal, greater subarrays

* multi-pivot quicksort

$\rightarrow$ meer werk per elementje, maar kleinere partitionering
afwijking maken

vb 3 pivots



comparison based sorteralgoritmen kunnen niet beter dan $\sim n \log_2 n$ zoals we hiervoor hebben aangeleerd
 → willen sorteren in lineaire tijd → krachtigere beschrijvingen

Bucket sort

de n elementen in k buckets verdelen $\sim n$

→ elk element in zijn bucket steken gaat in $O(1)$ tijd

elke bucket is gesorteerd volgens een comparison based algoritme

* worst case: alle elementen in 1 bucket [steile heuvel bucket]

→ probleem niet verminderd

→ goede keuze bucket bel!

* best case: n/k elementen per bucket

gelijke verdeling input bel!

→ komt ook een merkwuurde case waarbij we uitgaan van uniforme willekeurige input en dus een gelijke spreiding verwachten

→ comparison based algoritme $\sim n \log_2 n$ toepassen op elke bucket

$$\Rightarrow k \left[c \frac{n}{k} \log_2 \left(\frac{n}{k} \right) \right] = \sim n \log_2 \left(\frac{n}{k} \right)$$

$$= \sim n [\log_2(n) - \log_2(k)]$$

• als k een $O(1)$ is:

$$\Rightarrow \sim n \log_2 n$$

(open verschil met quick / merge sort)

• als k proportioneel is met n

(k schaalbaar lineair met n)

$$\Rightarrow n/k \text{ is een } O(1)$$

$$\Rightarrow \sim c' n$$

$$\text{met } c' = c \log_2(n/k)$$

heuze binnenste sorteralgoritme

- weinig elementen per bucket → insertion sort

(beginnen met sorteren elke keer een nieuw element 'id' bucket is toegevoegd)

- veel elementen per bucket → quick sort of mergesort

- als groepen voldoende is → geen sorteralgoritme binnenin nodig

geheugenprobleem: $\sim n$ extra ruimte

- array van k buckets

- elke bucket bevat een linked list van n/k elementen

Key-indexed / Counting sort

relatief veel gevallen sorteren maar met relatief weinig versch. waarden

= klein # keys → $R = \#$ versch. keys

→ tellen hoeveel elke key voorkomt → weten dat de eerste sorteer plaatsen die key gaan hebben, de volgende sorteer de volgende key, enz.

	0	1	2	3	4	5	6	7
a[i]	a	d	c	a	b	c	a	c

	a	b	c	d	e	-
count[i]	0	3	1	1	1	2
	a	b	c	d	e	-
	0	3	4	5	6	8

0 keys kleiner dan a
 → a begint op a[0]

4 keys < c → c begint op a[4]

	0	1	2	3	4	5	6	7
aux[i]	a	a	a	b	c	d	e	e

```

int N = a.length;
int[] count = new int[R+1];
for (int i=0; i<N; i++)
    count[a[i]+1]++;
for (int i=0; i<R; i++)
    count[i+1] += count[i];
for (int i=0; i<N; i++)
    aux[count[a[i]]++] = a[i];
for (int i=0; i<N; i++)
    a[i] = aux[i];
    
```

→ counting array = teller

tellen vld freq. : lineair in N
 [offset van 1]

geaccumuleerde som bepalen
 ~ lineair in R

hulparray om bij te houden hoeveel elementen al geplaatst zijn
 elementen al geplaatst zijn
 elementen terugplaatsen

→ algoritme geen elementen met elkaar vergelijken → impliciete orde
 complexiteit: in termen van array access (waarde opzoeken, wijzigen of bewerken)
 ~ $8N + 3R + 4$ (om N elementen met keys integers t/m 0 en $R-1$ te sorteren)

→ kan variëren afh. van hoe je telt, sowieso met $\sim N \log N \sim R$

complexiteit: extra geheugen ruimte

$\sim N + R$

(counting array + auxiliary array)

* in-place sorteeralgoritme

= sorteeralgoritme gaat vol array / eigen datastructuur zelf elementen verplaatsen
→ geen extra geheugen nodig

vs selection sort, insertion sort...

vs out-place sorteren (extra geheugen nodig)

vs merge sort, counting sort...

* stabiel sorteeralgoritme

= oorspronkelijke volgorde van elementen met dezelfde key/waarde

≠ NIET veranderd in plaats sorteren

→ oorspronkelijke volgorde behouden

⇒ hiërarchisch sorteren mogelijk [vorige sortering blijft behouden]

vs counting sort

↳ belangrijk voor heel grote data vs string karakter per karakter mogelijk
in vergelijkend traag gaat

- Least Significant Digit sort

(radix sort)

radix = basis (eenheden, tientallen en zo)

! alle getallen / strings moeten dezelfde lengte hebben
belijgt de karakters van rechts naar links

↳ begin sorteren met LSD (laatste karakter ~ eenheden)

dan volgende digit (nieuwe LSD, nu voorlaatste karakter ~ tientallen)
herhalen tot LSD eerste karakter is

gebruiken hiervoor telkens counting sort

(aangezien counting sort stabiel is blijft de oorspronkelijke volgorde via LSD's behouden)

$\sim (7N + 3R)W$

array access

(met W = woordlengte)

$= \sim 7NW + 3RW$

7 of 8 hangt af van initialisatie / hoe je telt
bel!!! is dat het lineair verloopt

$\sim N + R$ ruimte

- Most Significant Digit sort

! kan strings van verschillende lengte wel sorteren

~ soort van bucket met counting sort

belijgt karakters van links naar rechts

↳ karakter per karakter sorteren met counting sort beginnende met het eerste,
dan recursief toepassen op het tweede, dan recursief op het derde enz., uiteindelijk dan
tegen alle samenvoegen

↳ in groepjes ingedeeld op het 1^e karakter

in elk van deze groepjes nu counting sort opvoeren

in dit per groepje blijven herhalen

→ groepjes onafhankelijk (recursief) van elkaar sorteren

globaal gezien $\sim N \log_R N$ karakters onderzocht

strings

versch waarden per string

karakters / ingesloten alfabet

MSD belijgt het genoeg karakters
om de keys te sorteren

onderzochte karakters oft van keys

↳ in elke counting sort elk karakter belijgt
telkens een niveau R naar beneden

! kan sublineair zijn

↳ door ruimte in counting sort vaak opgevoeren

voor klein # digits / strings kan daar voor een grote overhead zorgen

→ voor klein # strings insertion sort gebruiken

in het voor een bepaald karakter, moet wel nog steeds stappen stap

3-way string quicksort

- sort op strings
- 3 partities: kleiner dan, groter dan of gelijk aan de pivot (niet 2 pivots)
- vergelijken van strings zoals beschreven in MSD (1^{ste} karakter eerst)
- maar dan quicksort rekent groepen ipv counting sort
- waardoor de R afhankelijkheid wegvalt
- $\sim 1,39N \log_2 N$ karakter compares (i.e.t. $\sim 1,39N \log_2 N$ string compares stand. quicksort)
- strings op basis van 1^{ste} karakter opdelen in 3 partities
 - * voor $>$, $<$ pivot: rekent hetzelfde karakter gebruiken om de groep verder te sorteren
 - * voor $=$ pivot: rekent het volgende karakter gebruiken om de groep verder te sorteren
- rekent het opnieuw vergelijken van het begin vld string gebruikt nu genoeg compares
- rekent maar wenn de string lang zijn

4.3 x 4.4

Stacks & Queues

- fundamentele datatypes: sets van objecten met basisoperaties: insert, remove, isEmpty, verzameling van elementen

Stacks

- element dat we verwijderen is het meest recent toegevoegde element $\sim$ stapel
- push (element toevoegen) en pop (element verwijderen) gebruiken
- nu maar achteraan (of voraan afh van implementatie) de verzameling
- void push(Object o), Object pop(), boolean isEmpty()
- nu moet zien dat voor iets anders opgesplitst is en anders afhalen dan dat taken object

* implementatie met gekoppelde lijsten

Linked list = datastruct. met elementen die aan elkaar gekoppeld zijn met verwijzingen/pointers - alle elementen naast elkaar $\rightarrow$ linked list

→ klasse Node met 2 datavelden

- item: object met effectieve data [type String, Integer, ...]
- next: pointer naar het volgende object [van type Node]

→ Stack = 1 variabele [Node] bijhouden die verwijst naar het allereerste element vld gekoppelde lijst, zeg first

→ pop()

first $\rightarrow$ [] $\rightarrow$ [] $\rightarrow$ []
 first $\rightarrow$ []
 now item to return { String item = first.item; // uit object first [klasse Node] is item teruggegeven
 first = first.next;
 return item; }
 en first is overschreven die pointer next
 (2^{de} element is nu top)

→ push(String item)

{ Node oldfirst = first;
 first = new Node();
 first.item = item;
 first.next = oldfirst; }

// first verwijst naar het nieuwe element
 // pointer toevoegen om van nieuwe 1^{ste} el nu 2^{de} te gaan

performantie:

pop $\sim$ tijd { N elementen pushen/poppen $\sim$ N tijd

push $\sim$ tijd

stack met N items: $\sim$ N geheugen

* implementatie met arrays

verruimen een geheugenblok met array $s[]$ voor de stack.

houden de plaats bij, zeg N , waar volgend vrij element is

→ $push(item): s[N] = item; \quad item \text{ toevoegen op plaats } s[N]$
 $pop(): \text{ element op plaats } s[N-1] \text{ verwijderen}$

↳ voor initieel int $N=0$:

$push(item) \{ s[N++] = item; \}$
 $pop() \{ \text{return } s[--N]; \}$

→ gemakkelijker want moeten nu geen pointer bijhouden, maar wel een geheugenblok verruimen

! resizen als array vol is (array heeft maar bepaalde capaciteit)

→ nieuwe array maken groter dan de vorige

moeten alle elementen vld oude array in de nieuwe kopiëren (KOST!)
 we willen geen te grote (veel geheugen), maar ook geen te kleine (te vaak resizen) nieuwe array

optie 1 $push$ op een volle array → vergroot lengte met 2
 pop " " " " → verklein " " 1

[beginnen met stack grootte 0]

