



LUDWIG-
MAXIMILIANS-
UNIVERSITÄT
MÜNCHEN

FAKULTÄT FÜR MATHEMATIK, INFORMATIK UND STATISTIK
INSTITUT FÜR INFORMATIK



Skriptum zur Vorlesung Algorithmen und Sequenzen

gehalten im Wintersemester 2006/07

am Lehrstuhl für Praktische Informatik und Bioinformatik

Volker Heun

6. Februar 2007

Version 1.34

Vorwort

Dieses Skript entstand parallel zur Vorlesung *Algorithmen auf Sequenzen* des Wintersemesters 2006/07 und baut auf dem vorherigen Skript der Vorlesungen des Wintersemesters 2003/04 und des Wintersemesters 2004/05 auf. Diese Vorlesung wurde an der Ludwig-Maximilians-Universität speziell für Studenten der Bioinformatik, aber auch für Studenten der Informatik, im Rahmen des gemeinsam von der Ludwig-Maximilians-Universität München und der Technischen Universität München veranstalteten Studiengangs Bioinformatik gehalten.

Das vorliegende Skript gibt den Inhalt aller drei Vorlesungen wieder, die sich jedoch inhaltlich ein wenig unterscheiden. Die Teile die im Wintersemester 2006/07 nicht Teil der Vorlesung waren, sind mit einem * markiert.

Diese Fassung ist zwar korrigiert, aber noch nicht prinzipiell überarbeitet worden, so dass das Skript an einigen Stellen etwas kurz und unpräzise ist und sicherlich auch noch eine Reihe von (Tipp)Fehlern enthält. Daher bin ich für jeden Hinweis darauf (an Volker.Heun@bio.ifi.lmu.de) dankbar.

An dieser Stelle möchte ich Sabine Spreer, die an der Erstellung des ersten Kapitels in L^AT_EX 2_ε maßgeblich beteiligt war, sowie Alois Huber und Hermann Klann danken, die an der Erstellung des zweiten mit sechsten Kapitels in L^AT_EX 2_ε maßgeblich beteiligt waren, außerdem Herrn Stempel für seine Hinweise auf Ungenauigkeiten und Tippfehler. Weiterhin möchte ich insbesondere meinen Mitarbeitern Johannes Fischer und Simon W. Ginzinger für Ihre Unterstützung bei der Veranstaltung danken, die somit das vorliegende Skript erst möglich gemacht haben.

München, im Wintersemester 2006/07

Volker Heun

Inhaltsverzeichnis

1	Optimal Scoring Subsequences	1
1.1	Maximal Scoring Subsequence	1
1.1.1	Problemstellung	1
1.1.2	Biologische Anwendungen	2
1.1.3	Naive Lösung	4
1.1.4	Lösen durch dynamische Programmierung	4
1.1.5	Divide-and-Conquer-Ansatz	5
1.1.6	Cleverer Lösung	6
1.1.7	Zusammenfassung	7
1.2	All Maximal Scoring Subsequences	8
1.2.1	Problemstellung	8
1.2.2	Elementare Eigenschaften der strukturellen Definition	10
1.2.3	Ein Algorithmus zur Lösung	18
1.2.4	Zeitkomplexität	22
1.3	Bounded All Maximum Scoring Subsequences	24
1.3.1	Problemstellung	24
1.3.2	Lösung mittels Dynamischer Programmierung	25
1.3.3	Effiziente Lösung mittels Dynamischer Programmierung	26
1.4	Bounded Maximal Scoring Subsequence	27
1.4.1	Problemstellung	27
1.4.2	Links-Negativität	27
1.4.3	Algorithmus zur Lösung des BMSS-Problems	29
1.5	Maximal Average Scoring Subsequence (*)	31

1.5.1	Problemstellung	31
1.5.2	Rechtsschiefe Folgen und fallend rechtsschiefe Partitionen . . .	32
1.5.3	Algorithmus zur Konstruktion rechtsschiefer Zeiger	36
1.5.4	Elementare Eigenschaften von MASS	37
1.5.5	Ein Algorithmus für MASS	39
2	Suffix-Trees Revisited	43
2.1	Definition von Suffix-Tries und Suffix-Trees	43
2.1.1	Σ -Bäume und (kompakte) Σ^+ -Bäume	43
2.1.2	Grundlegende Notationen und elementare Eigenschaften . . .	44
2.1.3	Suffix-Tries und Suffix-Trees	45
2.2	Repräsentationen von Bäumen	48
2.2.1	Darstellung der Kinder mit Feldern	48
2.2.2	Darstellung der Kinder mit Listen	49
2.2.3	Darstellung der Kinder mit balancierten Bäumen	50
2.2.4	Darstellung des Baumes mit einer Hash-Tabelle	50
2.2.5	Speicherplatzeffiziente Feld-Darstellung	52
2.3	WOTD-Algorithmus	53
2.3.1	Die Konstruktion	53
2.3.2	Zeitbedarf	55
2.4	Der Algorithmus von Ukkonen	56
2.4.1	Suffix-Links	56
2.4.2	Verschachtelte Suffixe und verzweigende Teilwörter	57
2.4.3	Idee von Ukkonens Algorithmus	58
2.4.4	Ukkonens Online Algorithmus	61
2.4.5	Zeitanalyse	66

3	Repeats	69
3.1	Exakte und maximale Repeats	69
3.1.1	Erkennung exakter Repeats	69
3.1.2	Charakterisierung maximaler Repeats	72
3.1.3	Erkennung maximaler Repeats	76
3.1.4	Revers-komplementäre Repeats	78
3.2	Tandem-Repeats	79
3.2.1	Was sind Tandem-Repeats	80
3.2.2	Eigenschaften von Tandem-Repeats	81
3.2.3	Algorithmus von Stoye und Gusfield	84
3.2.4	Laufzeitanalyse und Beschleunigung	86
3.2.5	Algorithmus von Main und Lorentz	89
3.2.6	Longest Common Extensions	90
3.2.7	Algorithmus von Landau und Schmidt	93
3.3	Vokabulare von Tandem-Repeats	96
3.3.1	Vokabulare und Überdeckungen	96
3.3.2	Skizze des Algorithmus von Gusfield und Stoye	99
3.3.3	Tandem-Repeats und Lempel-Ziv-Zerlegungen	99
3.3.4	Phase I: Bestimmung einer linkesten Überdeckung	103
3.3.5	Phase II: Dekorierung einer Teilmenge	109
3.3.6	Phase III: Vervollständigung der Dekorierung von $\mathcal{V}(t)$	114
4	Interludium: Lowest Common Ancestors	119
4.1	Algorithmus von Bender und Farach-Colton	119
4.1.1	Lowest Common Ancestor und Range Minimum Queries . . .	119
4.1.2	Euler-Tour eines gewurzelten Baumes	120
4.1.3	Reduktion LCA auf RMQ	121

