
Algorithmen und Datenstrukturen

Prof. Dr. Ralf Möller

Universität zu Lübeck

Institut für Informationssysteme

Tanya Braun (Übungen)

sowie viele Tutoren

Danksagung

Einige der nachfolgenden Präsentationen wurden mit ausdrücklicher Erlaubnis des Autors und mit umfangreichen Änderungen und Ergänzungen übernommen aus:

- „Effiziente Algorithmen und Datenstrukturen“ (Kapitel 4: Hashing)
gehalten von Christian Scheideler an der TUM
<http://www14.in.tum.de/lehre/2008WS/ea/index.html.de>
- „Algorithmen und Datenstrukturen“
gehalten von Sven Groppe an der UzL

Wörterbuch-Datenstruktur

S: Menge von Schlüssel-Wert-Paaren (**key**, **object**)

Operationen:

- **insert**(**k**, **e**, **s**): $s := s \cup \{(k, e)\}$
// Änderung nach außen sichtbar
- **delete**(**k**, **s**): $s := s \setminus \{(k, e)\}$, wobei **e** das Element ist, das unter dem Schlüssel **k** eingetragen ist
// Änderung von **s** nach außen sichtbar
- **lookup**(**k**, **s**): Falls es ein $(k, e) \in S$ gibt, dann gib **e** aus, sonst gib $\perp$ aus

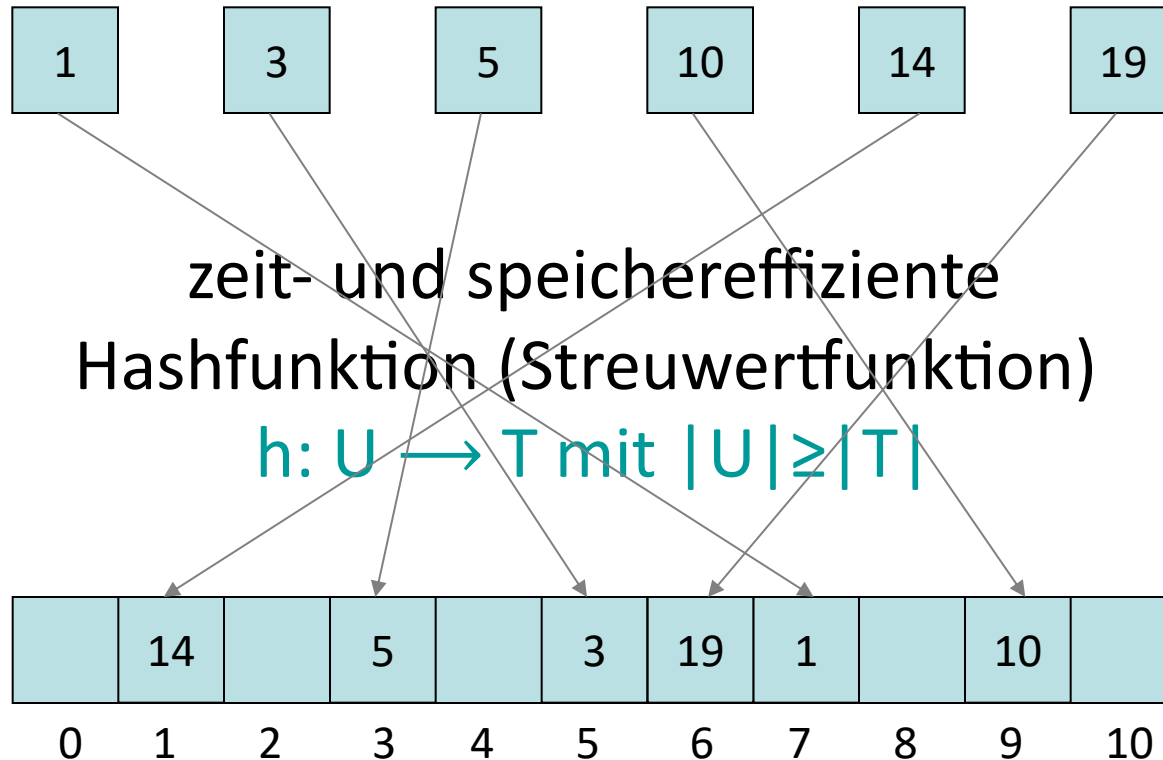
Wörterbücher

Unterschied zur Menge:

- Elemente (Einträge) bei Wörterbüchern als Attribut-Wert-Paare verstanden
- Iteration über Elemente eines Wörterbuchs in **willkürlicher** Reihenfolge **ohne** Angabe eines Bereichs
 - Über alle Attribute
 - Über alle Werte
 - Über alle Attribut-Wert-Paare

Hashing (Streuung)

Einige Elemente
aus einer Menge U :



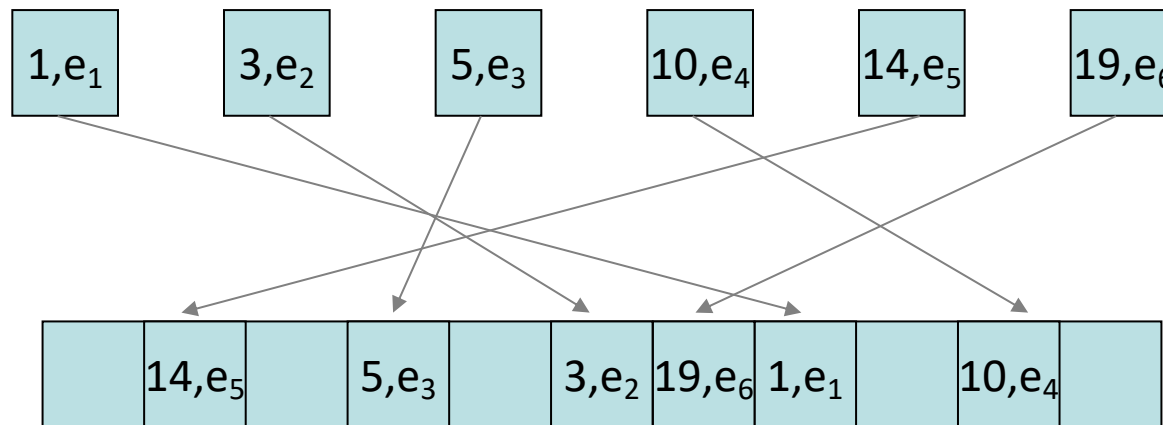
zeit- und speichereffiziente
Hashfunktion (Streuwertfunktion)

$$h: U \rightarrow T \text{ mit } |U| \geq |T|$$

Hashtabelle T

Assoziation durch Hashing

- Schlüssel k seien hier Zahlen
- Assoziation von e mit k



- Schlüssel k selbst können auch Objekte sein, es muss nur eine Abbildung auf T definiert sein bzw. werden

Hashing (perfekte Streuung, kein Kollisionen)

procedure **insert**(k, e, s):

$T := ht(s)$

$T[h(k)] := (k, e)$

procedure **delete**(k, s):

$T := ht(s)$

$T[h(k)] := \perp$

function **lookup**(k, s):

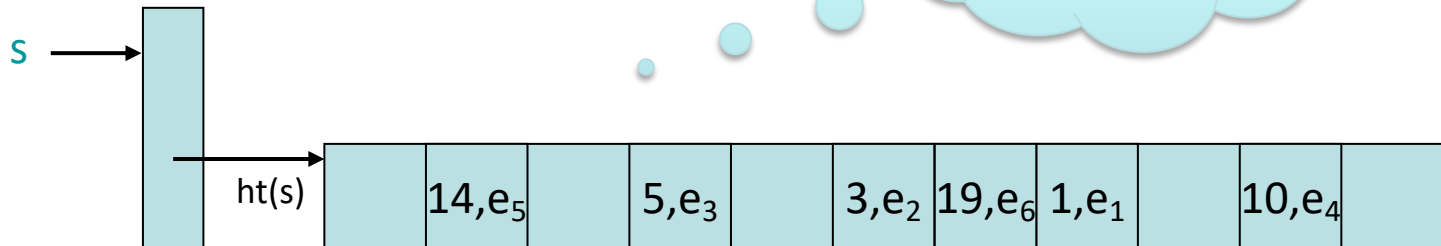
$T := ht(s); t := T[h(k)]$

return

if $t = \perp$ then $\perp$ else **second**(t)

Zu schön, um
wahr zu sein

Iteration?