→ kost om een stack met N items te hebben

$N + 2(1 + 2 + 3 + \dots + N-1) \sim N^2$ array accesses

↑
 kost voor elke
 nieuwe $push$ op
 de stack

↑
 $L = \sum_{k=1}^{N-1} k = \frac{(N-1)N}{2}$
 kost om telkens
 de array te resizen
 van k naar $k+1$ elementen

→ kost kopiëren zijn
 telkens 2 array accesses
 [copy & paste]

optie 2 $push$ op een volle array → verdubbel de lengte vld array

→ kost om de eerste N items op de stack te hebben

$N + 2(1 + 2 + 4 + 8 + 16 + \dots + N/2) \sim 3N$ array accesses

↑
 kost voor elke
 $push$

↑
 kost om telkens
 de lengte te
 verdubbelen

we moeten minder vaak resizen en hebben langer een vaste kost per element

array doen knippen om geheugen te besparen

we moeten oppassen dat de rij geen acceden w door het veelvuldige
 $push$ en pop in de buurt van de volle array

⇒ lengte gehalveerd want de array $\frac{1}{4}$ vol zit

Linked list vs dynamische array implementatie

linked list - worst case duurt elke operatie $O(1)$ tijd

- gebruikt extra tijd en ruimte die de pointers

dyn array - elke operatie duurt $O(1)$ amortiserende tijd

- minder geheugen verspilling

amortiseren =
 afbreken door
 regelmatige aflossing

- Queues

element dat we verwijderen is het minst recent toegevoegde element, ~ wichtig

* implementatie met Linked lists

we hebben nu zowel een pointer naar het 1^e als naar het laatste element

→ dequeue()

$d \text{ string } item = \text{first.item};$
 $\text{first} = \text{first.next};$
 $\text{return } item; \}$

$\text{first} = 1^{\text{e}} \text{ object toegevoegd}$
 → head

→ enqueue(string item)

$d \text{ Node oldlast} = \text{last};$
 $\text{Node last} = \text{new Node}();$
 $\text{last.item} = item;$
 $\text{last.next} = \text{null};$
 $\text{oldlast.next} = \text{last}; \}$

$\text{last} = \text{laatste object toegevoegd}$
 → tail

↳ allebei $O(1)$ tijd

→ queue van N ~ N tijd
 ~ N geheugen

* implementatie met array

array $q[]$ om items in q te slaan

→ enqueue (item) : $q[tail] = item$;

dequeue () : element op $q[head]$ verwijderen

→ in head en tail steeds updaten modulo de capaciteit ook hier dynamisch wijzen

round robin system

als je elementen toevoegt ald tail en je komt op het einde van array kun je je tail laten rondlopen dat je terug op pos 0 zit - doorlopen

- Hash tabellen (prim database)

elementen opslaan die we nu identificeren met een bep. sleutel

→ data records waar key (unieke identifier) plaatsen in array kunnen value (alle secundaire data) hebben

en ook die key ook opzoeken of verwijderen

el key is een integer

→ array met mapping sleutel 1 op 1 : value met key k op $a[k]$

wordt echter veel complexer: complexe sleutel, $max(k) > array.length, \dots$

→ hash functie : manier waarop we de sleutels op posities van tabel mappen

- gelineariseerde sleutel op juiste plaats

! injectie : meer keys ≠ sleutels dan plaatsen in array

→ dealen met collisions (zie later)

hash functie

we mappen sleutel k op een integer $\in [0, M-1]$ (met $M = table.length$) mappen

unwoudig

uniforme verdeling sleutels

zelfde sleutel → zelfde hash

⇒ modulair hashen

(implemte hash functie)

modulair hashen

hash value = key % M

pos. int

→ priemgetal (leidt tot minder collisions - minder hashes met dezelfde value)
stel M niet priem, dan als $k = base \cdot 10$ en $M = 10^k$ → veel minder variatie waarde in uitkomsten

willen juist veel variatie (minder collisions) → willen alle digits sleutel betrokken bij berekening

basisonderstelling: keys uniform verdeeld → hashes helpt ook zo uniform mogelijk verdeeld

collisions

tenzij M extreem groot is, zullen er altijd wel 2 sleutels op dezelfde index gehasht w

2 manieren om daarmee om te gaan

1) Separate chaining

hash tabel interpreteren als een array v. ge linkte lijsten

→ hash: mapt key op integer i in 0 en $M-1$ met $M = length$ array

index: element voraan in i de linked list zetten (tenzij die er al in zit en nog niet staat, mag blijven staan, wel value overwriten)

search: i de ge linkte lijst aflopen totdat element gevonden

willen de gem lengte v. linked list kort houden, want anders schakelen

(hash tabel uitbreiden) en dat is weer een kost

we verwachten dat elke list $\alpha = \frac{N}{M}$ (load factor) elementen heeft met $N = \# \text{elementen}$, $M = \# \text{plaatsen in hash tabel}$

uit observaties valt af te leiden dat de mate waarin de bezetting afwijkt van α binomiaal verdeeld is

$\Rightarrow$ kans dat k sleutels op dezelfde plaats gemapt is

$$\binom{N}{k} \left(\frac{1}{M}\right)^k \left(\frac{M-1}{M}\right)^{N-k}$$

↓
permutatie
versch opties
("combinatie")

↓
kans dat
 k keys in
1 hokje

↓
kans $N-k$
andere keys in
andere $M-1$
hokjes

$$\text{met } \binom{N}{k} = \frac{N!}{k!(N-k)!}$$

$\hookrightarrow$ zien dat de kansen ~~het~~ snel afnemen om k keys op eenzelfde plaats te mappen als k afwijkt van verwachte $k = \alpha$

$\Rightarrow$ # keys in elke linkheid list is ongeveer N/M [in de veronderstelling dat alle keys uniform mappen]

$\Rightarrow$ verwachte kost om elementen te zoeken $\sim N/M = cte$

want heel linkheid list aflopen

als M te groot $\rightarrow$ veel lange ketens

als M te klein $\rightarrow$ ketens te lang

↑
keuze typisch $M \sim N/5$

of $M \sim N/4$

$\Rightarrow$ cte $N/M \Rightarrow$ zoeken in cte tijd

verwachte kost om een element toe te voegen

* key is uniek aan element

al een element met die key is al hash tabel $\rightarrow$ oude element overschrijven met nieuwe el.
b.v. toevoegen eerst checken of die key er al in zit of nog niet $\Rightarrow \sim N/M$

* key is niet uniek aan element, maar waarde (secondaire data) wel

$\rightarrow$ element gewoon toevoegen vooraan aan de linkheid list $\Rightarrow$ cte tijd

verwachte kost om elementen te verwijderen

wordt vaak heel de tijd aflopen om het te verwijderen element te lokaliseren
 $\Rightarrow \sim N/M$

revisie $\rightarrow$ rehashing: nieuwe hash tabel met meer hash waarden vraagt erom om de nieuwe hash van elementen uit de oude hash tabel te berekenen in te plaatsen $\rightarrow$ KOSTELIJK

$\hookrightarrow$ doen we enkel als α boven bepaalde waarden

goal: $\alpha = N/M = cte$

! keys altijd bijhouden, hash is niet voldoende om specifieke element terug te vinden

2. Linear probing

lineair ~~proberen~~ ~~oeverzoeken~~ waar de volgende lege plaats is of array is

N keys in tabel met grootte $M > N$ opslaan

(als ~~het~~ gewenste slot bezet is, neem je gewoon het eerstvolgende vrije (linker) aflopen)

! parken nog meer elementen in je hash tabel dan er plaats is

zoeken in tabel is nu complexer (werken op voorhand niet waar het element gaat zitten)

$\rightarrow$ bereken hash

als $key == \text{search key} \rightarrow$ element gevonden

als hash een lege positie geeft $\rightarrow$ element niet in tabel

als $key \neq \text{search key} \rightarrow$ volgende slot proberen

herhalen totdat 1 vrij slot of 2 opties voldaan is

} kan mogelijk veel tijd kosten

nadeel: neiging om clusters te vormen

hoe langer de cluster, hoe meer kans dat je element achteraan komt

$\rightarrow$ nog langere clusters

+ mogelijkheid dat 2 clusters $\rightarrow$ samengevoegd tot een supercluster

$\hookrightarrow$ kost zal groter $\rightarrow$ (zoeken + toevoegen zal mogelijk langer duren)

$\Rightarrow$ willen clustervorming vermijden, b.v. door quadratic probing

voordeel: zal altijd plaats vinden zodra tabel niet vol
+ geen gelinkte lijst nodig

revisie $\rightarrow$ weer alle keys rehashen

wanneer en hoe revisen?

grootte linear probing: $\alpha = N/M < 1$ (door initiele aanname)

search hit (zoeken element in hash tabel, # pogingen nodig om te vinden)

$$\sim \frac{1}{2} \left(1 + \frac{1}{1-\alpha}\right)$$

March min: # pogingen om te bereiken dat een element niet in de hash tabel zit = # pogingen voor inserts

$$\sim \frac{1}{2} \left(1 + \frac{1}{(1-\alpha)^2}\right)$$

voor de keuze van $\alpha = 1/2$ is dit: March hits = $3/2$
March misses = $5/2$

→ goal voor inserten: $\alpha \leq 1/2 \Rightarrow$ array verdubbelen of halveren indien nodig
element verwijderen uit hash tabel
maken "tombstone" achterlaten op plaats van verwijderde element (een X)
dus marshering zorgt ervoor dat we bij het zoeken naar een element weten dat hier iets verwijderd is en dus niet denken dat een bepaald element niet in de tabel zit want het zit achter de tombstone bevindt.