4.1.4	Ein quadratischer Algorithmus für RMQ	122
4.1.5	Eine verbesserte Variante	123
4.1.6	Incremental Range Minimum Query (*)	124
4.1.7	Ein optimaler Algorithmus (*)	125
4.1.8	Optimale Lösung für RMQ (*)	127
4.1.9	Eine einfachere optimale Variante nach Alstrup et al.	129
4.2	Algorithmus von Schieber und Vishkin (*)	133
4.2.1	LCA-Queries auf vollständigen binären Bäumen	133
4.2.2	LCA-Queries auf beliebigen Bäumen	137
4.2.3	Vorverarbeitung	141
4.2.4	Beziehung zwischen Vorfahren in T und B	143
4.2.5	Berechnung einer LCA-Query	144
4.2.6	LCA-Query-Algorithmus	147
5	Suffix-Arrays	151
5.1	Grundlegende Eigenschaften von Suffix-Arrays	151
5.1.1	Was sind Suffix-Arrays?	151
5.1.2	Konstruktion aus Suffix-Bäumen	152
5.1.3	Algorithmus von Manber und Myers	153
5.2	Algorithmus von Ko und Aluru (*)	158
5.2.1	Typ S und Typ L Suffixe	158
5.2.2	Sortieren der Suffixe mit sortierten Typ S/L Suffixen	160
5.2.3	Sortierung der Typ S Suffixe	163
5.2.4	Der Algorithmus von Ko und Aluru und Laufzeitanalyse	169
5.3	Skew-Algorithmus von Kärkkäinen und Sanders	170
5.3.1	Tripel und Verkürzung der Zeichenreihe	170
5.3.2	Rekursives Sortieren der verkürzten Zeichenreihe	171

5.3.3	Sortieren der restlichen Suffixe	172
5.3.4	Mischen von A_{12} und A_0	173
5.3.5	Laufzeitanalyse	173
5.4	Suchen in Suffix-Arrays	174
5.4.1	Binäre Suche	174
5.4.2	Verbesserte binäre Suche	175
5.4.3	Effiziente Berechnung der LCP-Tabelle	178
5.5	Enhanced Suffix-Arrays (*)	180
5.5.1	LCP-Intervalle	180
5.5.2	Die Child-Tabelle und ihre Eigenschaften	184
5.5.3	Optimale Suche in Enhanced Suffix-Arrays	186
5.5.4	Berechnung der Child-Tabelle	188
5.5.5	Komprimierte Darstellung der Child-Tabelle	190
5.5.6	Simulation von Suffix-Baum-Algorithmen auf Suffix-Arrays . .	192
5.6	Extended Suffix Arrays	194
5.6.1	LCP-Intervalle	194
5.6.2	Navigation im Extended Suffix-Array	198
5.6.3	Optimale Suche in Extended Suffix-Arrays	199
5.6.4	Auffinden des Elters	199
5.6.5	Suffix-Links	201
5.6.6	LCA-Queries in Extended Suffix-Arrays	202
5.6.7	Speicherplatzbedarf	203

6	Genome Rearrangements	205
6.1	Modellbildung	205
6.1.1	Rearrangements und zugehörige Mutationen	205
6.1.2	Permutationen	207
6.2	Sorting by Reversals	209
6.2.1	Komplexität von Min-SBR	210
6.2.2	2-Approximation für Min-SBR	210
6.2.3	Algorithmus und Laufzeitanalyse	216
6.3	Eine bessere untere Schranke für Min-SBR	218
6.3.1	Breakpoint-Graphen	218
6.3.2	Elementare Beobachtungen	220
6.3.3	Die untere Schranke	222
6.4	Sorting by Oriented Reversals	224
6.4.1	Orientierte Permutationen	224
6.4.2	Reality-Desire-Diagramm	226
6.4.3	Der Overlap-Graph	228
6.4.4	Hurdles and Fortresses	233
6.4.5	Eine untere Schranke für Min-SOR	235
6.4.6	Sortierung orientierter Komponenten	237
6.4.7	Eliminierung von Hurdles	244
6.4.8	Algorithmus für Min-SOR	247
6.5	Sorting by Transpositions (*)	248
6.5.1	Elementare Definitionen	248
6.5.2	Untere Schranken für die Transpositions-Distanz	250
6.5.3	Orientierte und Unorientierte Desire-Edges	251
6.5.4	Eine 2-Approximation	252
6.6	Sorting by Transpositions and Reversals (*)	257

6.6.1	Problemstellung	257
6.6.2	Untere Schranken für die Transversal-Distanz	258
6.6.3	Eine 2-Approximation	260
6.7	Sorting by weighted Transversals (*)	266
6.7.1	Gewichtete Transversal-Distanz	266
6.7.2	Starke Orientierung und starke Hurdles	267
6.7.3	Eine untere Schranke für die Transversal-Distanz	268
6.7.4	Eine Approximation für Min-SWT	270
6.7.5	Approximationsgüte	272
6.8	Weitere Modelle (*)	274
6.8.1	Gewichtung der Operationen durch ihre Länge	274
6.8.2	Duplikationen	275
6.8.3	Multi-chromosomale Genome	276
6.8.4	Multiple Genome Rearrangements	276
A	Literaturhinweise	279
A.1	Lehrbücher zur Vorlesung	279
A.2	Skripten anderer Universitäten	279
A.3	Originalarbeiten	280
A.3.1	Optimal Scoring Subsequences	280
A.3.2	Suffix-Trees	281
A.3.3	Repeats	281
A.3.4	Lowest Common Ancestors and Range Minimum Queries . . .	282
A.3.5	Construction of Suffix-Arrays	282
A.3.6	Applications of Suffix-Arrays	283
A.3.7	Sorting by Reversals	284
A.3.8	Sorting by Oriented Reversals	284

A.3.9	Sorting by Transpositions	285
A.3.10	Sorting by Transversals	285
A.3.11	Erweiterungen zu Genome Rearrangements	286
B	Index	289

Optimal Scoring Subsequences

1

1.1 Maximal Scoring Subsequence

Ziel dieses Abschnittes ist es, (möglichst effiziente) Algorithmen für das Maximal Scoring Subsequence Problem vorzustellen. Dabei werden wir zunächst noch einmal kurz die wichtigsten Paradigmen zum Entwurf von Algorithmen wiederholen.

1.1.1 Problemstellung

MAXIMAL SCORING SUBSEQUENCE (MSS)

Eingabe: Eine Folge $(a_1, \dots, a_n) \in \mathbb{R}^n$.

Gesucht: Eine (zusammenhängende) Teilfolge $(a_i, \dots, a_j)$, die $\sigma(i, j)$ maximiert, wobei $\sigma(i, j) = \sum_{\ell=i}^j a_\ell$.

Bemerkung: Mit Teilfolgen sind in diesem Kapitel immer (sofern nicht anders erwähnt) zusammenhängende (d.h. konsekutive) Teilfolgen einer Folge gemeint (also anders als beispielsweise in der Analysis).

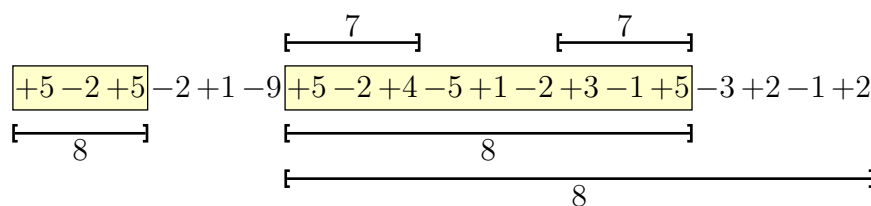


Abbildung 1.1: Beispiel: Maximal Scoring Subsequences

In Abbildung 1.1 ist ein Beispiel angegeben. Wie man dort sieht, kann es mehrere (und auch nicht-disjunkte) Lösungen geben.

Bemerkungen:

- Es sind mehrere Lösungen möglich.