Hashing: Übliches Anwendungsszenario

- Menge U der potentiellen Schlüssel u „groß“
- Anzahl der Feldelemente $\text{length}(T)$ „klein“
- D.h.: $|U| \gg \text{length}(T)$, aber nur „wenige“ $u \in U$ werden tatsächlich betrachtet
- Werte u können „groß“ sein (viele Bits)
 - Große Zahlen, Tupel mit vielen Komponenten, Bäume, ...
 - Eventuell nur Teile von u zur einfachen Bestimmung des Index für T betrachtet
 - Nur einige Zeichen einer Zeichenkette betrachtet
 - Bäume nur bis zu best. Tiefe betrachtet
 - Sonst Abbildungsvorgang h evtl. zu aufwendig
- Folge der großen Menge bzw. der teilweisen Betrachtung:
 - Verschiedene Elemente möglicherweise auf gleichen Index abgebildet (Kollision)

Hashfunktionen

- Hashfunktionen müssen i.A. anwendungsspezifisch definiert werden (oft für Basisdatentypen Standardimplementierungen angeboten)
- Hashwerte sollen möglichst gleichmäßig gestreut werden (sonst Kollisionen vorprogrammiert)

- Ein erstes Beispiel für $U = \text{Integer}$:

```
function h(u)  
    return u mod m  
wobei m=length(T)
```

Falls m keine Primzahl:
Schlüssel seien alle
Vielfache von 10 und
Tabellengröße sei 100
→ Viele Kollisionen

Warum i.A.
nicht?

- Kann man h auf komplexen Objekten über deren „Adresse“ realisieren?

Hashing zur Assoziation und zum Suchen

Analyse bei perfekter Streuung

- insert: $O(f(u, h))$ mit $f(u, h) = 1$
für $u \in \text{Integer}$ und h hinreichend einfach
- delete: $O(f(u, h))$ dito
- lookup: $O(f(u, h))$ dito

Problem: perfekte Streuung
Sogar ein Problem: gute Streuung

Fälle:

- Statisches Wörterbuch: nur lookup
- Dynamisches Wörterbuch: insert, delete und lookup

Hashfunktionen

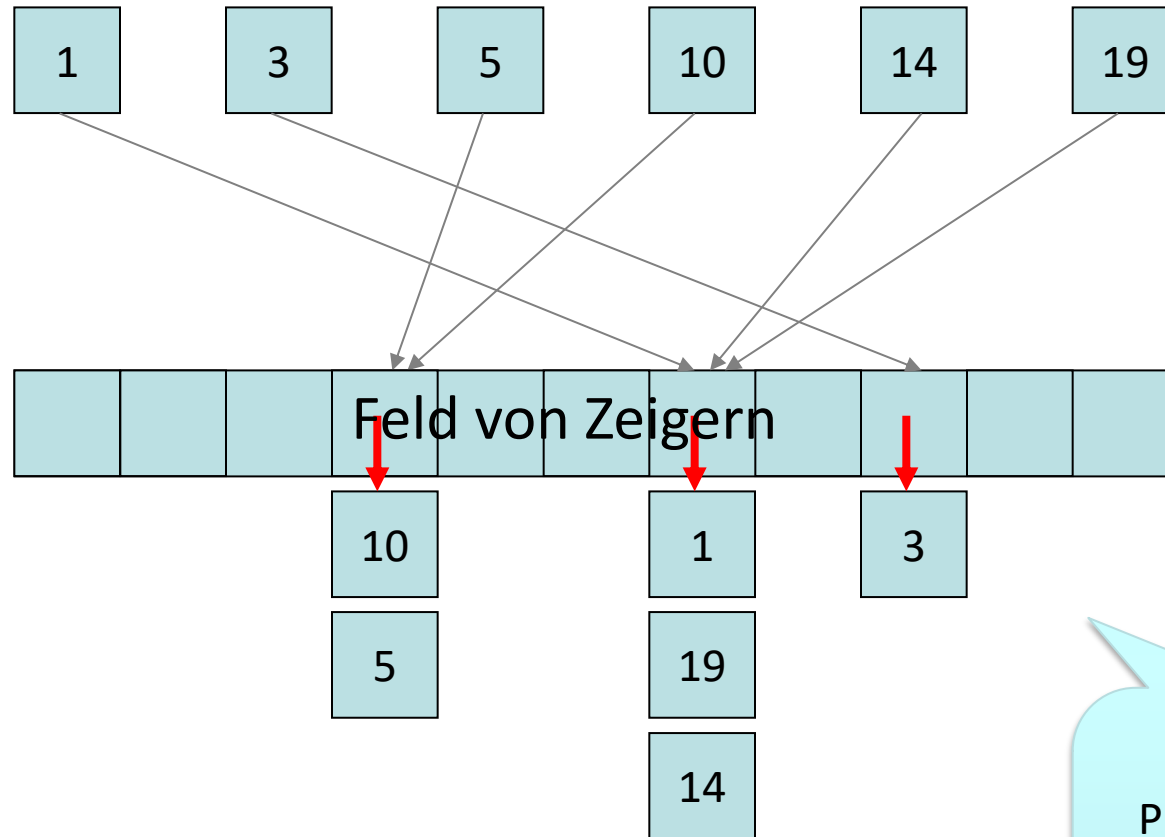
- Beim Erzeugen einer Hashtabelle kann man die initiale Größe angeben
- Man möchte den Nutzer nicht zwingen, eine Primzahl zu verwenden
- Andere Hashfunktion für $U = \text{Integer}$:

function $h(u)$

return $(u \bmod p) \bmod m$

wobei $p > m$ eine „interne“ Primzahl und $m = \text{length}(T)$
nicht notwendigerweise prim

Hashing mit Verkettung¹ (Kollisionslisten)



unsortierte verkettete Listen

Vereinfachte
Präsentation der
Tupel
(nur Schlüssel
dargestellt)

¹ Auch geschlossene Adressierung genannt.

Hashing mit Verkettung

T: Array $[0..m-1]$ of List

procedure **insert**(k, e, s):

$T := ht(s)$; insert_in_list((k, e), $T[h(k)]$)

procedure **delete**(k, s):

$T := ht(s)$; delete_from_list((k, e), $T[h(k)]$)

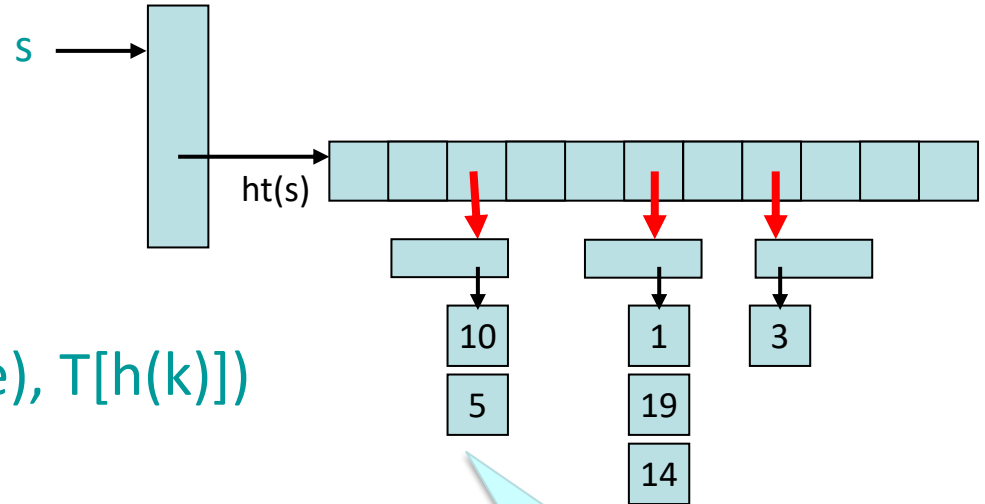
function **lookup**(k, s):

$T := ht(s)$

 for $(k', e) \in T[h(k)]$ do

 if $k=k'$ then return e

 return $\perp$

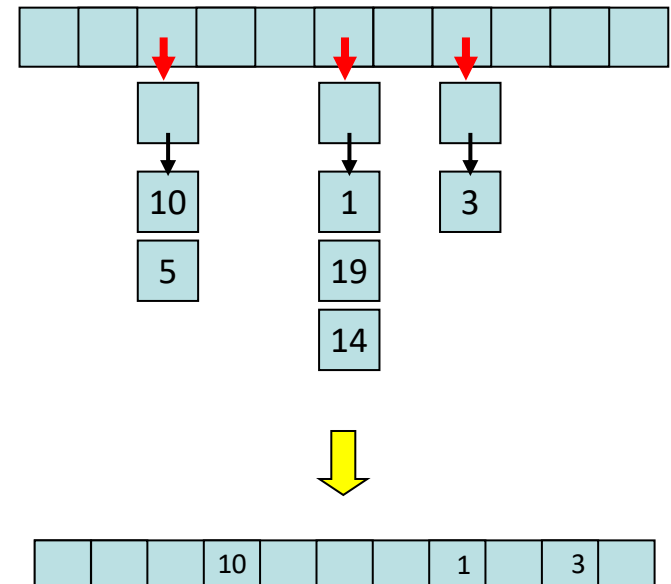


Vereinfachte
Präsentation der
Tupel
(nur Schlüssel
dargestellt)