Separate chaining

- delete makkelijk te implementeren
- clustering minder gevoelig voor een slechte hashfunctie
- performantie verduidelijk meer

vs

Linear probing

- delete vraagt aandacht
- clustering kan een probleem vormen
- minder nutteloze ruimte

2.4, 3.2, 3.3

- Priority queues

= datastructuur (queue) waar we elementen op zetten, maar zo dat het grootste element vooraan staat (niet maakt niet uit $\neq$ queue)

- wat de basisoperaties: insert(), max(), delmax(), empty()

operaties die het echt werk doen

* als η goedkend - makkelijk max bepalen

delmax = laatste element afdrukken ~ 1

max = " " opvragen ~ 1

insert: hele lijst aflopen om te weten waar in te voegen $\sim N$

* als η ongeordend - makkelijk toevoegen

insert: gewoon ^{aan het einde} toevoegen ~ 1

delmax, max: hele lijst aflopen $\sim N$

doel: willen insert en delmax aan elkaar gelijkstellen opdat zelfde complexiteit

→ gebruiken binaire boom als structuur

	insert	delmax / max	
ongeordeerd η	1	N	N → vnl toevoegen
geordeerd η	N	1	1 → vnl verwijderen
binaire heap	$\log N$	$\log N$	1 → evenwicht

- Binaire boom

→ complete binaire boom: elk element (buiten de bladen) telkens 2 vertakkingen
kezelmaat opgesteld vbo en vlnr

↳ diepte vlnr EBT: N knopen → $1 + \lfloor \log_2 N \rfloor$ niveaus / diepte / hoogte

operaties beperkt die diepte ($\log_2 N$)

↳ als nu elke knop van die boom een bepaalde waarde heeft, en η ,
waarvoor geldt dat de η vlnr onder η is dan die van
zijn kinderen, dan spreken we over een (max) heap

hier meer bepaald een binaire heap

binaire boom
= η d'knopen
die links en
rechts met
binaire bomen
verbonden
zijn

een binair heap kan weergegeven worden als een boom, maar ook als een array
 elke knoop w genummerd 2^w en vbo met de wortel root zijnde 1
 in het laatste blad N

deze nummers komen nu overeen met de positie van de keys in de array
 → * grootste key @ $a[1]$

- * elke ouder van knoop k is op $\lfloor k/2 \rfloor$
- * kinderen van knoop k zijn op $2k$ en $2k+1$

↳ kunnen geen uitspraken doen over de positie vld kleinste key
 buiten dat het zich ergens bevindt van $\lfloor \frac{N}{2} \rfloor + 1$ t.e.m. N

de binair heap kan nu gebruikt w als voorstelling voor een priority queue
 ouder van N is $\lfloor N/2 \rfloor$
 en kleinste element heeft
 geen kinderen

→ $\text{max}()$ & return $a[1]$;

inserten & $\text{delmax}()$ maken gebruik vld functies $\text{swim}()$ en $\text{sink}()$

• $\text{swim}()$

stel dat je een geldige heap structuur hebt, maar
 een knoop plots een waarde groter dan die van zijn ouder heeft
 ⇒ moeten ouder en kind van plaats wisselen

- moeten niet naar beneden kijken of de heap nog in orde is,
want $\text{oudkind} = \text{ouder} > \text{oudouder} = \text{kind}$
- moeten wel checken of de nieuwe ouder nu niet groter is
dan zijn ouder

→ als dit we zo is, moeten we iteratief dit proces
 herhalen totdat de ouder terug groter is
 of totdat we de root hebben bereikt.

private void swim(int h) knoop die jouwt zit (nummer)
 { while (h > 1 && less(h/2, h)) mogelijk
 d exch(h, h/2); $h/2$ is de ouder van knoop h
 h = h/2; } }

worstcase zit de wortel helemaal beneden ⇒ $\sim \log_2 N$ vgl'en
 best case zit alles juist → 1

↳ insert

public void insert(Key x) {
 pq[++N] = x;
 swim[N]; }

// knoop van onder toevoegen (volgende
 in heap en naar boven
 laten zwemmen)

worstcase $\log_2 N$ vgl'en

→ heap opbouwen

optie 1 elke element (vld N elementen) om de beet in de heap inserten
 heap w groter en groter

→ # vgl: $\log_2 1 + \log_2 2 + \log_2 3 + \dots + \log_2 (N-1) + \log_2 N$
 (worstcase) $= \log_2 (N!)$

$\approx N \log_2 (N)$ 2 Stirling

⇒ $\sim N \log_2 (N)$ tijd (vergelijkingen)

• $\text{sink}()$

stel dat een knoop plots kleiner is geworden dan minstens 1 van zijn kinderen

⇒ moeten ouder en grootste kind met elkaar wisselen

moeten dit blijven herhalen totdat de heap terug hersteld is

private void sink(int h)

{ while (2*h <= N)

 d int j = 2*h;

 if (j < N && less(j, j+1)) j++;

 if (!less(h, j)) break;

 exch(h, j);

 h = j;

// blijven herhalen zolang knoop
 nog kinderen heeft of tot break

// grootste vld 2 kinderen w j

// vgl'en met grootste vld 2 kinderen

// kind w met k checken ouder

```

private void sink (int k) {
  while (2*k <= N) {
    int j = 2*k;
    if (j < N && less(j, j+1)) j++;
    if (! less(k, j)) break;
    exch(k, j);
    k = j;
  }
}

```

2 vergelijkingen per loop (inderdaad met elkaar x grootste met ouder)
 + $\log_2 N$ niveaus

⇒ max $2 \log_2 N$ vergelijkingen

→ heap opbouwen

optie 2 elk element zit al in de heap/array

heap geldig maken door van onder naar boven
 (beginnende bij knoop $\lfloor N/2 \rfloor$, laatste ouder) niks toe te
 passen (tot knoop 2) en bij elke stap steeds ervoor te
 zorgen dat alles terug in een geldige heapstruct staat

→ heap bottom-up geconstrueerd

worst case & conservatieve schatting:

niks toepassen op $N/2$ knopen ~~dat~~ elk max $\log_2 N$ niv. kan zinken
 $\sim N/2 \cdot 2 \log_2 N \sim N \log_2 N$ vgl'en

echter overschatting want voorlaatste niv kan max 1 niv zinken

⇒ meer exact: $\sim 2N$ vgl'en

↳ delmax

1 element is het max → verwijder uit heap

en laatste element nu helemaal verplaatsen, zetten en laten zinken

public Key delMax() {

Key max = pq[1];

exch(1, N--);

sink(1)

pq[N+1] = null;

return max; }

worst case $2 \log_2 N$ vgl'en

// N met 1 verminderen omdat we eigenlijk
 nog maar N-1 knopen hebben

// waarde vorige laatste knoop (die nu de
 grootste is) op null zetten

→ heap sort

n elementen die we willen sorteren in een ^{array} priority queue zetten

→ in geldige heapstruct (opdat pq)

- telkens grootste element verwijderen en achteraan in een array zetten

- heap 1 korter maken (rearrange)

↳ n keer herhalen ⇒ ge sorteerde array als resultaat

nodige operaties: heapifying + (1 voor 1 delmax toepassen) N

$\sim 2N \log_2 N + \sim 2N \log_2 N$

⇒ $\sim 4N \log_2 N$

- Binair zoekbomen (BST)

kan heel onregelmatige structuur hebben (afh van diepte, telkens 2 vertakking maar mag
 ook leeg zijn)
 enige voorwaarde: elke knoop i of de linker subboom kleiner dan ~~alle~~ de

↳ symmetrische

knopp erboven

ouder

elke knoop i of rechter subboom groter dan ~~alle~~ knopp erboven

→ makkelijk om te doorzoeken

- niet heel boom doorzoeken, maar telkens gewoon

gebruikt kleine Node met volgende var:

in juiste subboom blijven kiezen

private class Node

{ private Key key;

private Value value;

private Node left, right;

public Node (Key key, Value value) {

this.key = key;

this.value = value; }

- search: geeft value terug van gegeven key of null als de key er niet is
 voor link: # vgl'en = diepte vld knoop
 voor min: # vgl'en = diepte van waar de knoop zou zijn
- insert: zoek de key in de boom
 - als de key vld boom zit → value wettten alsoer geen dubbele zijn toegelaten
 - als de key niet vld boom zit → nieuwe knoop toevoegen op null link
 [blad onderaan in juiste subtrah]

veel verschillende BST's zijn mogelijk voor dezelfde set keys
 # vgl'en = diepte vld knoop (onvoorspelbaar)

gemiddelde gedrag vld zoekboom als we willekeurig elementen toevoegen?

→ diepte $\sim 1.39 \log_2 N$ → # vgl'en voor search keys

|| valider echter onregelmatige diepte die willek. toevoegen

analyse analoog aan quicksort ^{pivot} met 2 partities (root knoop met 2 takken)

$$C(n) = 2 + \frac{1}{n+1} [C(0) + C(1) + C(2) + \dots + C(n-1)]$$

2 vergelijkingen telkens

- gelijk aan knoop
- groter / kleiner dan knoop

→ bij aflopen kan de verzameling waarin we aflopen 0 elementen hebben, 1 element, 2 elementen, ..., n-1 elementen

elk van die scenario's hebben een gelijke kans om voor te komen

$$\Rightarrow (n+1)C(n) = 2(n+1) + [C(0) + C(1) + \dots + C(n-1)]$$

n+1 scenario's: element gevonden, tak met 0 elementen, 1 element, ..., n-1 elementen

$$n C(n-1) = 2n + [C(0) + C(1) + \dots + C(n-2)]$$

$$\rightarrow \text{aftrekken: } (n+1)C(n) - n C(n-1) = 2 - C(n-1)$$

$$\Rightarrow (n+1)C(n) = 2 + (n-1)C(n-1)$$

$$\Rightarrow C(n) = \frac{2}{n+1} + \frac{n-1}{n+1} C(n-1)$$

→

- delete:
 - simple: maak de knoop leeg, maar waar hem staan gemarkeerd met een tombstone
 - delete min: ga zoveel mogelijk naar links tot je een null link link tegen komt
 kwang nu deze knoop door zijn rechterlink (en update de tellingen)
 - Hibbard deletion: zoek voor de knoop met key k die je wilt verwijderen
 - * als k geen kinden heeft: knoop gewoon verwijderen en door null link vervangen
 - * als k 1 kind heeft: knoop verwijderen en link vervangen door de link naar het kind van k
 - * als k 2 kinden heeft: zoek de opvolger van k → minimum in rechter subboom
 verwijder dit minimum en vul de waarde van deze knoop in op de plaats van k

gemiddeld $\sim \sqrt{N}$

• min operations:

findMin → key uiterst links

findMax → key uiterst rechts

floor (= grootste key $\leq k$) → als k == root: floor is root

k < root: floor moet in de linkerboom zitten

k > root: floor kan in de rechter subboom zitten of de root zijn

ceil (= kleinste key $\geq k$) → gelijkaardig aan hierboven maar dan omgekeerd

L problemen BST's

- slechte worst case performantie

- idealiter: gebalanceerde boom → theoretische pogingen hiervoor: 2-3-boom x rood-zwart-boom