- Ist eine Lösung in der anderen enthalten, so wählen wir als Lösung immer eine kürzester Länge. Die anderen ergeben sich aus Anhängen von Teilfolgen mit dem Score Null. Darüber hinaus haben solche Lösungen noch eine schöne Eigenschaft, wie wir gleich sehen werden.
- Es gibt keine echt überlappenden Lösungen. Angenommen, es gäbe echt überlappende Lösungen a' und a'' einer gegebenen Folge a (siehe dazu auch Abbildung 1.2). Die Sequenz gebildet aus der Vereinigung beider Sequenzen (respektive ihrer Indices) müsste dann einen höheren Score haben, da der Score der Endstücke ≥ 0 ist (sonst würde er in den betrachteten Teilfolgen nicht berücksichtigt werden).

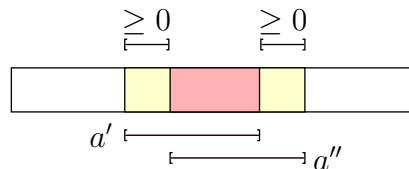


Abbildung 1.2: Skizze: Überlappende optimale Teilfolgen

1.1.2 Biologische Anwendungen

In diesem Abschnitt wollen wir kurz einige biologische Probleme vorstellen, die sich als Maximal Scoring Subsequence formulieren lassen.

Transmembranproteine: Bestimmung der transmembranen Regionen eines Proteins. Eingelagerte Proteine in der Membran sollten einen ähnlichen Aufbau wie die Membran selbst haben, da die Gesamtstruktur stabiler ist. Somit sollten transmembrane Regionen hydrophob sein.

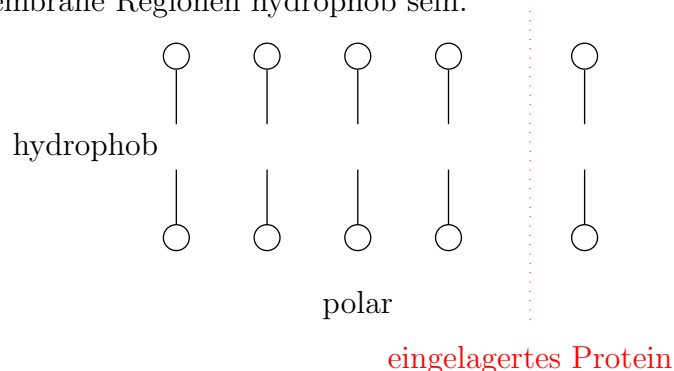


Abbildung 1.3: Beispiel: transmembrane Proteine

Mit einer geeigneten Gewichtung der Aminosäuren, können solche hydrophoben Regionen mit Hilfe der Lösung eines Maximal Scoring Subsequence Problems gefunden werden.

Für die einzelnen Aminosäuren werden die folgende Werte gemäß der Hydrophobizität der entsprechenden Aminosäure gewählt:

- hydrophobe Aminosäuren: ein Wert aus $\in [0 : 3]$;
- hydrophile Aminosäuren: ein Wert aus $\in [-5 : 0]$.

Lokalisierung GC-reicher DNS-Abschnitte: In GC-reichen Regionen der DNS finden sich häufig Gene. Das Auffinden solcher GC-reicher Regionen lässt sich als Maximal Scoring Subsequence Problem beschreiben:

- $C, G \rightarrow 1 - p$ für ein $p \in [0 : 1]$;
- $A, T \rightarrow -p$.

Zusätzlich können Längenbeschränkungen sinnvoll sein; obere, untere Schranke der Länge für z.B. Proteine, die man sucht.

Vergleichende Analyse von Genomen: Im Vergleich des Mensch- und Maus-Genoms liegen Sequenz-Ähnlichkeiten für Exons bei 85% und für Introns bei 35%. Mit Hilfe eines lokalen Sequenz-Alignments (Smith-Waterman) lassen sich solche Übereinstimmungen gut auffinden. Jedoch kann es bei den gefundenen Lösungen den so genannten Mosaik-Effekt geben, d.h. sehr ähnliche Sequenzen sind immer wieder von sehr unähnlichen, jedoch relativ kurzen Stücken unterbrochen.

Mit Hilfe eines geeigneten Maximal Scoring Subsequences Problems können solche Mosaik-Effekte aufgedeckt werden. Hierzu wird eine Variante des Maximal Scoring Subsequence Problems verwendet. Man normiert die erzielten Scores mit der Länge der zugehörigen Teilfolge. Somit lassen sich so genannte *poor regions* (stark positive Teile, die mit kurzen stark negativen Fragmenten unterbrochen sind) ausschließen.

Konservierte Regionen: Gut konservierte Regionen eines mehrfachen Sequenzalignments lassen sich durch Gewichtung der Spalten gemäß ihrer Ähnlichkeiten (beispielsweise SP-Maß einer Spalte) und einem anschließenden Auffinden von Maximal Scoring Subsequences bestimmen.

'Ungapped' local alignment: Auch lokales Alignment ohne Lücken (gaps) können aus der Dot-Matrix durch Anwenden von Algorithmen für das Maximal Scoring Subsequence Problem auf die Diagonalen effizient finden. Dieses Verfahren ist insbesondere dann effizient, wenn man mit Längenrestriktionen arbeiten will, oder den Score ebenfalls wieder mit der zugehörigen Länge der Folge normalisieren will.

1.1.3 Naive Lösung

Im Folgenden wollen wir eine Reihe von Algorithmen zur Lösung des Maximal Scoring Subsequence Problems vorstellen, die jeweils effizienter als die vorher vorgestellte Variante ist.

Beginnen wollen wir mit einem naiven Ansatz. Wir benötigen jedoch vorher noch eine Notation.

Notation 1.1 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine Folge reeller Werte, dann bezeichnet $\sigma(i, j) := \sum_{\ell=i}^j a_\ell$ für $i \leq j \in [1 : n]$.

Die naive Methode bestimmt zuerst alle Werte $\sigma(i, j)$ für alle $i \leq j \in [1 : n]$.

Für die Laufzeit (Anzahl Additionen, die proportional zur Laufzeit ist) ergibt sich dann pro Tabelleneintrag $\Theta(j - i)$, also insgesamt:

$$\sum_{i=1}^n \sum_{j=i}^n \Theta(j - i) = \Theta \left(\sum_{i=1}^n (n - i)^2 \right) = \Theta \left(\sum_{i=1}^n i^2 \right) = \Theta(n^3).$$

Ein alternative Begründung ist die folgende: In der Tabelle mit n^2 Einträgen benötigt jeder Eintrag maximal n Operationen. Hierbei erhalten wir jedoch nur eine obere Schranke und nicht die korrespondierende untere Schranke für die Laufzeit.

1.1.4 Lösen durch dynamische Programmierung

Ein anderer Ansatz ergibt sich aus einer trivialen Rekursionsgleichung für $\sigma(i, j)$. Es gilt folgende Rekursionsgleichung:

$$\sigma(i, j) = \begin{cases} a_i & \text{für } i = j \\ \sigma(i, k) + \sigma(k + 1, j) & \text{für ein } k \in [i : j - 1] \text{ sofern } i < j \end{cases}$$

In der Regel ist eine direkte Implementierung dieser Rekursionsgleichung zu aufwendig, da meist exponentiell viele Aufrufe erfolgen (siehe zum Beispiel rekursive Berechnung einer Fibonacci-Zahl)! In diesem Fall wären es sogar nur n^3 rekursive Aufrufe, was aber nicht besser als der naive Ansatz ist. Mit Hilfe der *dynamische Programmierung* können wir jedoch effizienter werden, da hier Werte mehrfach berechnet werden.

Die Tabellengröße ist $O(n^2)$. Jeder Eintrag kann mit der Rekursionsgleichung in Zeit $O(1)$ berechnet werden. Dabei wird die Tabelle diagonal von der Mitte nach rechts oben aufgefüllt (siehe auch Abbildung 1.4).