Analyse der Komplexität bei Verkettung

- Sei α die durchschnittliche Länge der Listen, dann
 - $\Theta(1+\alpha)$ für
 - erfolglose Suche und
 - Einfügen (erfordert Überprüfung, ob Element schon eingefügt ist)
 - $O(1+\alpha)$ für erfolgreiche Suche

Kollisionslisten im Folgenden
nur durch das erste Element
direkt im Feld dargestellt



Dynamische Hashtabelle

Problem: Hashtabelle kann zu groß oder zu klein sein

Lösung: Reallokation

- Wähle neue geeignete Tabellengröße
- Wähle neue Hashfunktion
- Übertrage Elemente auf die neue Tabelle
 - Jeweils mit Anwendung der (neuen) Hashfunktion
 - In den folgenden Darstellung ist dieses nicht gezeigt!

Dynamische Hashtabelle

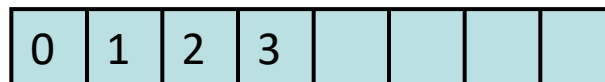
- Sei m die Größe des Feldes, n die Anzahl der Elemente
- Tabellenverdopplung ($n > m$):



- Tabellenhalbierung ($n \leq m/4$):



- Von



- Nächste Verdopplung: $> n$ insert Ops
- Nächste Halbierung: $> n/2$ delete Ops

Wegen Kollisionen
evtl. für Hashtabelle
schon ab $n > m/2$ nötig

Dynamische Hashtabelle

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

reallocate

$$\phi(s)=0$$

+

0	1	2	3	4			
---	---	---	---	---	--	--	--

insert

$$\phi(s)=2$$

0	1	2	3	4	5		
---	---	---	---	---	---	--	--

$$\phi(s)=4$$

0	1	2	3	4	5	6	
---	---	---	---	---	---	---	--

$$\phi(s)=6$$

0	1	2	3	4	5	6	7
---	---	---	---	---	---	---	---

$$\phi(s)=8$$

0	1	2	3	4	5	6	7				reallocate		
---	---	---	---	---	---	---	---	--	--	--	------------	--	--

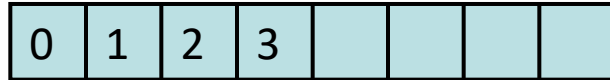
$$\phi(s)=0$$

+

0	1	2	3	4	5	6	7	8			insert		
---	---	---	---	---	---	---	---	---	--	--	--------	--	--

$$\phi(s)=2$$

Dynamische Hashtabelle



$$\phi(s)=0$$



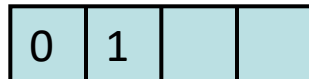
$$\phi(s)=2$$



delete

$$\phi(s)=4$$

+



reallocate

$$\phi(s)=0$$

Generelle Formel für $\phi(s)$:

(w_s : Feldgröße von s , n_s : Anzahl Einträge)

$$\phi(s) = 2 \lfloor w_s/2 - n_s \rfloor$$

Dynamische Hashtabelle

Generelle Formel für $\phi(s)$:

(w_s : Feldgröße von s , n_s : Anzahl Einträge)

$$\phi(s) = 2 \lfloor w_s/2 - n_s \rfloor$$

Behauptung:

Sei $\Delta\phi = \phi(s') - \phi(s)$ für $s \rightarrow s'$. Für die **amortisierten** Laufzeiten gilt:

- insert: $t_{\text{ins}} + \Delta\phi \in O(1)$
- delete: $t_{\text{del}} + \Delta\phi \in O(1)$

Dynamische Hashtabelle

Problem: Tabellengröße m sollte prim sein
(für gute Verteilung der Schlüssel)
Wie finden wir Primzahlen?

Lösung:

- Für jedes k gibt es Primzahl in $[k^3, (k+1)^3]$
- Wähle Primzahlen m , so dass $m \in [k^3, (k+1)^3]$
- Jede nichtprime Zahl in $[k^3, (k+1)^3]$ muss Teiler $< \sqrt{(k+1)^3}$ haben
→ erlaubt effiziente Primzahlfindung

Offene Adressierung

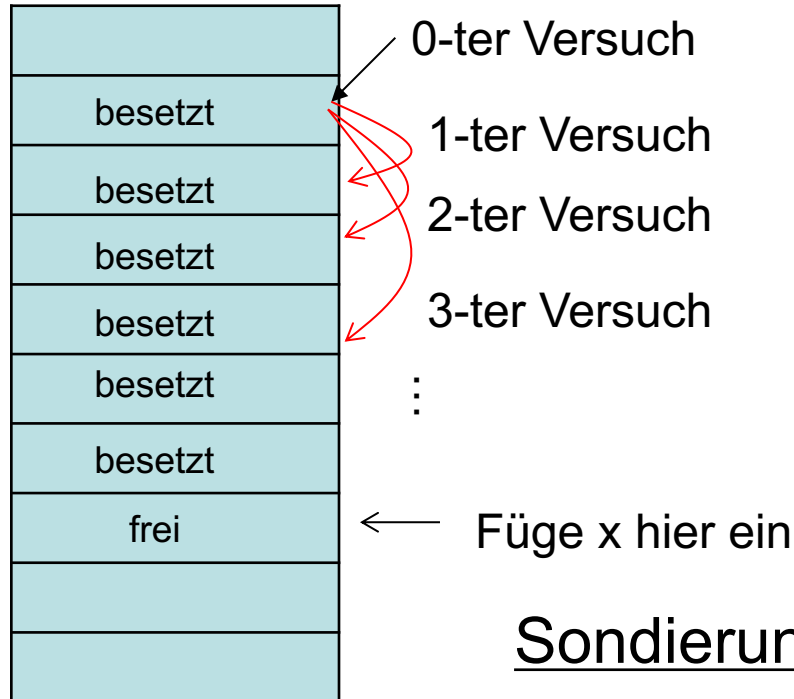
- Bei Kollision speichere das Element „woanders“ in der Hashtabelle
- Vorteile gegenüber Verkettung
 - Keine Verzeigerung
 - Schneller, da Speicherallokation für Zeiger relativ langsam
- Nachteile
 - Langsamer bei Einfügungen
 - Eventuell sind mehrere Versuche notwendig, bis ein freier Platz in der Hashtabelle gefunden worden ist
 - Tabelle muss größer sein (maximaler Füllfaktor kleiner) als bei Verkettung, um konstante Komplexität bei den Basisoperationen zu erreichen

Offene Adressierung

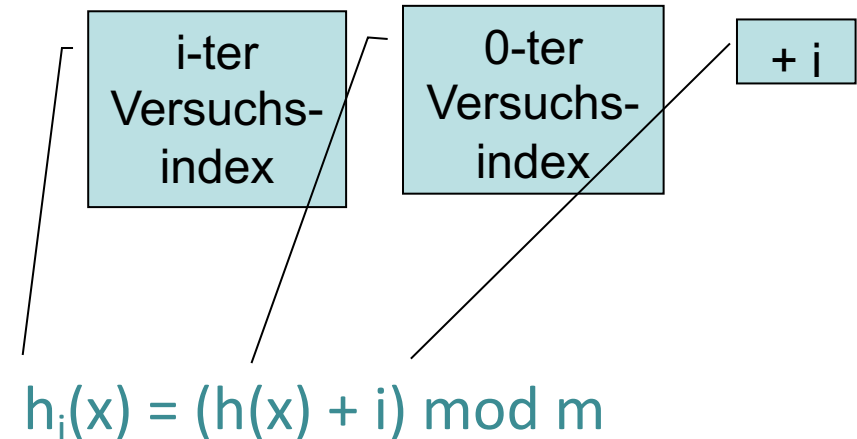
- Eine *Sondierungssequenz* ist eine Sequenz von Indizes in der Hashtabelle für die Suche nach einem Element
 - $h_0(x), h_1(x), \dots$
 - Sollte jeden Tabelleneintrag genau einmal besuchen
 - Sollte wiederholbar sein
 - so dass wir wiederfinden können, was wir eingefügt haben
- Hashfunktion
 - $h_i(x) = (h(x) + f(i)) \bmod m$
 - $f(0) = 0$ Position des 0-ten Versuches
 - $f(i)$ „Distanz des i-ten Versuches relativ zum 0-ten Versuch“