- 2-3 bomen

andere type search tree

1 of 2 keys per knoop toegestaan

→ 2-knoop: 1 key & 2 kinderen

3-knoop: 2 keys & 3 kinderen

(kleiner dan L, groter dan R, + m L en R)

→ PERFECT GEBALANCEERO: elk pad van ^{root} knoop tot blad heeft dezelfde lengte

↳ mogelijk 2 testen doen om te weten in welke tak we moeten afdalen

→ search: 2 testen / vgl'en om nrv dalen te kiezen, maar diepte nrv wel vast

• insert: element K toevoegen van onder (tot boom (bladen))

Gebouw: L in welke subtak bepaald of vrio te lezen

+ als de eindknoop waarin we rechtekomen een 2-knoop is

→ in 3-knoop (K gewoon toevoegen)

* als eindknoop een 3-knoop, met een 2-knoop als onder

→ tegelijk 4-knoop maken met 2-knoop als onder

→ middelste waarde (vld 3) vld 4-knoop in

de 2-knoop onder toevoegen en de ~~4~~ 4-knoop

opsplitsen in 2 2-knopen → onder in 3-knoop

+ als eindknoop 3-knoop met 3-knoop als onder

→ zelfde idee als hierboven, maar dit dan 2x herhalen

totdat de onderknoop een 3-knoop is (en geen 4-knoop meer)

* als eindknoop 3-knoop en we komen op de root die ook een 3-knoop is

→ niveau lg maken

boom laten groeien naar boven toe

↳ elke informatie behoudt de perfecte balans en symmetrische orde

⇒ elk pad van root tot blad heeft dezelfde lengte

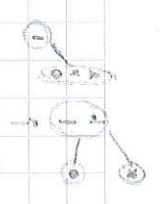
→ overal dezelfde diepte

↳ worst case: allemaal 2-knopen → $\log_2 N$ diep (boven grens)

best case: allemaal 3-knopen → $\log_3 N \approx 0,631 \log_2 N$ diep (ondergrens)

→ we garanderen dat root & insert algoritmes altijd logaritmisch opaan

⇒ search, insert, delete & search hit $\sim c \log_2 N$ met $0,631 \leq c \leq 1$



- Rood-zwart bomen (RB trees)

= 2-3 bomen vertalen in binaire bomen (BST)

→ vertalingen in binaire boom kleur geven

* rood: (a) en (b) vormen samen een 3-knoop → nrv links rechte vertaling

* zwart: gewone vertaling

↳ een veel id implementatie vld knoop vlt knoop houdt kleur vertaling bij

"wqds"

of meer

- geen knoop mag 2 rode vertalingen hebben || enkel 2-3 knopen in 2-3 boom

- elk pad vld root naar ~~et~~ en blad heeft || 2-3 bomen overal vaste

hetzelfde # zwarte links (black-balance)

|| diepte

- rode vertalingen altijd nrv links

↳ variaties op mogelijk (zie rotaties)

→ we alle operaties die we op 2-3 bomen toepassen (bv insert functie)

ook op RB bomen toepassen (rotatie link maken nodig)

↳ als 2 rode in 1 rij (2 opeenvolgende links (of rechts) subtaken)

dan mrv rotatie het middenste element "naar boven brengen"

omv rotaties en dan kleur terug nrv zwart veranderen → flip!

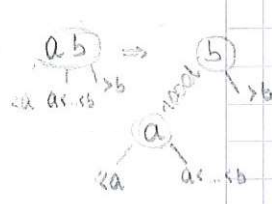
! naar links rechte → ook zwart en rood manipuleren

indien middelste
zwarte verbinding met
ouder
→ wru toe

max diepte RB boom: $2 \log_2 N$

gem $\sim \log_2 N$

rode vertalingen elg



Greedy algorithms

= gulzige algoritmen $\rightarrow$ zetten telkens een stap die op dat moment het meest optimaal lijkt (hopewel dat je zo een goede outcome hebt)
 $\rightarrow$ probleem nooit globaal bekijken, maar opdelen in stappen en die lokaal telkens bekijken
 $\Rightarrow$ geen garantie op globaal het beste algoritme

v.b. coin exchange problem: muntgeld teruggeven in zo weinig mogelijk muntjes/briefjes
 greedy algoritme: bij elke stap steeds het grootste mogelijke briefje/muntje teruggeven (zonder overshoot)

wekt voor de meeste muntstukken als meest optimale oplossing
 echter voor bvb. fictieve "€": 1€, 7€, 10€ wekt dit niet meest optimaal

$\hookrightarrow$ zal 15€ teruggeven als 10€ + 5.1€ ipv meest optimale 2.7€ + 1€

Maple problem: 3 ingeschreven, niet-overlappende wielen in een drieschak plaatsen zodat de wielen zo'n groot mogelijk oppervlak bezetten

$\rightarrow$ greedy algoritme geeft hier de meest optimale opt



2 toepassingën bestuderen

1. Activity Scheduling

zo veel mogelijk niet-overlappende activiteiten plannen in beperkte resource

constraint: begin- en eindtijd van elke activiteit

$\rightarrow$ greedy algorithm oplossen, maar wel criterium gebruiken we om "de beste" activiteit lokaal te kiezen?

Heuristisch 1: activiteit die het 1^e begint als 1^e plannen
 2^e activiteit $\rightarrow$ activiteit die na 1^e het vroegst begint enz.

! geen garantie op meeste # activiteiten i/h ruimtebeslag

v.b. het kan zijn dat de 1^e activiteit 10h duurt, waarvan eigenlijk ook 4 andere activiteiten gepland konden w.

Heuristisch 2: kortste (niet-overlappende) activiteit steeds kiezen

! geen garantie meeste # activiteiten

v.b. kortste activiteit v. 3 act. kan met de andere 2 overlappen waardoor je maar 4 act. kan plannen ipv 2 (iets langwer.)

Heuristisch 3: activiteit met minst # conflicten steeds kiezen

! geen garantie



v.b. kiezen paarw. activiteit niet maar kunnen dan maar 3 act. plannen ipv 4

Heuristisch 4: steeds de activiteit die als 1^e eindigt en niet met de vorige overlapt, plannen

T.B.!

$\Rightarrow$ OPTIMALE OPLOSSING (volgens # activiteiten)

MEEST

$\hookrightarrow$ misschien niet uniek maar wel opt.

Bewijs met het ongreijmde

stel $A = \{a_1, a_2, \dots, a_n\}$ de set van activiteiten gekozen volgens heuristisch 4.

[Deze activiteiten zijn gesorteerd volgens eindtijd (en dus ook volgens begin tijd).]

Vuondstel nu dat er een andere set bestaat, kup

$A' = \{a'_1, a'_2, \dots, a'_m\}$ met $m \geq n$

die $\neq$ meest optimaal is.

$\hookrightarrow$ optimaal $\rightarrow$ meeste # act

We vergelijken nu A' en A . willen bewijzen dat $m = n$

Vuondstel dat $a_1 \neq a'_1$,

dan eindigt a_1 voor (of tegelijkertijd met) a'_1 [heuristisch 4]

dus dan kan a_1 NIET overlappen met a'_2 .

Dan is $\{a_1, a'_2, a'_3, \dots, a'_m\}$ ook een geldige optimale opt.

Via inductie kunnen we zo alle a'_i 's vervangen door a_i , $i=1, \dots, n$ beinvloed den

$\Rightarrow \{a_1, a_2, \dots, a_n, a'_{n+1}, \dots, a'_m\}$ is dan ook een geldige optimale oplossing

Dit is echter in contradictie met heuristisch 4, want we zouden geen activiteiten meer mogen zijn die ~~overlappen~~ niet overlappen met a_n (constructie A), dus $a'_{n+1}, \dots, a'_m$ zijn tegelijkertijd

$\Rightarrow m = n \Rightarrow$ heuristisch 4 geeft een optimale opt. $\square$

allemaal suboptimaal

2 File Compression

je string, voorbeeld B in binaire data, in een zo optimaal mogelijk # bits comprimeren, zeg $C(B)$

dan decomprimeren kan je dan van $C(B)$ terug B verkrijgen

↳ compressie = $\frac{\text{bits in } C(B)}{\text{bits in } B}$ → willen dit zo klein mogelijk

! het ultieme compressie algoritme bestaat echter NIET

IB → geen enkel algoritme kan elke bitstring comprimeren

Bewijs 1 dit contradictie

Stel dat er een compressie algoritme bestaat dat elke bitstring kan comprimeren tot een kleiner # bits.

Dan kan een gegeven bitstring B_0 gecompressieerd $\bar{w}$ tot B_1 ,

maar dan kan B_1 ook gecompressieerd $\bar{w}$ tot B_2

en dit kan herhaald blijven $\bar{w}$ totdat we uiteindelijk met een bitstring eindigen van lengte 1

→ Te kortzichtig!

Als we alle bitstrings kan compr. tot 1^* , kan we ze niet meer decomp. □

Bewijs 2 dit te tellen hoeveel mogelijke geconn. files je kan houden met een vast # bits

Stel dat er een algoritme bestaat dat alle mogelijke bitstrings van 1000 bits kan comprimeren.