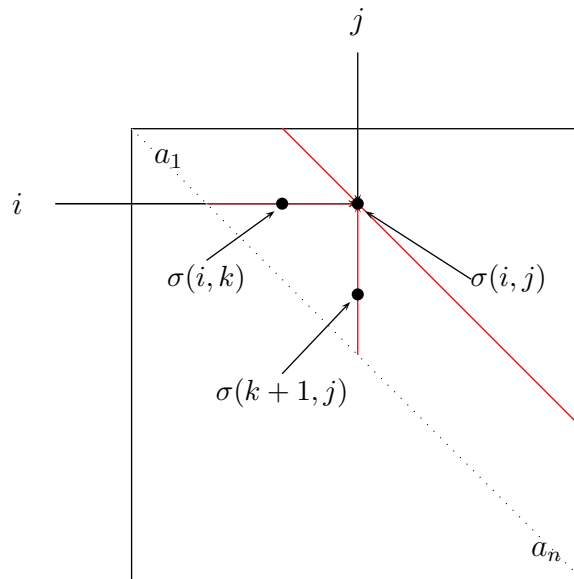


Abbildung 1.4: Skizze: Auffüllen der dynamischen Programmierungstabelle

1.1.5 Divide-and-Conquer-Ansatz

Eine andere Lösungsmöglichkeit erhalten wir einem Divide-and-Conquer-Ansatz, wie in Abbildung 1.5 illustriert. Dabei wird die Folge in zwei etwa gleich lange Folgen aufgeteilt und die Lösung in diesen rekursiv ermittelt.

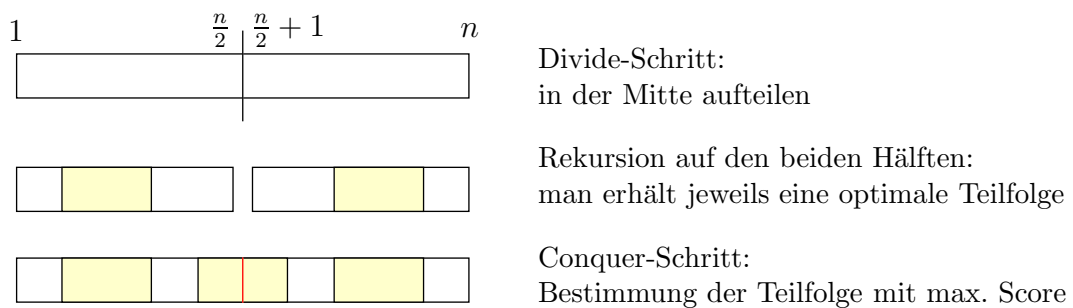


Abbildung 1.5: Skizze: Divide-and-Conquer bei Maximal Scoring Subsequences

Man kann dabei aber auch die optimale Teilfolge in der Mitte zerschneiden. Daher muss man zusätzlich von der Mitte aus testen, wie von dort nach rechts bzw. links eine optimale Teilfolge aussieht (siehe auch Abbildung 1.6).

Dazu bestimmen wir jeweils das Optimum der Hälften, d.h

$$\max \{ \sigma(i, n/2) \mid i \in [1 : n/2 + 1] \} \quad \text{und} \quad \max \{ \sigma(n/2 + 1, j) \mid j \in [n/2 : n] \}.$$

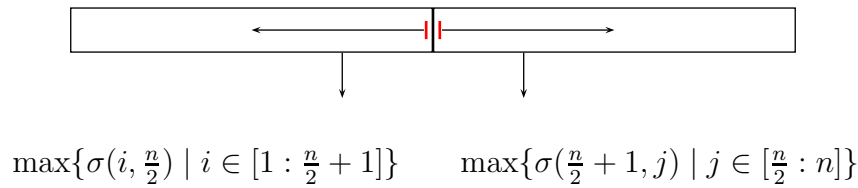


Abbildung 1.6: Skizze: Conquer-Step

Man überlegt sich leicht, dass die optimale Teilfolge, die die Positionen $n/2$ und $n/2 + 1$ überdeckt, aus der Konkatenation der beiden berechneten optimalen Teilfolgen in den jeweiligen Hälften bestehen muss.

Für die Laufzeit erhalten wir sofort die folgende Rekursionsgleichung, die identisch zur Laufzeitanalyse von Mergesort ist, da das Bestimmen einer optimalen Teilfolge über die Mitte hinweg in Zeit $O(n)$ (wie oben gesehen) geschehen kann:

$$T(n) = 2 \cdot T\left(\frac{n}{2}\right) + O(n) = O(n \log(n)).$$

1.1.6 Clevere Lösung

Wenn wir wie beim Divide-and-Conquer-Ansatz das Problem nicht in der Mitte aufteilen, sondern am rechten Rand, so können wir (ganz im Gegensatz zum Problem des Sortierens) eine effizientere Lösung finden (siehe auch Abbildung 1.7).

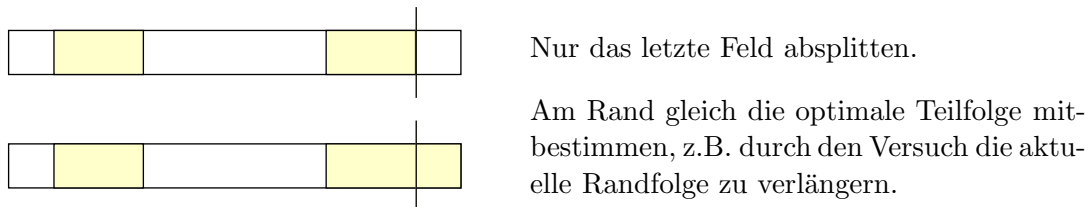


Abbildung 1.7: Skizze: Asymmetrische Divide-and-Conquer

Das einzige Problem besteht darin, dass die Folge am rechten Ende einen negativen Score erhält. Dann wähle die leere Folge mit dem Score 0.

Man kann diese Idee rekursiv als Divide and -Conquer-Step implementieren oder iterativ wie in Abbildung 1.8 angegeben auflösen.

Da im Wesentlichen einmal linear über die Folge gelaufen wird und für jedes Element nur konstante Kosten (Additionen, Maximumsbildungen) anfallen, erhalten wir offensichtlich eine Laufzeit von $O(n)$.

```

MSS (real  $a[]$ , int  $n$ )
{
    int  $max = 0$ ;    $\ell = 1$ ;    $r = 0$ ;
    int  $rmax = 0$ ;    $rstart = 1$ ;

    for ( $i = 1$ ;  $i \leq n$ ;  $i++$ )
    {
        if ( $rmax + a_i > 0$ )
             $rmax := rmax + a_i$ ;
        else
        {
             $rmax = 0$ ;
             $rstart = i + 1$ ;
        }

        if ( $rmax > max$ )
        {
             $max = rmax$ ;
             $\ell = rstart$ ;
             $r = i$ ;
        }
    }
}

```

Abbildung 1.8: Algorithmus: Die clevere Lösung

Da dies eine offensichtlich optimale Lösung ist, halten wir das Ergebnis im folgenden Satz fest.

Theorem 1.2 *Eine Teilfolge mit maximalem Wert einer gegebenen reellen Folge lässt sich in Linearzeit mit konstantem zusätzlichen Platzbedarf bestimmen.*

1.1.7 Zusammenfassung

In der folgenden Tabelle sind alle Resultate der vorgestellten Algorithmen noch einmal zusammengefasst.

In der folgenden Tabelle sind noch die Längen von Folgen angegeben, die in einer Sekunde bzw. einer Minute auf einem gewöhnlichen Rechner (AMD-Athlon, 2GHz, 2GB Hauptspeicher) verarbeitet werden können.