Einfügung von x: Lineares Sondieren

Linear probing:



- $f(i)$ ist eine lineare Funktion von i , z.B. $f(i) = i$

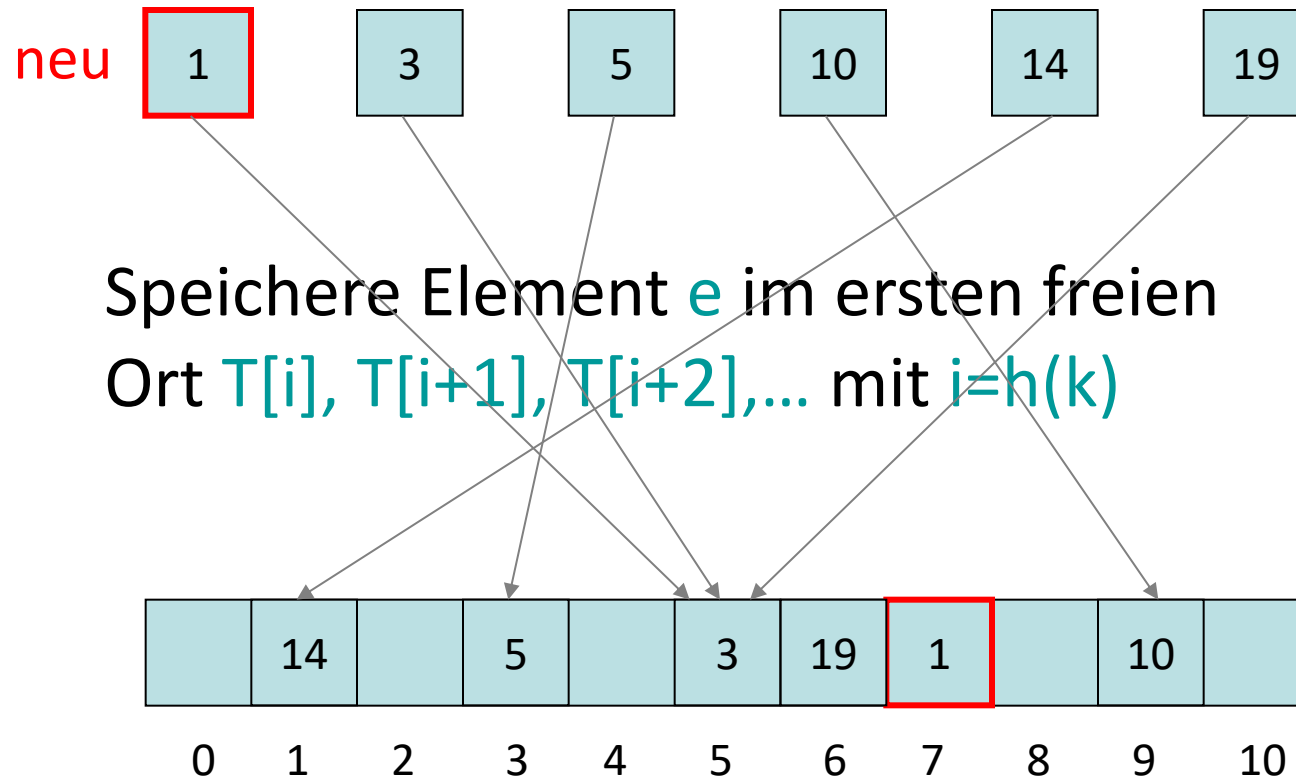


Sondierungssequenz: +0, +1, +2, +3, +4, ...

Fahre fort bis ein freier Platz gefunden ist

#fehlgeschlagene Versuche als eine Messgröße der Performanz

Hashing with Linearer Sondierung (Linear Probing)



Hashing with Linearer Sondierung

T: Array [0..m-1] of pairs (key, Element) // m>n

```
procedure insert(k, e, s)
  T := ht(s); i := h(k)
  while T[i] <> ⊥ & first(T[i]) <> k do
    i := (i+1) mod m
  T[i] := (k, e)
```

```
function lookup(k, s):
  i := h(k); T := ht(s)
  while T[i] <> ⊥ & first(T[i]) <> k do
    i := (i+1) mod m
  return second(T[i])
```

Hashing with Linearer Sondierung

Problem: Löschen von Elementen

Lösungen:

1. Verbiete Löschungen
2. Markiere Position als gelöscht mit speziellem Zeichen (ungleich $\perp$)
3. Stelle die folgende **Invariante** sicher:
Für jedes $e \in S$ mit idealer Position $i=h(k)$ und aktueller Position j gilt

$T[i], T[i+1], \dots, T[j]$ sind besetzt

Nachteile der Linearen Sondierung

- Sondierungssequenzen werden mit der Zeit länger
 - Schlüssel tendieren zur Häufung in einem Teil der Tabelle
 - Schlüssel, die in den Cluster gehasht werden, am Ende des Clusters gespeichert (→ vergrößern damit den Cluster)
 - Seiteneffekt
 - Andere Schlüssel sind auch betroffen, falls sie in die Nachbarschaft gehasht werden

Analyse der offenen Adressierung

- Sei $\alpha = n/m$ mit n Anzahl eingefügter Elemente und m Größe der Hashtabelle
 - α wird auch **Füllfaktor** der Hashtabelle genannt
- Anzustreben ist $\alpha \leq 1$
- Unterscheide erfolglose und erfolgreiche Suche

Analyse der erfolglosen Suche

Behauptung: Im typischen Fall $O(1/(1-\alpha))$

- Bei 50% Füllung ca. 2 Sondierungen nötig
- Bei 90% Füllung ca. 10 Sondierungen nötig

Ohne Beweis

Analyse der erfolgreichen Suche

Behauptung: Im durchschnittlichen Fall $O(1/\alpha \ln(1/(1-\alpha)))$

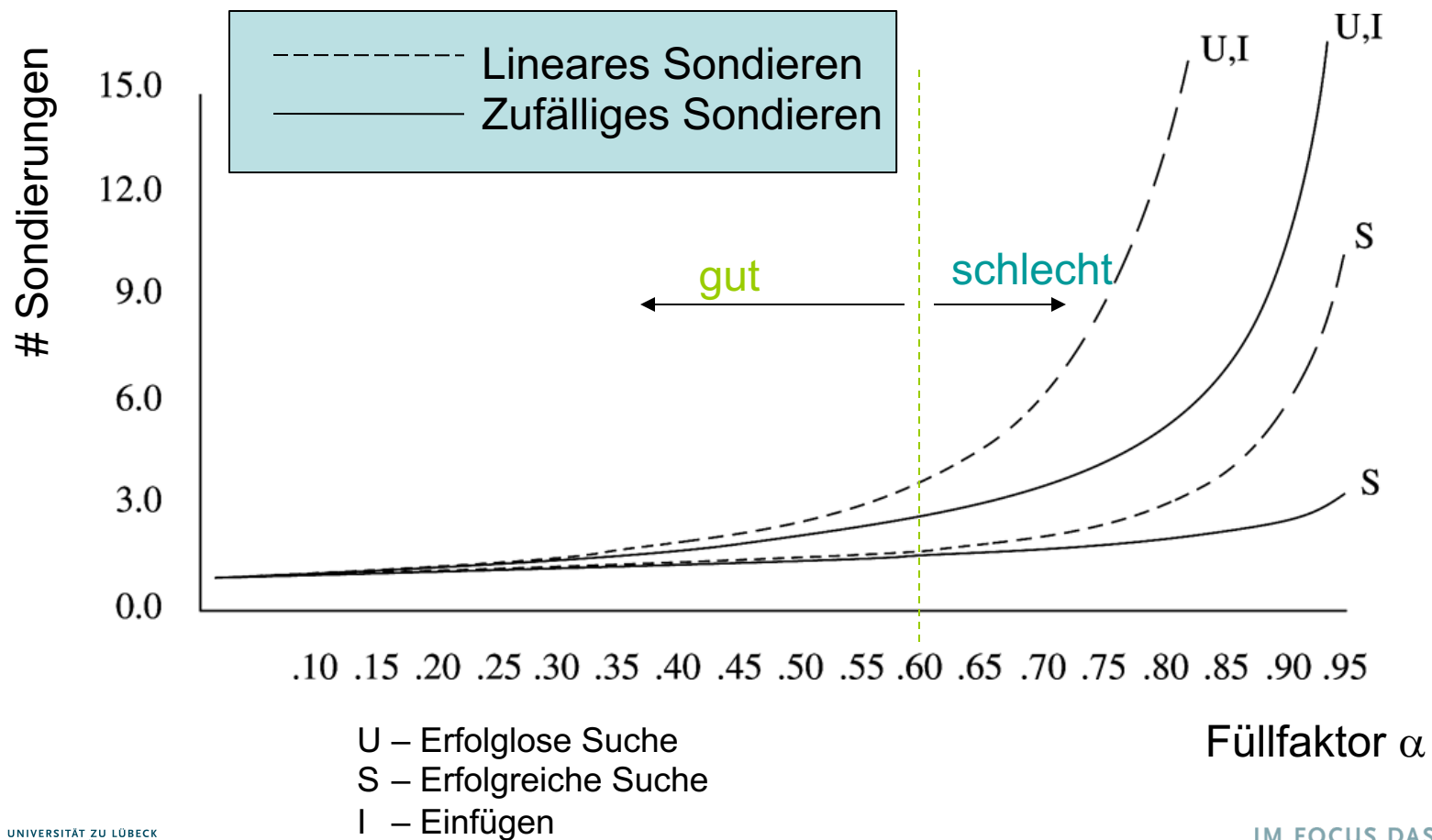
- Bei 50% Füllung ca. 1,39 Sondierungen nötig
- Bei 90% Füllung ca. 2,56 Sondierungen nötig