2^{1000} mogelijke bitstrings van 1000 bits (elke bit 2 verschillende mogelijke waarden)

Het algoritme zal gecompressieerde strings kan teruggeven van

1 bit lang → 2 x'n strings, 2 bits lang → 4 strings, ..., 999 bits lang → 2^{999} strings

→ we zullen nu

$$1 + 2 + 4 + \dots + 2^{998} + 2^{999} = \sum_{i=0}^{999} 2^i = 2^{1000} - 1$$

strings van lengte 1 tot 999 bits
alle mogelijke strings voor te stellen met minder dan 1000 bits

strings uit het algoritme kunnen volgen

Met < 999 bits kan er namelijk maar $2^{999} - 1$ strings

gecodeerd $\bar{w}$ → informatie verloren

→ gaat niet □

voorbeelden van algoritmen om bitstrings te comprimeren

• Run length encoding

tellen hoeveel opeenvolgende nullen we hebben in bitstring in begin dan het # " enen, dan weer het # nullen en zo voort

→ in plaats van alle nullen en enen bij te houden houden

we dus het #nullen #enen #nullen #enen ... bij

alle van die getallen houden we bij in x-bits (vb 4-bit telling voor max 15)

→ we kan zo bv's een string van 40 bits bijhouden in 16 bits

hieruit kan we ook steeds onze oorsp. string terug opbouwen

aanpakken we onze getallen steeds moeten voorstellen als bits vte vaste lengte

dan het zijn dat we soms 0 nullen of 0 enen moeten voorstellen vb 17 nullen maar

→ vrees van getallen → niet optimaal

maar 4 bit representatie of beginnen met een ipr nullen

↳ alternatieven

• variable length coding

"woorden" kan een verschil lengte hebben, maar dan moet je onderscheid kan maken tns de versch karakters om te weten wann een karakter in word string stopt/eindigt → delimiter introduceren

vb spatie/gap bij morse

vb morse of UTF-8

waarbij veel voorkomende karakters door een kleiner # bits $\bar{w}$ voorgesteld

www.lmsintl.com

we gaan hier uit van lossless compressie
- dus zonder verlies van informatie

kan 1 bit niet lossen decomp. tot versch bitstring

code = voorstelling kar.

Huffman coding

↳ meest optimale manier om bitstruks te comprimeren

= prefixrijge code → geen enkele code is een prefix / begin van een andere code
geen enkele code is de start van een andere code
⇒ geen nood aan een delimiter

↳ voor te stellen door een complete binaire boom

waarbij elke vertakking overeen komt met een 0 of 1 (links: 0, rechts: 1)

en elk karakter correspondeert met een blad v/d boom

NIET met andere knopen die zich niet op het einde bevinden

↳ beste code ~~opt~~ voor elk karakter opdat de tekst in zgn geheel zo kort mogelijk gecodeerd is

→ Huffman coding

prefixrijge code construeren als greedy algoritme

1) 2 minima van pq afhalen

2) combineren onder een nieuwe knoop

freq. knoop = $\sum$ freq. bladen

3) 1 & 2 herhalen totdat alle knopen met de laagste mogelijke freq. gecombineerd zgn

4) 1 & 2 herhalen, maar nu ook de al eerder gemaakte

vertakte knopen gebruiken in de combinatie

→ eindresultaat = code boom

tel de frequentie van elk karakter in de te coderen sequentie

→ 2 minst frequente karakters inkleed in een nieuwe knoop

nu dit telkens herhalen (minst frequente karakters eerst samen combineren) totdat alles gecombineerd is in 1 code boom



implementatie vereist een priority queue om telkens de 2 karakters met de laagste freq te vinden → ~n log n tijd nodig om code op te stellen voor n karakters
! BOOM moet COMPLETE zijn: als een knoop vertakkingen heeft, moeten het er 2 zgn, anders kan de prefixrijgeheid niet gegarandeerd w

Om te bewijzen dat de Huffman codeboom optimaal is, zullen we eerst enkele lemma's beschouwen

Lemma 1 Een optimale codeboom is compleet

B. altijd 2 vertakkingen als er vertakkingen w'reen knoop zijn
andus kan de boom collapse tot een optimale boom
⇒ compleet

Lemma 2 Er bestaat een optimale codeboom in dewelke dat de 2 minst frequente karakters broer en zus zijn op maximale diepte

B. uit lemma 1 volgt dat elke knoop op het laagste niveau een knoef of zus heeft.

zgn nu dat x en y de minst frequente karakters zijn, dan moeten zij op maximale diepte v/d codeboom zitten. Want stel dat dit niet zo was, dan zou er een letter w zgn op maximale diepte die meer frequent voorkomt (veronderstelling x, y) en dat zou niet optimaal zgn.

Stel nu dat x en y geen siblings zouden zgn, dan kan dit gecorrigeerd w door y te wisselen met x's huidige sibling. Dit veranderd de code, maar niet het # bits

⇒ kan altijd zo een boom maken.

□

Stelling De Huffman coding zal leiden tot de meest optimale codeboom.

B. de inductie

basistap: 2 karakters w volgens Huffman

in deze codeboom voorgesteld

elk karakter w hier voorgesteld die 1 bit

⇒ optimaal



inductiestap

veronderstel nu dat we de optimale codeboom kan opstellen voor minstens dan x karakters

we tonen nu aan dat de codeboom voor x karakters dan ook opt. is

(zie volgende pagina)

niet enkel
kan ook anders,
maar dat is er
keken een

optimale boom voor $n-1$ kar.
 uitbreiden tot n karakters $\rightarrow T_H$

Stel dat T_H de Huffman codeboom is voor karakters $s_1, \dots, s_n$ met respectievelijk freq. $f_1, \dots, f_n$.
 Om T_H op te stellen beginnen we met de 2 symbolen met de minste freq te kiezen, zeg s_i en s_j .
 Deze 2 combineren we dan in een gezamenlijke knoop, zeg s^* met freq $f_i + f_j$.
 Nu zijn er nog $n-1$ symbolen die we moeten combineren volgens Huffman.
 Aangenomen we veronderstellen dat Huffman een codeboom die optimaal is geeft voor alle karakters t.e.m. $n-1$, vullen we door Huffman toe te passen op de overige karakters de optimale boom T_H^* .

Dan is de totale lengte van resulterende binaire string (= tot # bits)

$$W(T_H) = W(T_H^*) + (f_i + f_j)$$

elke lengte n minstens 2 langer, $f_i + f_j$ keer

grootte binaire string $n-1$ symb, met s^*

Stel nu dat er een andere ~~optimale~~ boom T_{opt} bestaat voor $(s_1, f_1), \dots, (s_n, f_n)$ die optimaal is

Mit lemma 1 en 2 volgt nu dat de minst freq knopen s_i en s_j op het diepste niv. als broer en zus voorkomen

Maak nu een nieuwe boom T^* door s_i en s_j te mergen in hun ouder $(s^*, f_i + f_j) \rightarrow$ nieuw symb

T^* bevat nu $n-1$ symbolen

Dan is de totale lengte van meest opt boom

$$W(T_{opt}) = W(T^*) + (f_i + f_j)$$

Mit de inductiehypothese (T_H^* is opt. voor $n-1$ kar.) vinden we dan dat

$$W(T_H^*) \leq W(T^*)$$

En aangezien $W(T_H) = W(T_H^*) + (f_i + f_j)$

$$\text{vinden we dat } W(T_H) \leq W(T^*) + (f_i + f_j)$$

$$\Rightarrow W(T_H) \leq W(T_{opt})$$

Dus $W(T_H) = W(T_{opt})$ (de optimale boom) en dus T_H is een optimale codeboom voor n karakters.

Mit de inductie volgt dus dat de Huffman boom optimaal is voor elk # karakters ≥ 2 . □

0, 1 van je bits bepalen niet je code, nu bvb vervangen w die a,b een code is verschillend voor een andere code als je voor je letters een andere lengte van codes krijgt



Minimum opspannende bomen

F.v.I.

grafen: knooppunten (vertices) met onderlinge verbindingen (bogen / edges), wentwel met een gewicht
 ↳ paden t/m 2 knopen?
 ring vld graaf? [Euleraanschuip: alle bogen exact 2 keer gebruikt
 Hamiltoniaanschuip: alle knopen exact 1 keer gebruikt]
 isomorfisme?
 connectiviteit?

- Minimum opspannende boom

= een ^{subgraaf} opspannende boom voor graf G die alle knopen verbindt, zonder een ring te maken met minimaal kost (som gewichten gekozen bogen G (postfig))

- MST is niet uniek

voorstelling grafen als gegevensstructuur - nodig om operaties uit te voeren

* lineaire array of linked list vld bogen

→ houdt de koppels vld verbonden knopen bij

⊕ alle bogen zijn geschreven vanuit 1 knoop in 1 knoop
 ↳ lineair $\sim e$ → geen goede struct.

* 2D-array: 2D $V \times V$ boolean array

→ verbinding t/m i en j : x_{ij} = true (1)

geen boog t/m i en j : x_{ij} = false (0)

⊕ bogen t/m 2 knopen → cks tgd $\sim e \sim 1$

⊕ alle bogen in 1 knoop → lineair $\sim v$

⊖ veel geheugen, ook als heel sparse

⊖ houden hele matrix bij ondanks symm (i, j) = (j, i)

⊖ updaten = updaten van heel de matrix

* lijst of array die voor elke knoop v een gekozen lijst heeft die bijhoudt naar welke andere knopen er een verbinding is

⊕ alle bogen in 1 knoop → cks tgd

⊕ geheugen relatief goed

⊕ redundantie, maar nodig om voorstellen hierboven goed uit te voeren

↳ kopiëren referenties naar obj die bogen voorstellen, niet naar bogen zelf

- MST algoritmes: allemaal greedy algoritmes die gebruik maken vld

cut property

↳ een mede = een partitie vld vertices (knoopen) in 2 niet-lege verzamelingen
 → deelt knopen graf op in 2 groepen

mer de zogenoemde "crossing edges" die knopen vld 2 versch groepen met elkaar verbinden

↳ uitproperty = van alle bogen die door een mede gaan (alle crossing edges), meer diegene met het ^{kleinsten} minimaal gewicht deel uitmaken vld MST

Bewijs uit het ongerijpde

Stel dat e is de min-weight crossing edge

Stel dat e niet in de MST valt, maar dat de verbinding gemaakt is door een andere crossing edge, zeg f .

Dan $w(f) \geq w(e) \Rightarrow w(T \text{ met } f) \geq w(T \text{ met } e)$

waarom T de opgespannen boom is verhoogd door f te verwijderen en e toe te voegen, en T de MST met f is.

Dit is echter in tegenspraak met de def van een MST.

⇒ e moet vld MST zitten.

□

gaan uit vld samenhangende graf met duidelijke gewichten per boog

→ greedy MST algoritme

veruit vanuit alle bogen, waarvan er nog geen id MST zitten (MST init. leeg)
maak een mede door bogen die nog niet id MST zitten
→ voeg de doorgeteste boog met min. gewicht toe ald MST
herhaal totdat er $V-1$ bogen id MST zitten
↳ hoe kiezen we nu deze medes? → voorsh alg.