Algorithmus	Zeit	Platz	Bemerkung
naive Programmierung	$O(n^3)$	$O(n^2)$	Tabelle füllen
DP	$O(n^2)$	$O(n^2)$	geschickter füllen
Divide and Conquer	$O(n \log(n))$	$O(n)$	die Eingabelänge ist n
Cleverer Lösung	$O(n)$	$n + O(1)$	die Eingabelänge ist n

Abbildung 1.9: Tabelle: Laufzeiten für die MSS Lösungsansätze

Seq-Len	1sec.		1min.	Space
Naive	800	$\xrightarrow{\times 4}$	3.200	$O(n^2)$
Dyn.Prog.	4.200	$\xrightarrow{\times 8}$	32.000	$O(n^2)$
D&C	1.500.000	$\xrightarrow{\times 50}$	75.000.000	$O(n)$
Clever	20.000.000	$\xrightarrow{\times 60}$	1.200.000.000	$n + O(1)$

Abbildung 1.10: Tabelle: zu verarbeitenden Problemgrößen

1.2 All Maximal Scoring Subsequences

Nun wollen wir uns mit der Frage beschäftigen, wenn wir nicht nur eine beste, sondern alle besten bzw. alle möglichen Teilfolgen mit positive Score und zwar nach absteigenden Score erhalten wollen. Zuerst einmal müssen wir uns über die Fragestellung klar werden, d.h. was ist überhaupt die gesuchte Lösung.

17.10.06

1.2.1 Problemstellung

Geben wir zunächst die formale Problemstellung an und diskutieren anschließend die damit verbundenen Probleme.

ALL MAXIMAL SCORING SUBSEQUENCES (AMSS)

Eingabe: Eine Folge $(a_1, \dots, a_n) \in \mathbb{R}^n$.

Gesucht: Alle disjunkten Teilfolgen, die ihren Score maximieren.

Zunächst einmal muss man sich überlegen, was es heißen soll, dass *alle disjunkten Teilfolgen ihren Score maximieren*. Betrachten wir dazu das Beispiel in Abbildung 1.11. Die lange Folge ist keine Lösung, da wir keine überlappenden Folgen haben wollen.

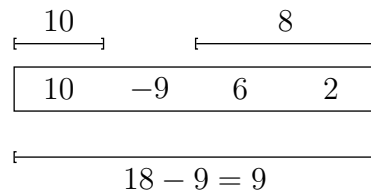


Abbildung 1.11: Beispiel: Maximal bewertete Teilfolgen

Wir geben zwei mögliche Definitionen an, wie man alle maximalen Teilfolgen einer Folge definieren kann. Im Folgenden werden wir zeigen, dass die beiden Definitionen äquivalent sind. A priori ist dies überhaupt nicht klar.

Definition 1.3 (Strukturelle Definition) Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine reelle Folge. Eine Teilfolge $a' = (a_i, \dots, a_j)$ von a heißt maximal bewertet (engl. maximal scoring), wenn die beiden folgenden Eigenschaften erfüllt sind:

- (E1) alle echten Teilfolgen von a' haben einen kleineren Score;
- (E2) keine echte Oberfolge von a' in a erfüllt E1.

Die Bedingung E1 ist nach unserer vorhergehenden Diskussion klar, da wir keine Teilfolge mit größerem Score, die Teilfolge einer anderen Teilfolge ist, als Lösung verlieren wollen.

Als Bedingung E2 würde man vermutlich zunächst erwarten, dass jede Oberfolge ebenfalls einen kleineren Score als die gegebene Teilfolge besitzen soll. Die ist jedoch zu naiv, wie das vorherige Beispiel mit $a = (10, -9, 6, 2)$ zeigt. Damit wäre die Teilfolge $a' = (6, 2)$ keine maximal bewertete Teilfolge, da die Oberfolge a einen höheren Score besitzt: $9 > 8$. Als Lösung würde man jedoch sicherlich $\text{MSS}(a) = \{a_1, a_3 \cdot a_4\}$ erwarten. Diese beiden Folgen erfüllen jedoch sowohl E1 als auch E2.

Neben dieser strukturellen Definition kann man auch noch eine eher algorithmisch angelehnte Definition angeben.

Definition 1.4 (Prozedurale Definition) Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine reelle Folge. Eine kürzeste Teilfolge $(a_i, \dots, a_j)$ von a mit maximalem Score heißt maximal bewertet (engl. maximal scoring). Teilfolgen aus $(a_1, \dots, a_{i-1})$ bzw. $(a_{j+1}, \dots, a_n)$, die für diese maximal bewertet sind, sind auch für a maximal bewertet.

Man überlegt sich leicht, dass die Menge aller maximal bewerteten Teilfolgen nach der prozeduralen Definition eindeutig ist. Wir werden später noch sehen, dass dies

auch für die strukturelle Definition gilt (das folgt aus der noch zu zeigenden Äquivalenz der beiden Definitionen).

Notation 1.5 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine reelle Folge, dann bezeichnet $MSS(a)$ die Menge aller maximal bewerteten Teilfolgen von a .

Aus der prozeduralen Definition kann sofort ein rekursiver Algorithmus zur Bestimmung aller maximal bewerteter Teilfolgen abgeleitet werden, der in der folgenden Skizze in Abbildung 1.12 veranschaulicht ist. Man bestimmt zunächst eine Maximal

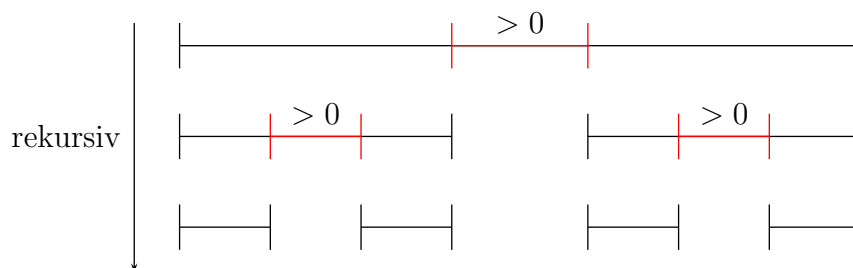


Abbildung 1.12: Skizze: Rekursiver Ansatz für AMSS

Scoring Subsequence und entfernt diese aus der Folge. Für die beiden entstehenden Folgen wird dieser Algorithmus rekursiv aufgerufen.

Die Laufzeit dieses Algorithmus erfüllt folgende Rekursionsgleichung, da das Auffinden einer Maximal Scoring Subsequence, wie wir im letzten Abschnitt gesehen haben, in Zeit $O(n)$ durchführbar ist:

$$T(n) = O(n) + T(n_1) + T(n_2) \quad \text{mit} \quad n_1 + n_2 < n.$$

Ähnlich zu Quicksort ergibt sich folgende Analyse. Im worst-case benötigt dieser Algorithmus offensichtlich maximal Zeit $O(n^2)$.

Im average-case kann man aus der Analyse von Quicksort herleiten, dass auch dieser Algorithmus einen Zeitbedarf von $O(n \log(n))$ hat. Hierzu muss man jedoch eine geeignete Wahrscheinlichkeitsverteilung annehmen, die in der Rechnung äquivalent zu der von Quicksort ist, die aber nicht unbedingt realistisch sein muss!

1.2.2 Elementare Eigenschaften der strukturellen Definition

In diesem Abschnitt werden wir einige grundlegende Eigenschaften maximal bewerteter Teilfolgen nach der strukturellen Definition herleiten, die es uns erlauben werden, die Äquivalenz der beiden Definitionen maximal bewerteter Teilfolgen zu zeigen.