Ohne Beweis

Zufälliges Sondieren

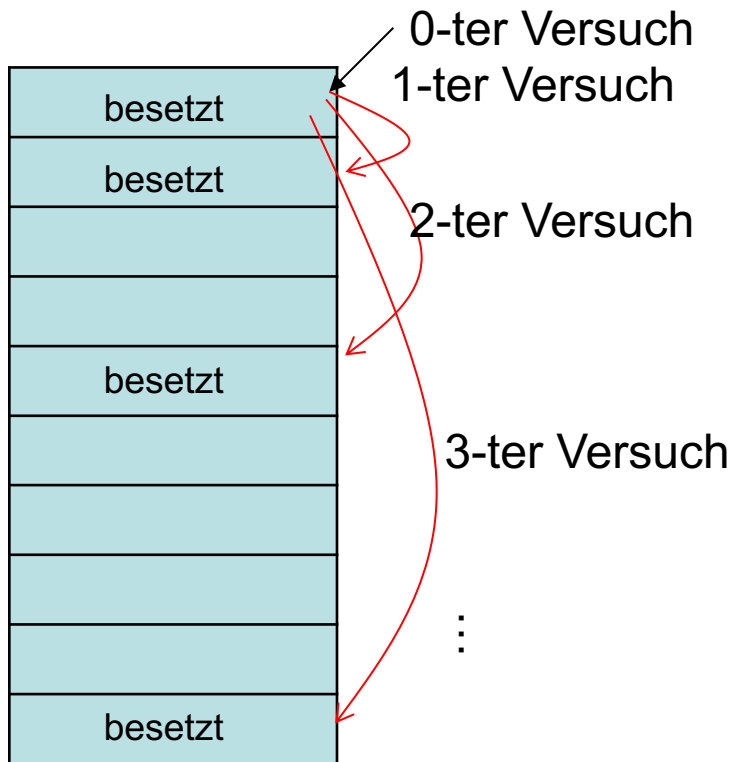
1. Wähle den jeweils nächsten Feldindex nach einer (reproduzierbaren) Zufallsfolge
 - Rechenaufwendig
2. Für jeden Schlüssel k wähle genügend lange zufällige Versatzfolge $f(i)$ und speichere Folge $f(i)$ zur Verwendung bei erneutem Hash von k
 - Speicheraufwendig
 - Bootstrap-Problem
 - Assoziation $\text{Key} \rightarrow \text{Indexfolge}$
 - Realisiert mittels Hashing?

Vergleich mit zufälligem Sondieren



Quadratisches Sondieren

Quadratisches Sondieren:



Fahre fort bis ein freier Platz gefunden ist

#fehlgeschlagene Versuche ist eine Meßgröße für Performanz

- Vermeidet primäres Clustering
- $f(i)$ ist quadratisch in i z.B., $f(i) = i^2$
 - $h_i(x) = (h(x) + i^2) \bmod m$
 - Sondierungssequenz:
+0, +1, +4, +9, +16, ...
 - Allgemeiner:
 $f(i) = c_1 \cdot i + c_2 \cdot i^2$

Löschen von Einträgen bei offener Adressierung

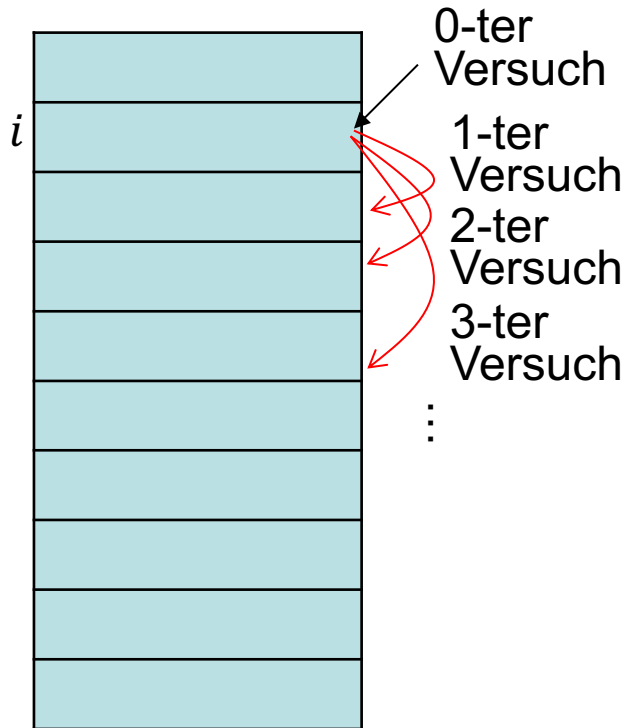
- Direktes Löschen unterbricht Sondierungskette
- Mögliche Lösung:
 - a) Spezieller Eintrag “gelöscht”. Kann zwar wieder belegt werden, unterbricht aber Sondierungsketten nicht.
Nachteil bei vielen Löschungen:
Lange Sondierungszeiten
 - b) Umorganisieren. Kompliziert, sowie hoher Aufwand

Analyse Quadratisches Sondieren

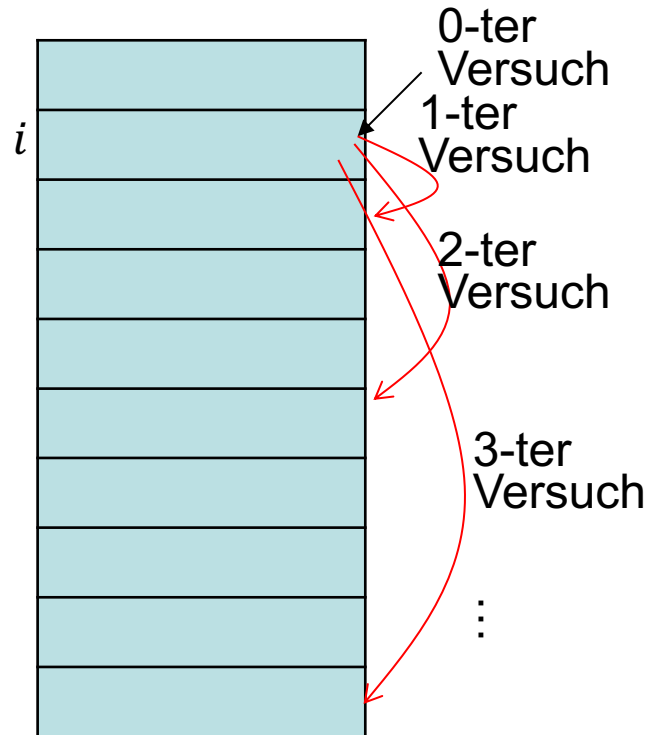
- Schwierig
- Theorem
 - Wenn die Tabellengröße eine Primzahl ist und der Füllfaktor höchstens $\frac{1}{2}$ ist, dann findet Quadratisches Sondieren immer einen freien Platz
 - Ansonsten kann es sein, dass Quadratisches Sondieren keinen freien Platz findet, obwohl vorhanden
- Damit $\alpha_{\max} \leq \frac{1}{2}$ für quadratisches Sondieren