→ Algoritme van Prim

- start met 1 beg. knoop, zeg vertex 0, id MST T
laat T uitgebreid groeien

↳ groeien als medes de medes die alle bogen

(v, w) met $v \in \text{MST}$, $w \notin \text{MST}$ ~~aanneem dat~~ doorvoegen

steeds de boog die aan 1 kant in en aan de andere kant uit T zit en min. is

↳ Lazy Prim: implementatie als priority queue ($\sim \log_2 n$)

telken bogen (v, w) in P en vragen zo heel gemakkelijk de

boog met het minimale gewicht op

! echter: niet alle bogen in P blijven geldig om toe te voegen

want oudere bogen kan knopen veroorzaken

Weten oude bogen
min. zitten en
blijven gevonden
nieuwe (v, w) knopen

→ checken bij selecteren min of beide knopen vld boog niet al in T zitten
indien wel zo → P opnieuw organiseren zonder die boog ($\sim \log_2 n$)

en opnieuw min. selecteren

↳ dit kan voor overhead zorgen

maar kost om telkens weer voll. nieuwe P op te stellen te hoog

analyse

worst case queue kan max E bogen bevatten → P vaak veel kleiner dan E

→ boog deleten / toevoegen $\sim \log_2 E$

ruimte nodig voor P $\sim E$

⇒ alle bogen overlopen: tijd $\sim E \log_2 E$

ruimte $\sim E$

Eager Prim → better Prim (overbodige bogen niet langer bijgehouden)

houden P van knopen bij die door een boog met T verbonden zijn

als prioriteit houden we het gewicht vld kortste boog die de

knop met T verbindt bij [goedkoopste verbinding]

min element v verwijderen en bijhorende knop toevoegen aan T v is (v, w)

- P updaten: voeg nu alle knopen x toe ald P die verbonden

zijn met de toegevoegde knop v

* als $x \in T$: neger

als $x \in T$

* als $x \in P$: voeg x toe aan P

* als $x \in P$: update de prioriteit
van x als (v, x) korter is dan de
al bekende verbinding

→ voordeel max V elementen in P , al wel iets meer tijd die updaten

analyse: ruimte $\sim V$

tijd $\sim E \log_2 V$

iets beter, vooral voor dichte grafen,
voor sparse echen $E \sim V$

→ algoritme van Kruskal

bogen beschouwen in oplopende volgorde (vb alle bogen id P en P steeds updaten)

voeg nu telkens de volgende boog (globaal min boog) toe aan MST T

TENZIJ dit een knoop zou veroorzaken

→ directe algoritme nodig

detectie / algoritme gebaseerd op union find:

oude tot welke substructuur elke vertex al behoort

in begin allemaal -1 want nog niets geselecteerd

elke keer dat je een boog toevoegt, moet je telkens eerst de substructuur van knopen opvragen

als de knopen in verschillende ^{bomen} ~~bomen~~ ^{mogelijk} → mogen met elkaar verbonden worden 1 vlot 2 knopen waardoor te geven vlot andere knoop en de waarde vlot andere knoop te verlagen met 1 ($\rightarrow -2$)

analyse Kruskal: in principe elke edge te selecteren

→ tijd $\sim E \log_2 E$

in algemeen hoger dan Prim

toepassing: hiërarchische clusterstructuur opbouwen

Kortste pad algoritmen

[in gewichte grafen $\rightarrow$ boogen hebben naast een gewicht ook een richting]

we gaan ervan uit dat er een kortste pad bestaat van knoop ^{start} ~~s~~ tot elke knoop ^{andere} ~~v~~

→ goal: dit pad vinden

- we zien dat er steeds een kortste pad boom (SPT) opt bestaat

→ SPT voor te stellen als 2 arrays: edgeTo[] en distTo[]

vertex-indexed $\rightarrow$ edgeTo[] = [null, 5 $\rightarrow$ 1, 0 $\rightarrow$ 2, 7 $\rightarrow$ 3, 0 $\rightarrow$ 4, 4 $\rightarrow$ 5, 3 $\rightarrow$ 6, 2 $\rightarrow$ 7]

edgeTo[v] = laatste boog op kortste pad ~~naar~~ van s naar v

distTo[] = [0, 1,05, 0,26, 0,97, 0,38, 0,73, 1,19, 0,10]

distTo[v] = lengte vln kortste pad van s naar v

- clue kortste pad algoritmen: edge / boog relaxatie

stel we hebben pad s $\rightarrow$ w $\rightarrow$ v

en we hebben een pad s $\rightarrow$ w

als de verbinding v $\rightarrow$ w een korter pad oplevert dan de huidige verbinding s $\rightarrow$ w

$\Rightarrow$ laatste boog die in s $\rightarrow$ w naar w wijst (edgeTo[w]) wipdoen

en kortste pad boom updaten met v $\rightarrow$ w

not als e = v $\rightarrow$ w korter pad naar w via v, update distTo[w] en edgeTo[w]

relaxatie boog

• Heel algemeen SPT algoritme

initialisatie: distTo[s] = 0 en distTo[w] = ∞ $\forall w \neq s$

herhaal dan: relaxeer eender welke boog $\rightarrow$ hij is het kortste afstand tot een knoop is, die boog toevoegen aan SPT, (eventueel) andere boog updaten die betere verbinding maken die nieuwe knoop

Bewijs Doorheen het algoritme is distTo[w] de lengte vln simpel pad van s naar v en edgeTo[w] de laatste boog in dat pad

Elke succesvolle relaxatie verlaagt distTo[w] voor een bep w.

distTo[w] kan hooguit een eindig # keer verminderen $\rightarrow$ stopt ooit $\rightarrow$ distTo[w] min $\forall w$ $\square$

$\hookrightarrow$ welke boogen we relaxeren hangt af vln algoritme

$\hookrightarrow$ optimaliteitsvoorwaarden:

distTo[] zijn de kortste pad afstanden vanuit s

and * voor elke knoop v: distTo[v] de lengte is van een pad van s naar v

* voor elke boog, e = (v,w): distTo[w] $\leq$ distTo[v] + e.weight()

(geen verdere relaxatie mogelijk)

• Dijkstra's algoritme

belijk alle vertices in stijgende orde van afstand tot s $\rightarrow$ niet-SPT vertex met de laagste distTo[] value

voor de laagste vertex toe aan SPT en relaxeer alle boogen verblijvende met die vertex

(herhaal voor alle vertices)

Bewijs Dijkstra geldig

Elke ~~steer~~ ~~boog~~ $e = (v \rightarrow w)$ $\rightarrow$ precies 1 keer geupdate (wanneer v is toegevoegd a/d SPT), dus $\text{distTo}[w] \leq \text{distTo}[v] + e.\text{weight}()$

Deze ongelijkheid blijft geldig, totdat het algoritme stopt omdat

- $\text{distTo}[w]$ niet kan stijgen; we hebben enkel niet neg. gewichten

en $\text{distTo}[v]$ niet kan veranderen; $\& v$ is niet meer geupdate

Dus ^{tot} bij het eindigen is de kortste pad optimaliteitsvoorwaarde gegarandeerd $\square$

analyse

tijdscomplexiteit $\approx E \log_2 V$

want stuur (net zoals Prim) op een PA van V knopen

en elke boog kan mogelijk eens geupdate w in SPT $\rightarrow$ update PA E keer

het is met Dijkstra ook mogelijk om het langste pad te bepalen

door alle gewichten negatief te maken en het kortste pad te zoeken

! echter: Dijkstra werkt niet met neg. gewichten, want dan kan je door een lus blijven lopen alsood je gewicht in de lus blijft afnemen

$\rightarrow$ alternatief algoritme dat wel met neg. gewichten werkt

vb Belman-Ford

$\approx$ Dijkstra, maar houdt ook het # keer dat een knoop geselecteerd

^{gewoon} $\rightarrow$ bij omdat je nooit in een loop reusht kan komen $\rightarrow$ extra geheugen

Als we min acyclische graaf hebben (kunnen mogelijk de toegelaten wcht bogen)

$\rightarrow$ wel met neg. gewichten werken met Dijkstra

Werkisch pad = kortste pad van begin tot einde

hier verknijping $\rightarrow$ overal verknijping

voorbeeld wielhoes: kortste pad voor beste wielhoes is 2 munten

beste pad alg. werken echter met optellen

en wielhoes met vermengvuldigen

$\rightarrow$ algoritme vid gewichten nemen

want min som leg $\approx$ min vermengvuldiging

Substring search

→ Wilt van algoritme die gegeven een zeer lange string (tekst) van lengte N en bepaald patroon van lengte M proberen te vinden met $M \leq N$

substring = stukje vlt tekst

3 + 1 mogelijke algoritmen

0) Brute force

elke karakter in patroon vergelijken met een karakter in je tekst, dan er het volgende era.

public static int search(String pat, String txt)

{ int M = pat.length();

int N = txt.length();

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

for (j = 0; j < M; j++)

{ if (txt.charAt(i+j) != pat.charAt(j)) break; }

if (j == M) return i; }

return N; }

} patroon op gelijke hoogte zetten met i vlt tekst
volg lang karakteren overeen komen volgend vlt'en stop wenn mis → i++ en opnieuw

j = pointer

↳ kan enorm veel werk zijn als we telkens op hetzelfde een mismatch hebben of als tekst en patroon beide repetitief zijn

want vaak telkens M testen op elke positie vlt patroon

$$M(N - M + 1) = MN - MM + M \sim M \cdot N \text{ als } M \ll N$$

testpos. patroon

karakteren mogelijkheden

ander middel brute force: je moet minstens een deel van M lang vlt tekst in geheugen houden voor als je merkt dat je terug naar startpos + 1

→ ~M extra geheugen

alternatieve implementatie - met back-up

public static int search(String pat, String txt)

{ int i, N = txt.length();

int j, M = pat.length();

for (i = 0; i < N && j < M; i++) {

if (txt.charAt(i) == pat.charAt(j)) j++;

else { i -= j; j = 0; }

if (j == M) return i - M;

else return N; }

→ i is index vlt al gematchte vlt tekst

→ j is # al gematchte karakteren

→ back-up mechanisme

1) Knuth-Morris-Pratt (KMP)

gebruik maken van kennis over hoe ver gevonden

+ " " gefaalde match (mismatch)