Darauf basierend werden wir im nächsten Abschnitt einen effizienten Algorithmus für die Lösung des All Maximal Scoring Subsequences Problem vorstellen.

Lemma 1.6 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Für jede Teilfolge $a' = (a_i, \dots, a_j)$ von a sind äquivalent:

1) a' erfüllt E1.

2) Es gilt für das Minimum aller Präfixe von a in a'

$$\sigma(1, i-1) = \min\{\sigma(1, k) \mid k \in [i-1, j]\}$$

und für das Maximum aller Präfixe von a in a'

$$\sigma(1, j) = \max\{\sigma(1, k) \mid k \in [i-1, j]\}$$

und dass diese eindeutig sind!

3) $\forall k \in [i : j] : \sigma(i, k) > 0 \wedge \sigma(k, j) > 0$.

Beweis: $1 \Rightarrow 2$: Der Beweis erfolgt durch Widerspruch:

1. Fall: Sei $k \in [i-1 : j-1]$ mit $\sigma(1, k) \geq \sigma(1, j)$.

Dann ergibt sich die folgende Situation, die in Abbildung 1.13 dargestellt ist.

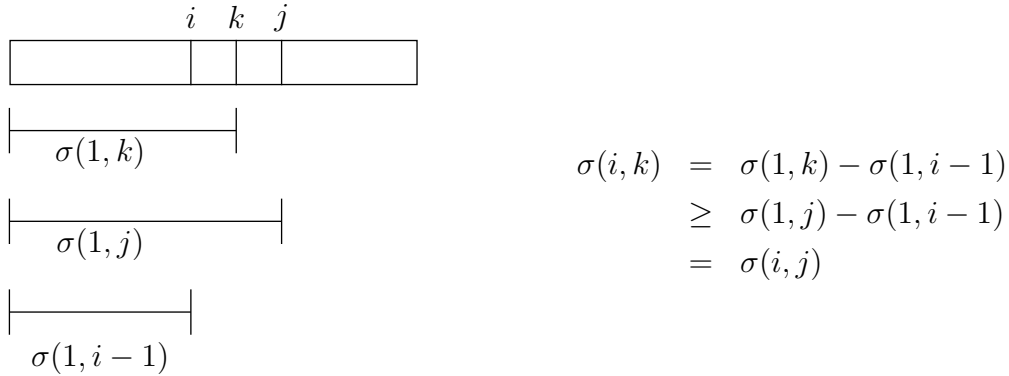


Abbildung 1.13: Skizze: Fall 1: $\sigma(1, k) \geq \sigma(1, j)$

Daraus ergibt sich ein Widerspruch zur Annahme der Eigenschaft E1, da die echte Teilfolge $(a_i, \dots, a_k)$ von $(a_i, \dots, a_j)$ keinen kleineren Score besitzt.

Anmerkung: Der Beweis gilt auch für $k = i-1$. In diesem Fall wäre $\sigma(i, j) \leq 0$, was nicht sein kann.

2. Fall: Sei jetzt $k \in [i : j]$ mit $\sigma(1, k) \leq \sigma(1, i - 1)$.

Dann ergibt sich die folgende Situation, die in Abbildung 1.14 dargestellt ist.

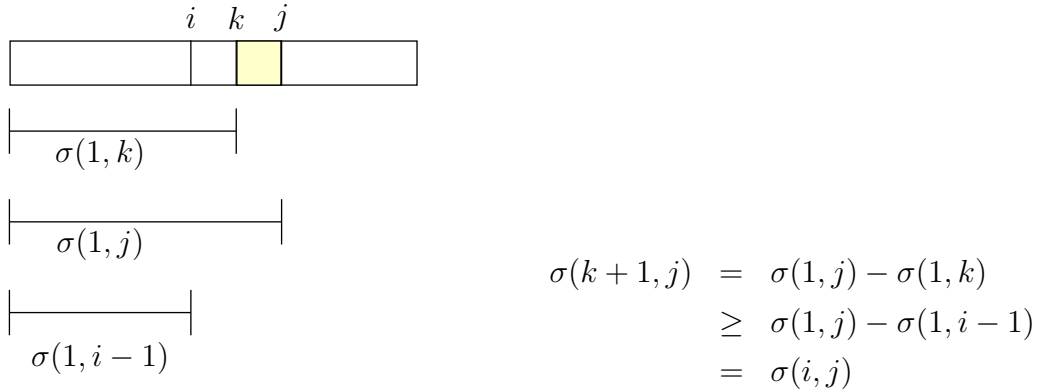


Abbildung 1.14: Skizze: $\sigma(1, k) \leq \sigma(1, i - 1)$

Daraus ergibt sich ein Widerspruch zur Annahme der Eigenschaft E1, da die echte Teilfolge $(a_{k+1}, \dots, a_j)$ von $(a_i, \dots, a_j)$ keinen kleineren Score besitzt.

Anmerkung: Der Beweis gilt auch für $k = j$. In diesem Fall wäre $\sigma(i, j) \leq 0$, was nicht sein kann.

2 $\Rightarrow$ 3 : Da das Minimum eindeutig ist, folgt $\sigma(1, i-1) < \sigma(1, k)$ für alle $k \in [i : j]$ und somit:

$$\sigma(i, k) = \sigma(1, k) - \underbrace{\sigma(1, i-1)}_{< \sigma(1, k)} > 0.$$

Da das Maximum eindeutig ist, folgt $\sigma(1, j) > \sigma(1, k-1)$ für alle $k \in [i : j]$ und somit:

$$\sigma(k, j) = \underbrace{\sigma(1, j)}_{> \sigma(1, k-1)} - \sigma(1, k-1) > 0.$$

3 $\Rightarrow$ 1 : Sei $a' = (a_k, \dots, a_\ell)$ mit $i \leq k \leq \ell \leq j$ sowie $i \neq k$ oder $\ell \neq j$. Dann gilt:

$$\sigma(k, \ell) = \sigma(i, j) - \underbrace{\sigma(i, k-1)}_{\geq 0} - \underbrace{\sigma(\ell+1, j)}_{\geq 0} > \sigma(i, j).$$

Hinweis: $\sigma(i, k-1)$ und $\sigma(\ell+1, j)$ sind jeweils ≥ 0 ; eines davon muss > 0 sein, da sonst $a' = a''$ gelten würde.

Damit ist gezeigt, dass E1 gilt. ■

Lemma 1.7 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Die maximal bewerteten Teilfolgen von a sind paarweise disjunkt.

Beweis: Wir führen auch hier den Beweis durch Widerspruch. Nehmen wir zunächst an, dass a' eine Teilfolge von a'' ist. Dies kann jedoch nicht sein, da dann a'' eine Oberfolge von a' ist, die E1 erfüllt (da a'' eine maximal bewertete Teilfolge ist). Also erfüllt a' nicht E2 und kann somit keine maximal bewertete Teilfolge sein und wir erhalten den gewünschten Widerspruch.

Seien a' und a'' zwei maximal bewertete Teilfolgen von a , die sich überlappen. Dies ist in Abbildung 1.15 illustriert.