Review Hashing

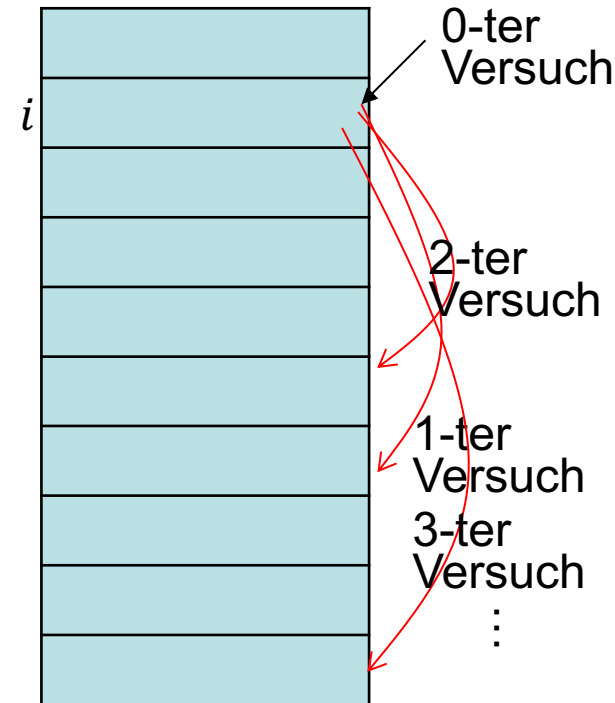
Lineares Sondieren:



Quadratisches Sondieren:



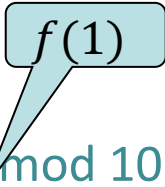
Doppel-Hashing*:



*(bestimmt mit einer zweiten Hashfunktion)

Doppel-Hashing

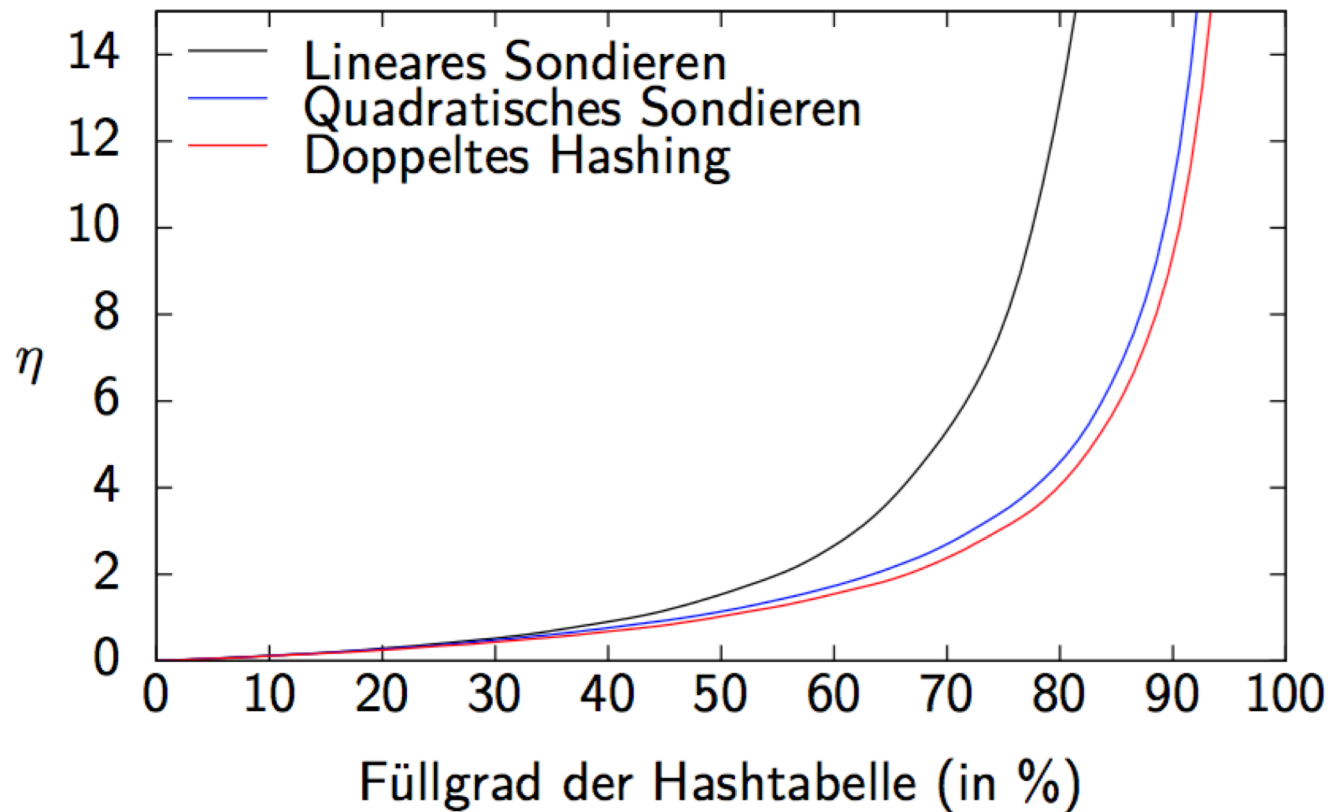
- Hashfunktion
 - $h_i(x) = (h(x) + f(i)) \bmod m$
 - $f(0) = 0$ Position des 0-ten Versuches
 - $f(i)$ „Distanz des i-ten Versuches relativ zum 0-ten Versuch“
- Benutze eine zweite Hashfunktion für alle Versuche außer dem ersten $f(i) = i \cdot h'(x)$
- Gute Wahl von h' ?
 - Sollte niemals 0 ergeben
 - $h'(x) = p - (x \bmod p)$ mit p Primzahl $< m$
- Beispiel für $m=10$.
 - $h_0(49) = (h(49) + f(0)) \bmod 10 = 9$
 - $h_1(49) = (h(49) + \underline{1 \cdot (7 - 49 \bmod 7)}) \bmod 10 = 6$ für $p=7$



$f(1)$

Praktische Effizienz von doppeltem Hashing

- ▶ Hashtabelle mit 538 051 Einträgen (Endfüllgrad 99,95%)
- ▶ *Mittlere* Anzahl Kollisionen η pro Einfügen in die Hashtabelle:

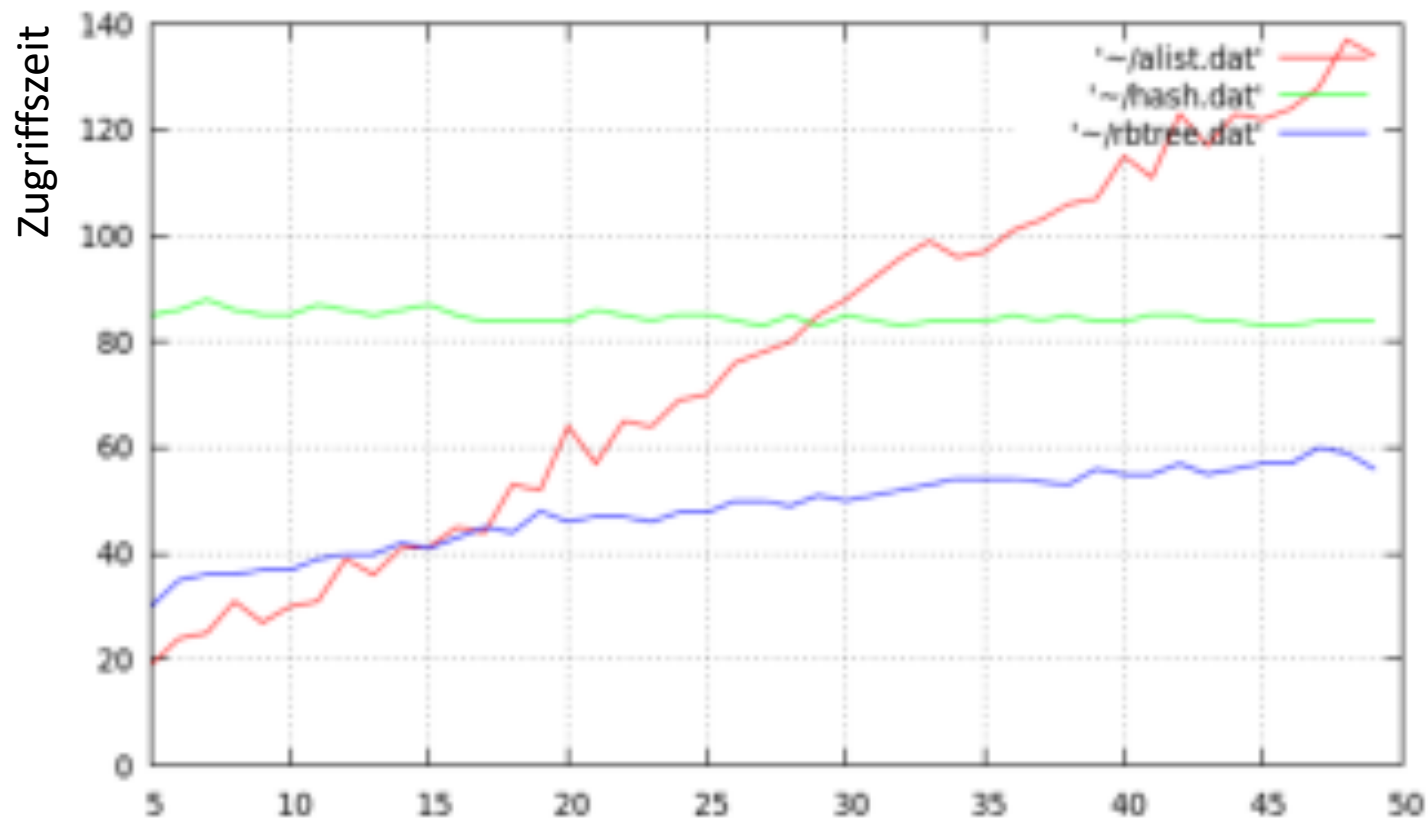


Analyse Hashing

- Messreihen zeigen, dass Doppel-Hashing fast genauso gut ist wie zufälliges Sondieren
 - Zweite Hashfunktion benötigt zusätzliche Zeit
- Doppeltes Hashing ist langsamer als Überlaufketten und lineares Sondieren bei dünn besetzten Tabellen, jedoch wesentlich schneller als lineares Sondieren, wenn der Füllgrad der Tabelle zunimmt
- Generell: Hashing ist bedeutend schneller als Suchen in Bäumen,
 - kann aber „ruckeln“ (durch Reallokation) und
 - unterstützt keine Bereichs-Iteration

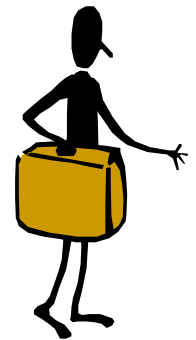
Vergleiche

- Schlüsselwortliste (rot) vs. Hashtabelle (grün) vs. Baum (blau)



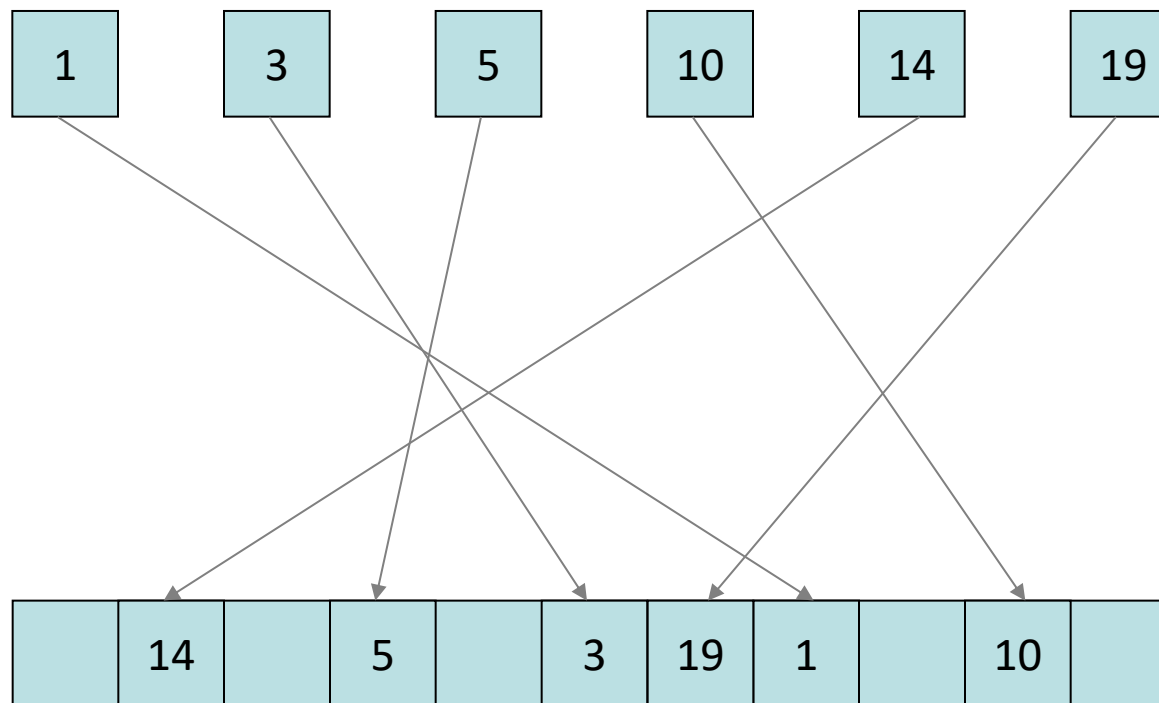
Zusammenfassung: Hashing

- Basisoperation (Suchen, Einfügen, Löschen) in $O(1)$
- Güte des Hashverfahrens beeinflusst durch
 - Hashfunktion
 - Verfahren zur Kollisionsbehandlung
 - Verkettung
 - Offene Adressierung
 - Lineares/Quadratisches Sondieren/Doppel-Hashing
 - Füllfaktor
 - Dynamisches Wachsen
- Statistisches vs. Dynamisches Hashen
- Universelles Hashing