+ " " structuur vlt patroon

→ DETERMINISTISCHE

→ patroon wiskundig naar voren als eindige toestandsmachine (DFA)

→ tekst gewoon als input over DFA laten lopen en van toestand naar toestand springen totdat we in een aanvaardbare eindtoestand terecht komen of we daar heel de tekst zijn gelopen (verworpen)

↳ overeenkomst met patroon → verder springen

mismatch → terug springen: niet telkens terug naar begin toestand (nultoestand DFA)

vanafte # karakteren vlt pat[i] (van pat[i]) dat

overeenkomst met de suffix (van txt[i])

→ naar die toestand DFA springen

toestand DFA komt overeen met # kar. al gematcht

back-up nodig

back-up voor

DFA komt in een computer voor als een 2D-array

kolom $\rightarrow$ ~~toestanden waaraan je terecht kunt komen~~

rij $\rightarrow$ mogelijke input (char.)

kolom $\rightarrow$ plaats in patroon

} toestand waaraan je terecht komt

public int search (String txt) {

int i, j, N = txt.length();

for (i = 0; j = 0; i < N && j < M; i++)

j = dfa [txt.charAt(i)][j]

} max N array accesses

if (j == M) return i - M;

else return N; }

gun backup nodig want moeten niet terug naar vorige j

moeten wel op voorhand de DFA hebben gemaakt / berekend

$\rightarrow$ DFA opstellen: tabel invullen

• match transition

stel in toets j (n j was al gemaakt) en next char $c \neq$ patn.charAt(j)
 $\rightarrow$ ga naar toets j+1

• mismatch transition

stel in toets j en next char $c \neq$ patn.charAt(j)

$\rightarrow$ de laatste j waarvan vld input waren patn[0..j-1] gevolgd door c
proberen deze nu door de DFA te laten lopen

want weten dat patn[0..j-1] c tot nu toe goed

recursief $\Rightarrow$ begint zich herhalen tot we een match tegenkomen

analyse KMP

- karakter accenser: machine construeren (elk kar. patroon accenser)
+ patroon door machine in tekst zoeken (lineair)

$\Rightarrow$ max $M + N$ char. accesses

- dfa [][] opstellen: tabel invullen = R (# kar. / grootte alfabet) $\times$ M (# toets / lengte patn.)
 $\rightarrow R \cdot M$ getallen berekenen

$\Rightarrow$ tijd $\propto$ geheugen $\sim R \cdot M$

2) Boyer-Moore

check het patroon van achter naar voor: komt het laatste (M) kar. vln patroon overeen met het kar. op pos M i/d tekst

- als mismatch

* als dat kar. overeenkomt met een kar. uit het patroon

$\rightarrow$ schuiven patroon zo op dat ze op dezelfde hoogte staan

* als dat kar. niet overeenkomt met een kar. uit het patroon (dubbele mismatch)

$\rightarrow$ schuiven het patroon M pos. vooruit

- als match

$\rightarrow$ vergelijk nu het voorlaatste karakter op dezelfde manier

totdat je van voor bent beland in je patroon of op het einde van je tekst zit

$\hookrightarrow$ gebruiken hiervoor i/e gegevenstructuur

nodig voor offset right[c] $\rightarrow$ geeft meest rechtse voorkomen van kar. c i/h patroon terug
en -1 als het niet i/h patroon zit

positie steeds updaten als je het karakter i/h patn. tegenkomt

public int search (String txt) {

int N = txt.length(), M = patn.length();

int skip;

for (int i = 0; i <= N - M; i += skip) {

skip = 0;

for (int j = M - 1; j >= 0; j--) {

if (patn.charAt(j) != txt.charAt(i+j)) {

skip = Math.max(1, j - right[txt.charAt(i+j)]);

break; }

match $\rightarrow$ if (skip == 0) return i;

return N;

$\hookrightarrow$ gemiddeld karakter uit de tekst om ervoor te zorgen dat we NOOIT terugspelingen (altijd positieve skip)

Analyse Boyer-Moore

Worst case: telkens maar 1 pos opschuiven $\rightarrow$ te lang in brute force situatie

$$\sim M \cdot N$$

Best case: telkens M pos hu opschuiven

$$\sim N/M$$

L kan combineren met KMP of andere hybrides maken om repetitie te vermijden

$\rightarrow$ worst case $\sim N$

3) Rabin-Karp

int karakter te vergelijken $\rightarrow$ hashcodes gebruiken

$\rightarrow$ modulaire hashen

zeg $R = \#$ karakter in alfabet

$\Rightarrow$ patroon = geheel i.e. geheel systeem met basis R

van dit patroon de hash berekenen $\rightarrow 2 \ 3 \ 5 \ 5 \ 2 \ 9 \ 7 = 613$

voor elke i ($i = 0, \dots, N-M$) berekenen we de hash van

het geheel gevormd door de kar. i t.e.m. $i+M-1$ in de tekst

als nu hash patroon = hash tekst substring

$\rightarrow$ mogelijke match, moeten checken of dit ook effectief een match is

$\rightarrow$ voordeel aan modulaire hashen $\rightarrow$ rekenregels modulo rekenen

$\Rightarrow$ incrementeel hash berekenen: geleidelijk hash updaten telkens dat we 1 stap opschuiven

$$(a+b) \bmod a = [(a \bmod a) + (b \bmod a)] \bmod a$$

$$(a \cdot b) \bmod a = [(a \bmod a) \cdot (b \bmod a)] \bmod a$$

aan de hand hiervan de hash v/patroon berekenen

erst voor het 1^{ste} kar, 2^{de} kar enzj, 3^{de} kar...

in de hash v/d substring (id tekst

voor de eerste M : hash opbouwen zoals hash patroon

volgens: steeds 1 kar vooraan weg x achteraan enzj

$\rightarrow$ stel het H-digst getal met basis R als volgt voor

$$x_i = t_i R^{M-1} + t_{i+1} R^{M-2} + \dots + t_{i+M-1} R^0 \quad \text{met } t_i = \text{txt.charAt}(i)$$

$$\text{dan } x_{i+1} = t_{i+1} R^{M-1} + t_{i+2} R^{M-2} + \dots + t_{i+M} R^0$$

$$\Rightarrow x_{i+1} = x_i \cdot R - t_i R^M + t_{i+M}$$

shift left $\uparrow$ $\uparrow$ $\uparrow$
 push naar achteraan $\uparrow$ $\uparrow$
 add new rightmost digit

$$(-a) \bmod a$$

$$= a - a$$

$$\Rightarrow x_{i+1} \bmod a = [x_i \bmod a + t_i \bmod a \cdot (-R^M \bmod a)] \cdot (R \bmod a)^M + t_{i+M} \bmod a$$

x vooronderstel dat $R < a, t_i < a \Rightarrow R \bmod a = R$ x $t_i \bmod a = t_i$

2 varianten Rabin

• Las Vegas variant

als de hash gelijk is $\rightarrow$ moet patroon karakter per karakter nog vergelijken met de substring

100% correct

heel waarschijnlijk lineaire tijd $\sim N$

worst case zeker $\sim M \cdot N$

• Monte Carlo variant

gaan ervan uit dat als hash gelijk $\rightarrow$ patroon v/substring ook gelijk
 kans op mismatch = $1/a$ $\rightarrow$ geen 100% zekerheid, maar
 zekerheid groter want a groter

altijd in lineaire tijd $\sim N$

! modulo berekenen traag dan één karakter compare

maar nu wel meerdere pati herkennen

hash berekenen
en vergelijken
beter en sneller
dan het patroon
karakter per
karakter met de
tekst te vergelijken

Dynamic programming

Slides Dynamic programming

- manier van algoritmes implementeren / ontwerpen
- = look-up tabel (berekeningen in tabellen simuleren, zie Excel)
- informatie van algoritme tijdelijk opslaan (tussentijdse resultaten tijdelijk opslaan) om dat later te hergebruiken om complexe recursieproblemen te vermijden

basistheorie

probleem opdelen in subproblemen

de optimale opt. voor deze subproblemen → optimale opt voor het globale probleem
deze subproblemen delen vaak dezelfde subsubproblemen, enz.

↳ normaal: recursie → opdelen tot in h.v. gewal → sub... subproblemen

maar echter: veel van die subsubproblemen zijn identiek

⇒ ipv steeds hetzelfde op te lossen, gaan we die nu 1 keer oplossen en het resultaat in een tabel opslaan om te kunnen hergebruiken

→ sub... subprobleem dat we herkennen oplossen i/d tabel

⇒ opt. blijft hetzelfde maar efficiënter op tijdscomplexiteit vlak

vb: rij van Fibonacci

recursief programma

```
public static int fibonacci(int n) {
```

```
    int fib;
```

```
    if (n == 0) fib = 0;
```

```
    else if (n == 1) fib = 1;
```

```
    else fib = fibonacci(n-1) + fibonacci(n-2);
```

```
    return fib; }
```

$$\begin{cases} T(0) = 1 \\ T(1) = 1 \\ T(n) = T(n-1) + T(n-2) \\ \Rightarrow T(n) > 2T(n-2) = 2^{n/2} \end{cases}$$

↳ slecht, maar heel tijdsineff. want doen veel dubbel werk (met name hoge n)

⇒ beter: bottom-up

```
public static int fibonacci(int n) {
```

```
    int[] fib = new int[n];
```

```
    fib[0] = 1; fib[1] = 1;
```

```
    for (i = 2; i <= n; i++)
```

```
        fib[i] = fib[i-1] + fib[i-2];
```

```
    return fib[n]; }
```