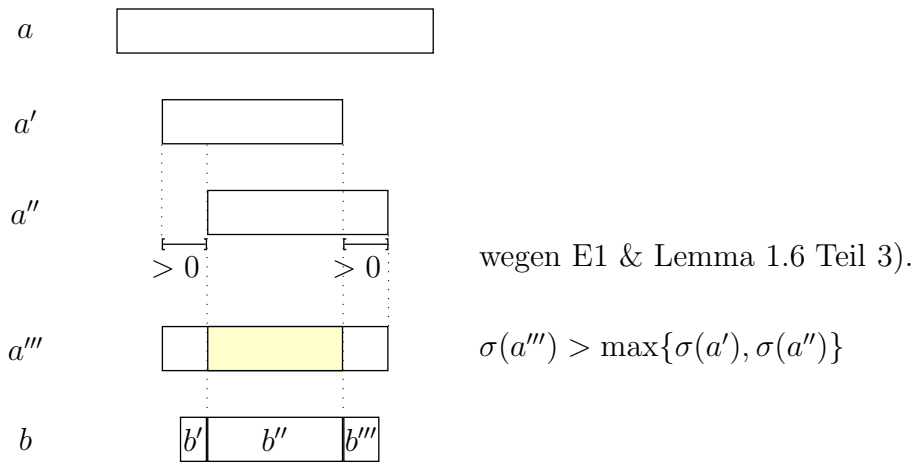


Abbildung 1.15: Skizze: Zwei überlappende maximal bewerteten Teilfolgen

Wir zeigen jetzt, dass die Folge a''' , die als Vereinigung der Folgen a' und a'' (respektive ihrer Indizes) definiert ist, die Eigenschaft E1 erfüllt. Sei dazu b eine beliebige Teilfolge von a''' . Ist b eine Teilfolge von a' (bzw. a'') dann gilt aufgrund der Eigenschaft E1 von a' (bzw. a'') $\sigma(b) < \sigma(a') < \sigma(a''')$ (bzw. $\sigma(b) < \sigma(a') < \sigma(a''')$).

Sei also nun b keine Teilfolge von a' oder a'' . Sei weiter $b = b' \cdot b'' \cdot b'''$, so dass $b' \cdot b''$ ein Suffix von a' und $b'' \cdot b'''$ ein Präfix von a'' ist. Dann gilt:

$$\begin{aligned}
 \sigma(b) &= \sigma(b' \cdot b'') + \sigma(b'' \cdot b''') - \sigma(b'') \\
 &\quad \text{da } \sigma(b' \cdot b'') < \sigma(a') \text{ wegen E1 von } a' \\
 &\quad \text{und } \sigma(b'' \cdot b''') < \sigma(a'') \text{ wegen E1 von } a'' \\
 &< \sigma(a') + \sigma(a'') - \sigma(b'') \\
 &= \sigma(a''')
 \end{aligned}$$

Dies ergibt den gewünschten Widerspruch zur Eigenschaft E2 von a' und a'' , da a''' E1 erfüllt und eine Oberfolge von a' und a'' ist. ■

Lemma 1.8 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Jede Teilfolge $a' = (a_k, \dots, a_\ell)$, die E1 erfüllt, ist Teilfolge einer maximal bewerteten Teilfolge.

Beweis: Wir führen den Beweis durch Widerspruch. Dazu sei $a' = (a_k, \dots, a_\ell)$ ein längstes Gegenbeispiel. a' erfüllt also E1 mit maximaler Länge, d.h. $\ell - k$ ist unter allen Gegenbeispielen maximal. Dies ist in Abbildung 1.16 illustriert.

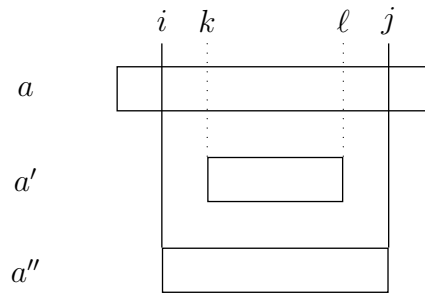


Abbildung 1.16: Skizze: a' ist ein Gegenbeispiel maximaler Länge

a' erfüllt E1, aber **nicht** E2 (sonst wäre a' kein Gegenbeispiel). Somit existiert eine echte Oberfolge $a'' = (a_i, \dots, a_j)$ von a' , die E1 erfüllt. Würde die Folge a'' auch E2 erfüllen, dann wäre a'' eine maximal bewertete Teilfolge, die auch eine Oberfolge von a' ist, d.h. a' wäre kein Gegenbeispiel.

Also erfüllt a'' E1, aber nicht E2. Besäße a'' eine maximal bewertete Oberfolge, so wäre diese auch eine maximal bewertete Oberfolge von a' und a' somit kein Gegenbeispiel.

Also besitzt a'' keine maximal bewertete Oberfolge und erfüllt E1. Somit ist a'' ein längeres Gegenbeispiel als a' , da $j - i > \ell - k$. Dies führt zu einem Widerspruch zur Annahme. ■

19.10.06

Korollar 1.9 Jedes positive Element ist in einer maximal bewerteten Teilfolge enthalten.

Beweis: Dies folgt aus dem vorherigen Lemma, da jede einelementige Folge mit einem positiven Wert die Eigenschaft E1 erfüllt. ■

Korollar 1.10 Innerhalb jeder Teilfolge, die mit keiner maximal bewerteten Teilfolge überlappt, sind die aufaddierten Scores monoton fallend.

Beweis: Nach dem vorherigen Korollar müssen alle Elemente einer solchen Folge nichtpositiv sein. ■

Lemma 1.11 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Betrachte $\mu = \min\{\sigma(1, k) \mid k \in [1 : n]\}$ (bzw. $\mu = \max\{\sigma(1, k) \mid k \in [1 : n]\}$). Sei k der größte (bzw. kleinste) Wert mit $\sigma(1, k) = \mu$. Dann befindet sich k am linken (bzw. rechten) Ende einer maximal bewerteten Teilfolge oder am Ende von a .

Beweis: Betrachte $\sigma(1, i)$ als Funktion von $i \in [0 : n]$ wie in Abbildung 1.17 dargestellt.

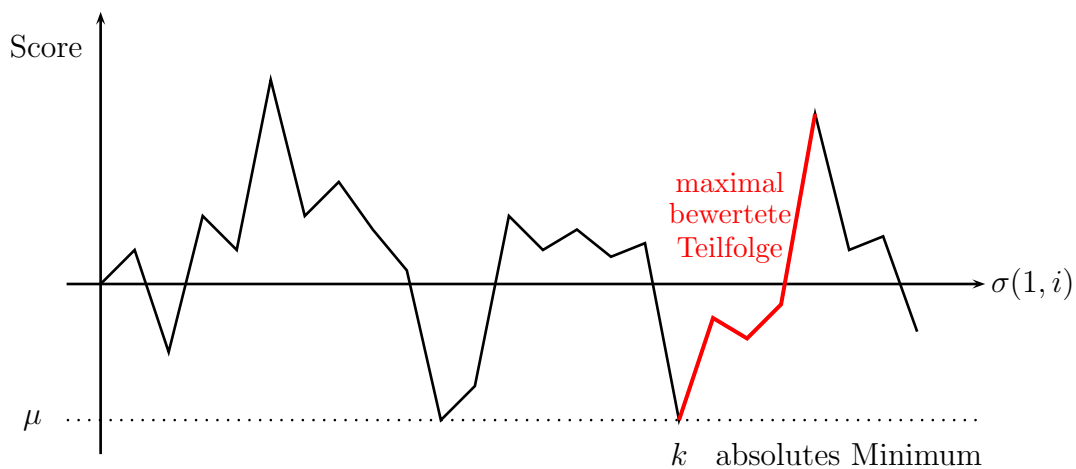


Abbildung 1.17: Skizze: Die Funktion $\sigma(1, i)$

Allgemein gilt $\sigma(i, j) = \sigma(1, j) - \sigma(1, i - 1)$.