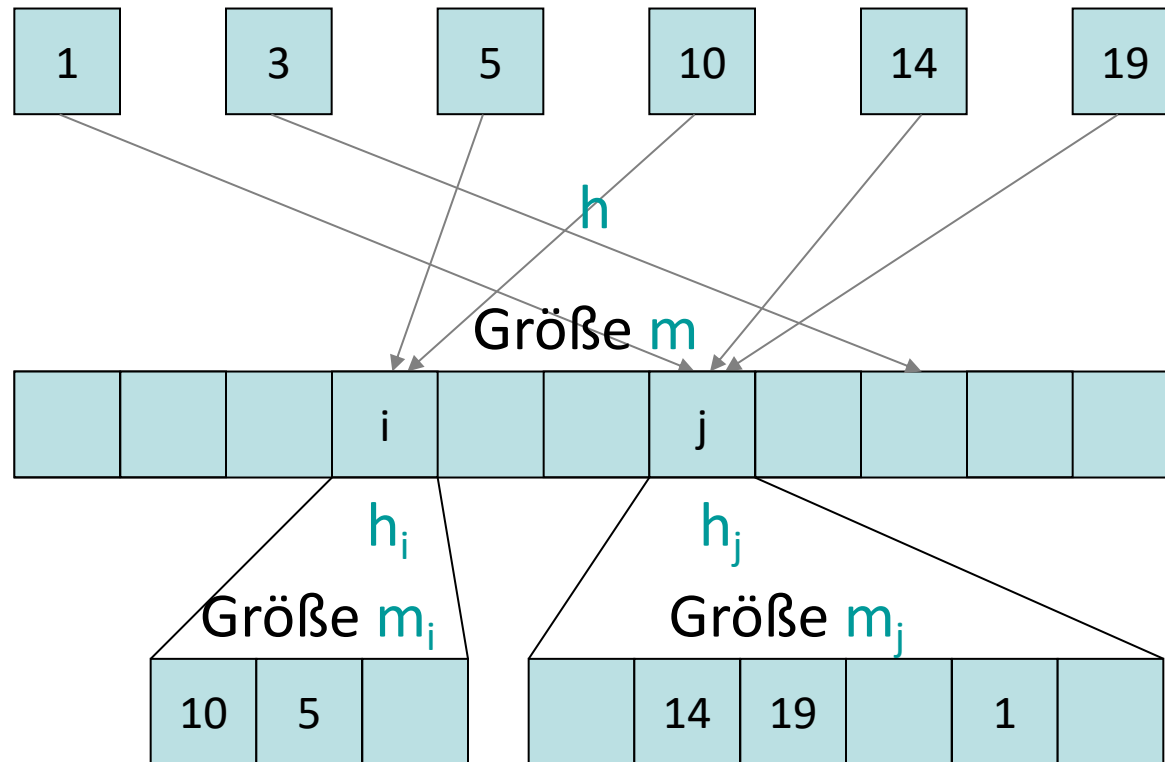


Statisches Wörterbuch

Ziel: perfekte Hashtabelle

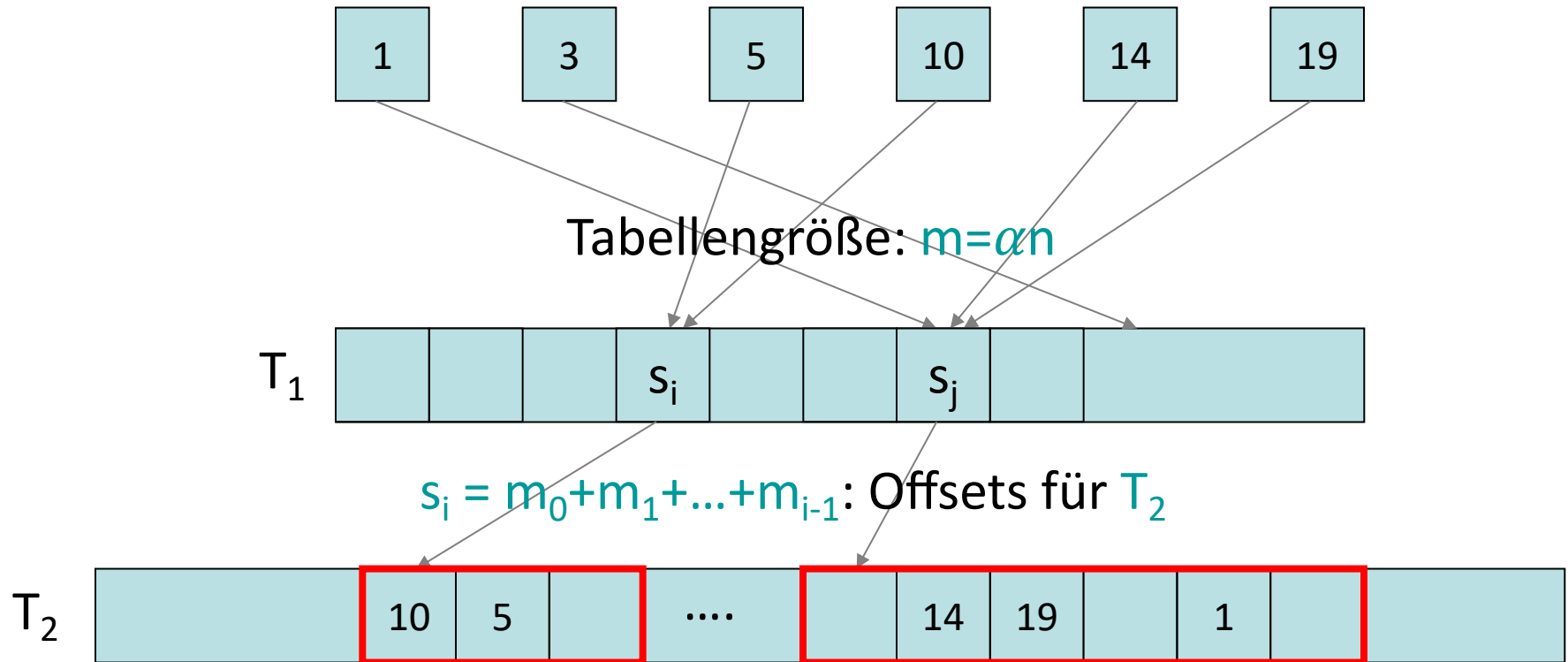


Statisches Wörterbuch (FKS-Hashing)



Keine Kollisionen in Subtabellen

Statisches Wörterbuch



Statisches Wörterbuch

Behauptung: Für jede Menge von n Schlüsseln gibt es eine perfekte Hashfunktion der Größe $\Theta(n)$, die in erwarteter Zeit $\Theta(n)$ konstruiert werden kann.

Sind perfekte Hashfunktionen auch dynamisch konstruierbar??

Hashing: Prüfsummen und Verschlüsselung

- Bei **Prüfsummen** verwendet man Hashwerte, um Übertragungsfehler zu erkennen
 - Bei guter Hashfunktion sind Kollisionen selten,
 - Änderung weniger Bits einer Nachricht (Übertragungsfehler) sollte mögl. anderen Hashwert zur Folge haben
- In der **Kryptologie** werden spezielle kryptologische Hashfunktionen verwendet, bei denen zusätzlich gefordert wird, dass es **praktisch unmöglich ist, Kollisionen absichtlich zu finden** (→ SHAx, MD5)
 - Inverse Funktion $h^{-1}: T \rightarrow U$ „schwer“ zu berechnen
 - Ausprobieren über $x=h(h^{-1}(x))$ ist „aufwendig“ da $|U|$ „groß“

Vermeidung schwieriger Eingaben

- **Annahme:** Pro Typ **nur eine Hash-Funktion** verwendet
- Wenn man Eingaben, die per Hashing verarbeitet werden, geschickt wählt, kann man **Kollisionen** durch geschickte Wahl der Eingaben **provozieren** (ohne gleiche Eingaben zu machen)
- **Problem:** Performanz sinkt (wird u.U. linear)
 - „Denial-of-Service“-Angriff möglich
- **Lösung:** Wähle Hashfunktion zufällig aus Menge von Hashfunktionen, die unabhängig von Schlüsseln sind
→ Universelles Hashing

Änderung der Hashfunktion bei Hashtabelle

- Hash-Funktion ändern beim Vergrößern oder Verkleinern einer Hashtabelle (Rehash)
- Messen der mittleren #Sondierungen und ggf. ein spontanes Rehash mit anderer Hash-Funktion
 - (latente Gefahr eines DOS-Angriffs abgemildert)
- Hierzu notwendig:
 - Auswahlmöglichkeit von h aus Menge von universell verwendbaren Hashfunktionen H

Universelles Hashing¹

- Eine Menge von Hashfunktionen H heißt universell, wenn für beliebig wählbare Schlüssel $x, y \in U$ mit $x \neq y$ gilt:

$$\frac{|\{h \in H \mid h(x) = h(y)\}|}{|H|} \leq 1/m$$

- Also: Wenn eine Hashfunktion h aus H zufällig gewählt wird, ist die relative Häufigkeit von Kollisionen kleiner als $1/m$, wobei m die Größe der Hashtabelle ist
- Beispiel: $h_{a,b}(k) = ((ak + b) \bmod p) \bmod m$
- Die Funktionen in H haben Parameter a, b
- Wir sprechen auch von einer Familie von Hashfunktionen

¹ Manchmal auch universales Hashing genannt: Für alle Schlüsselsequenzen geeignet, also universal einsetzbar

Universelles Hashing

Wir wählen eine Primzahl p , so dass jeder Schlüssel k kleiner als p ist.

$$Z_p = \{0, 1, \dots, p-1\}$$

$$Z_p^* = \{1, \dots, p-1\}$$

Da das Universum erheblich größer als die Tabelle T sein soll, muss gelten:

$$p > m$$

Für jedes Paar (a, b) von Zahlen mit $a \in Z_p^*$ und $b \in Z_p$ definieren wir wie folgt eine Hash-Funktion:

$$h_{a,b}(k) = ((ak + b) \bmod p) \bmod m$$

Beispiel:

$$p = 17 \quad m = 6$$

$$\longrightarrow h_{3,4}(8) = 5$$

Universelles Hashing

Beh: Die Klasse $H_{p,m} = \{h_{a,b} \mid a \in Z_p^* \wedge b \in Z_p\}$
mit $h_{a,b}(k) = ((ak + b) \bmod p) \bmod m$
von Hash-Funktionen ist universell.

Ohne Beweis

Carter, Larry; Wegman, Mark N. "Universal Classes of Hash Functions".
Journal of Computer and System Sciences. 18 (2): 143–154, 1979.

Zusammenfassung: Hashing

- Basisoperation (Suchen, Einfügen, Löschen) in $O(1)$
- Güte des Hashverfahrens beeinflusst durch
 - Hashfunktion
 - Verfahren zur Kollisionsbehandlung
 - Verkettung
 - Offene Adressierung
 - Lineares/Quadratisches Sondieren/Doppel-Hashing
 - Füllfaktor
 - Dynamisches Wachsen
- Statistisches vs. Dynamisches Hashen
- Universelles Hashing