↳ lineaire tabel die we maar per plaats invullen
→ maken i elementen berekenen
⇒ lineaire tijd, nn

kunnen recursieve oproepen die deelt us op te slaan en vervolgens dubbel te laten afhangen

vb: rod cutting: staaf in stukken snijden (snoegen van 1 blokje) mer afn. vln vld. vlnge
blokjes waaruit het stuk bestaat, heeft elk stuk een bep. waarde
willen de waarde vld staaf totaal zo hoog mogelijk hebben

naïeve opt.: staaf in 2 delen, alle n-1 mogelijke versnoeringen uitvoerend
en dan recursief het overgebleven stuk weer volgens alle mogelijke versnoeringen in 2 delen → zo recursief verder doen
↳ echter 2^{n-1} mogelijkheden die, vooral naarmate de staaf korter w., steeds beginnen overlappen met andere situaties
→ exponentieel

top-down: tabel bijhouden die de hoof waarde aangaat voor het versnoeden vln staaf van bep. lengte
1e keer staaf van die lengte? → recursief berekenen zoals hierboven, dan resultaat opslaan i/d tabel om later gewoon te kunnen opzoeken
→ kwadratisch

bottom-up: gelijkwaardig aan fibonacci
erst waarde voor staaf van lengte 1 berekenen, dan de handel daarvoor staaf van lengte 2 → herhalen opbouwen tot lengte n
↳ zelf ingrijpen, geen recursie, gewoon invullen
→ kwadratisch

vb rij munten (even aantal) met elk een waarde
speler mogen afwisselend een munt nemen, maar telkens enkel de linker-
rechtse of linkerse munt (2 spelers)
speler die op het einde de hoogste waarde heeft, wint
→ greedy algoritme? Neen - kan dat je zo een munt met hoge waarde voor je
tegenstander nagaat

→ opdelen in subproblemen

- nummer alle munten van 1 t.e.m. n
- houd de waarde bij i/c array: $V[1], \dots, V[2], \dots, \dots, V[n]$

→ voor een gegeven rij van munt i t.e.m. munt j

liest speler 1 ofwel $V[i]$ ofwel $V[j]$

wat is nu de maximale waarde die je uit de reeks i...j kan halen
als we het spel (met 2 spelers die elk optimaal spelen) uitspelen? M_{ij}

$M_{i,i+1} = \max(V[i], V[i+1])$ voor $j = i+1$: 2 munten over → hier de grootste van $V[i]$ en $V[j]$

voor $j > i+1$: meer dan 2 munten over

$$M_{ij} = \max \left\{ \min(M_{i+1,j-1}, M_{i+2,j}) + V[i], \min(M_{i,j-2}, M_{i+1,j-1}) + V[j] \right\}$$

↓ speler 1 kiest munt i
in speler 2 dan j of i+1
max want wil optimale voor jezelf
↓ speler 1 kiest munt j
in speler 2 j-1 of i
min want speler 2 wilt ook optimale voor zichzelf

→ tabel opstellen voor $M_{i,j}$ uit te vullen

aangaven van # munten → $j-i+1$ in een

- introductie: $j = i+1 \rightarrow j-i+1 = 2$

invullen met $\max(V[i], V[i+1])$

aangaven $i < j \rightarrow$ moeten enkel bovenste helft invullen

- vul $M_{i,j}$ in voor alle $j-i+1 = 4$

- " " " " " " $j-i+1 = 6$

→ herhalen tot $M_{1,n}$ bereikt

opvullen ~ kwadratische tijd vs exponentieel van pure recursief uitwerken
~ $n^2/4$ want moeten door dynamische progr. maar dan invullen

Longest Common Subsequence (LCS)

hoe gelijkaardig zijn 2 (lange) strings van karakters?

→ gemeenschappelijke subsequence = reekjes van karakters in de strings

(≠ substring!) die op elkaar volgen (zelfde volgorde, relatief)

onderlinge volgorde, maar niet noodzakelijk exact achter elkaar staan
→ $X[i_1]X[i_2] \dots X[i_k], Y[j_1]Y[j_2] \dots Y[j_k]$ met $j_1 < j_2 < \dots < j_k$ $j=1, \dots, k-1$

→ langste gemeenschappelijke subsequence

- naïeve opt: alle mogelijke subseq van string X (lengte n) opnoemen

→ 2^n mogelijke subseq

die nachecken of ze ook een subseq zijn van string Y (lengte m)

→ ~ m tijd voor 1 subseq te checken (lineair door string lopen)

→ totaal ~ $2^n m$ (exp)

+ heel veel overlappende problemen want veel van de 2^n subseq overlappen

→ dynamic programming

dynamische programmering om LCS op te lossen

string X met lengte n, string Y met lengte m

→ subprobleem: vangt mogelijke gemeensch. subseq in delen van X en Y

→ LCS van $X[0..i]$ en $Y[0..j] = L[i,j]$

→ voorstellen de matrix en uiteindelijk $L[m-1, n-1]$ berekenen

* Stel $X[i] = Y[j] = c$

⇒ LCS van $X[0..i]$ en $Y[0..j]$ eindigt met karakter c

Bewijs: uit het ongemak

Stel $L[i,j]$ eindigt op een ander karakter dan c, dan kunnen we $L[i,j]$ uitbreiden met c want X eindigde met ~~en~~ ^{in Y} → tegenspraak

Stel $L[i,j]$ eindigt met een c, maar die c in X en/of Y zit op een andere pos dan i,j. Dan kan we de pos. van c forceren naar de laatste pos van X en Y. □

⇒ $L[i,j] = L[i-1, j-1] + 1$ als $X[i] = Y[j]$

* Stel $X[i] \neq Y[j]$

⇒ LCS kan niet tegelijk op $X[i]$ en $Y[j]$ eindigen,

LCS kan op $X[i], Y[i]$ of op een van beiden eindigen

⇒ $L[i,j] = \max\{L[i-1, j], L[i, j-1]\}$ als $X[i] \neq Y[j]$

knippen in X

knippen in Y

→ zien welke vld & LCS nu nog langste is

* Randwaarden: $L[i, -1] = L[-1, j] = 0$

→ -1 wijst op lege string

→ Algoritme: tabel L opvullen beginnend bij $L[0,0]$ en zo verder werken totdat we $L[m-1, n-1]$ hebben bereikt

~ n.m

L -1 0 1 2 3 4 5 → j (Y)

-1	0	0	0	0	0	0
0	0	0	1	1	1	1
1	0	0	1	2	2	2
2	0	0	1	2	2	3
3	0	1	1	2	2	2

i (X)

LCS is niet uniek, maar de lengte ervan wel

○ duidt pos. aan waarvoor de eindwaarde hetzelfde zijn

1) rijen en kolommen met index -1 ingevuld met 0

2) invullen vlnz en vbnz

$X[0]Y[0] \rightarrow$ gelijk $\Rightarrow 1$

niet gelijk $\Rightarrow 0$

$X[0]Y[i]$ vullen aanvullen zoals hierboven $X[0]Y[m-1]$

zoals erom in Y: $Y[j] = X[0] \rightarrow$ verander naar 1 en blijft voor de rest ook staan want blijft in opeenvolgende Y zitten

3) gelijk (○) $\rightarrow$ diagonaal +1
versch $\rightarrow$ max (links, boven)

- Optimale binaire zoekbomen

≠ gebalanceerde binaire boom

voor een gegeven set $k_1 < k_2 < \dots < k_n$ van n gesorteerde stukjes (elke k_i met waaarsch. p_i)

Wat is nu de BST met minimum verwachte zoekkost?

$\text{cost}_{k_i} = \# \text{vgl'en nodig om } k_i \text{ te vinden}$

$= \text{diepte}(k_i) + 1$

→ wortel op diepte 0

↳ gemiddelde kost voor $k_i = \text{cost}(\text{vgl m frequentie/kans voorkomen})$

$\text{cost}_{k_i, \text{ama}} = (\text{diepte}(k_i) + 1) \cdot p_i$

⇒ totale kost vld zoekboom

$\text{cost}_{\text{BST}} = \sum_{i=1}^n (\text{diepte}(k_i) + 1) \cdot p_i$

$= \sum_{i=1}^n 1 + \sum_{i=1}^n \text{diepte}(k_i) \cdot p_i$

→ optimale zoekboom heeft kleinste kost

- niet per se kleinste lengte

- niet per se meest freq. kly als wortel

vinden dhr dynamische programmering

optimale zoekboom vinden voor $k_1, \dots, k_j$ ($1 \leq i \leq j \leq n$)

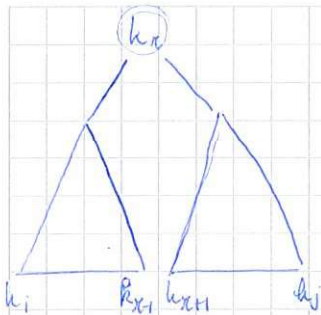
→ telkens uitbreiden door 2 subs berekende optimale zoekbomen te combineren aangezien voor optimale BST T geldt dat als T subboom T' met $k_1, \dots, k_j$ bevat, dan is T' de optimale BST voor die keys

→ optimale kost voor boom over $k_1, \dots, k_j = c(i,j)$

tabel invullen met als $i=j \Rightarrow c(i,j)=1$

$j < i \Rightarrow c(i,j)=0$

$j > i \Rightarrow c(i,j) = p_i + c(i, i-1) + w(i, i-1) + c(i+1, j) + w(i+1, j)$



$i \leq x \leq j$

stel dat h_x de wortel is van optimale BST voor $h_i, \dots, h_j$
dan hebben we links een subboom van $h_i, \dots, h_{x-1}$
en rechts van $h_{x+1}, \dots, h_j$
elk met hun eigen optimale wortel, $c(i, x-1)$ en $c(x+1, j)$
door het combineren van 2 bomen

- diepte van alle knopen in subboom met 1 verhoogt

→ kost subboom vermeerderd met de som van alle waarden in subboom, zeg w

$$\Rightarrow C(i, j) = p_x + C(i, x-1) + w(i, x-1) + C(x+1, j) + w(x+1, j)$$

↓
kost h_x

↓
som waarde linker subboom

$$\text{en aangegeven } w(i, j) = w(i, x-1) + p_x + w(x+1, j)$$

$$\Rightarrow C(i, j) = p_x + C(i, x-1) + w(i, x-1) + C(x+1, j) + w(x+1, j)$$

$$= C(i, x-1) + C(x+1, j) + w(i, j)$$

$$\rightarrow \text{optim: } C(i, j) = \min_{i \leq x \leq j} \{ C(i, x-1) + C(x+1, j) + w(i, j) \}$$

moeten 3 2D-arrayen berekenen

$C[1 \dots n, 1 \dots n]$ → kost opslaan

$w[1 \dots n, 1 \dots n]$ → w niet telkens opnieuw willen berekenen $w[i, j] = w[i, j-1] + p_j$

$root[i, j]$ → optimale x voor subboom $i \dots j$

→ onderstructuur: $C[1, n]$ en $root[1, n]$