Fall 1: Position k befindet sich in keiner maximal bewerteten Teilfolge. Nach Korollar 1.10 ist k am Ende einer nicht maximal bewerteten Teilfolge, da der folgende Wert a_{k+1} (sofern vorhanden) aufgrund der Definition von k positiv sein muss. Also am Ende von a oder an Position $k + 1$ beginnt eine maximal bewertete Teilfolge.

Fall 2: Position k befindet sich in einer maximal bewerteten Teilfolge a' . Angenommen a' beginnt an Position $i < k$. Dann gilt aber nach Definition von k , dass $\sigma(i, k - 1) \leq 0$ ist, was im Widerspruch zu Lemma 1.6 Charakterisierung 3 steht. Also beginnt an Position k eine maximal bewertete Teilfolge.

Den Beweis für das Maximum von $\sigma(1, i)$ sei dem Leser überlassen. ■

Im vorigen Beweis kann der zweite Teil nur eintreten, wenn $\sigma(1, 1)$ das eindeutige Minimum von $\sigma(1, i)$ ist.

Wir werden von nun an die folgende algorithmische Idee verfolgen: Sei a' eine Teilfolge von a . Wir werden versuchen a' zu einer maximal bewerteten Teilfolge zu verlängern.

Definition 1.12 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine Folge reeller Zahlen. Eine Teilfolge a' von a heißt a -MSS, wenn $a' \in \text{MSS}(a)$.

Lemma 1.13 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Sei a'' eine Teilfolge von a' und a' eine Teilfolge von a . Ist a'' eine a -MSS, dann ist a'' auch a' -MSS.

Beweis: Sei a'' eine a -MSS und sei sowohl a' eine Teilfolge von a als auch eine Oberfolge von a'' , wie in Abbildung 1.18 dargestellt.

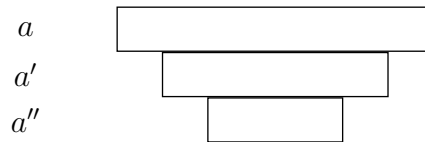


Abbildung 1.18: Skizze: a'' ist a -MSS

Da a'' eine a -MSS ist, erfüllt a'' die Eigenschaften E1 und E2 bezüglich a . Somit erfüllt die Folge a'' diese Eigenschaften auch bezüglich der Oberfolge a' . ■

Lemma 1.14 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Sei $a' = (a_1, \dots, a_i)$ und $a_k \leq 0$ für alle $k \in [i+1 : n]$. Alle maximal bewerteten Teilfolgen von a' sind auch maximal bewertete Teilfolgen von a und umgekehrt.

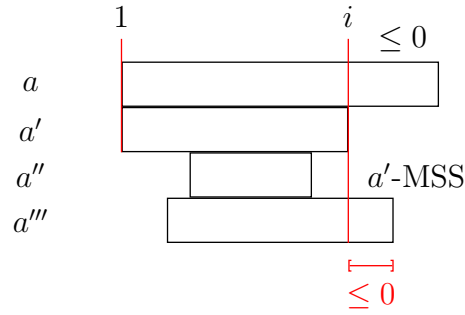
Beweis: Wir beweisen beide Implikationen getrennt.

⇐: Dies ist die Aussage von Lemma 1.13, da nach Lemma 1.6 Charakterisierung 3 jede a -MSS eine Teilfolge von a' sein muss.

⇒: Sei a'' eine a' -MSS, wie in Abbildung 1.19 illustriert.

Angenommen, a'' sei keine a -MSS. Dann muss es eine Oberfolge a''' von a'' geben, die E1 erfüllt. Da a'' eine a' -MSS ist, muss diese Oberfolge a''' in den hinteren Teil (ab Position $i+1$) hineinragen. Eine solche Folge a''' kann nach Lemma 1.6 Charakterisierung 3 nicht E1 erfüllen. Dies führt zu dem gewünschten Widerspruch. ■

Lemma 1.15 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Sei $a' = (a_i, \dots, a_j)$ eine a -MSS und sei $a^L = (a_1, \dots, a_{i-1})$ und $a^R = (a_{j+1}, \dots, a_n)$. Dann ist eine Teilfolge a'' von a^L (bzw. a^R) genau dann eine a^L -MSS (bzw. a^R -MSS), wenn $a'' \neq a'$ eine a -MSS ist.

Abbildung 1.19: Skizze: a'' ist eine a' -MSS

Beweis: Wir beweisen beide Implikationen getrennt.

$\Leftarrow$: Nach Lemma 1.7 sind maximal bewertete Teilfolgen disjunkt. Somit ist jede andere maximal bewertete Teilfolge eine Teilfolge von a^L oder a^R . Mit Lemma 1.13 ist nun jede maximal bewertete Teilfolge von a ungleich a' entweder eine maximal bewertete Teilfolge von a^L oder a^R .

$\Rightarrow$: Wir führen den Beweis durch Widerspruch. Sei dazu a'' eine a^L -MSS, aber keine a -MSS. Somit erfüllt a'' die Eigenschaft E1. Mit Lemma 1.8 folgt, dass es eine Oberfolge $\bar{a}$ von a'' gibt, die eine a -MSS ist.

Angenommen $\bar{a}$ wäre eine Teilfolge von a^L . Dann wäre nach Lemma 1.13 $\bar{a}$ auch eine a^L -MSS. Somit wären $\bar{a}$ und a'' überlappende a^L -MSS, was Lemma 1.7 widerspricht.

Somit müssen sich $\bar{a}$ und a' überlappen. Da aber beide a -MSS sind, ist dies ebenfalls ein Widerspruch zu Lemma 1.7. ■

Damit können wir nun die Äquivalenz der strukturellen und prozeduralen Definition maximal bewerteter Teilfolgen zeigen.

Theorem 1.16 *Die strukturelle und prozedurale Definition von maximal bewerteten Teilfolgen stimmen überein.*

Beweisidee: Sei $MSS(a)$ die Menge aller Teilfolgen von a , die maximal bewertete Teilfolgen sind (gemäß der strukturellen Definition, also die die Eigenschaften E1 und E2 erfüllen).

Sei $a' = (a_i, \dots, a_j) \in MSS(a)$ mit maximalem Score. Nach Lemma 1.15 gilt:

$$MSS(a) = \{a'\} \cup MSS(a_1, \dots, a_{i-1}) \cup MSS(a_{j+1}, \dots, a_n).$$

Der vollständige Beweis lässt sich jetzt formal mittels Induktion führen. ■

1.2.3 Ein Algorithmus zur Lösung

Wir beschreiben jetzt einen Algorithmus zur Ermittlung aller maximal bewerteten Teilfolgen.

- Die Eingabe ist die Folge $a = (a_1, \dots, a_n) \in \mathbb{R}^n$.
- Die Elemente werden von links nach rechts verarbeitet. Dabei merken wir uns disjunkte Teilfolgen $I_1, \dots, I_k$ eines Präfixes a' von a , die maximal bewertete Teilfolgen des bereits abgearbeiteten Präfixes $a' = (a_1, \dots, a_m)$ von a sein werden.
- $I_i = (a_{\ell_i}, \dots, a_{r_i})$, d.h. $I_i \hat{=} (\ell_i, r_i)$.
- Setze $L_i = \sigma(1, \ell_i - 1)$ und $R_i = \sigma(1, r_i)$.
- Damit gilt $\sigma(I_i) := \sigma(\ell_i, r_i) = R_i - L_i$.

Das können wir uns wie in Abbildung 1.20 veranschaulichen.

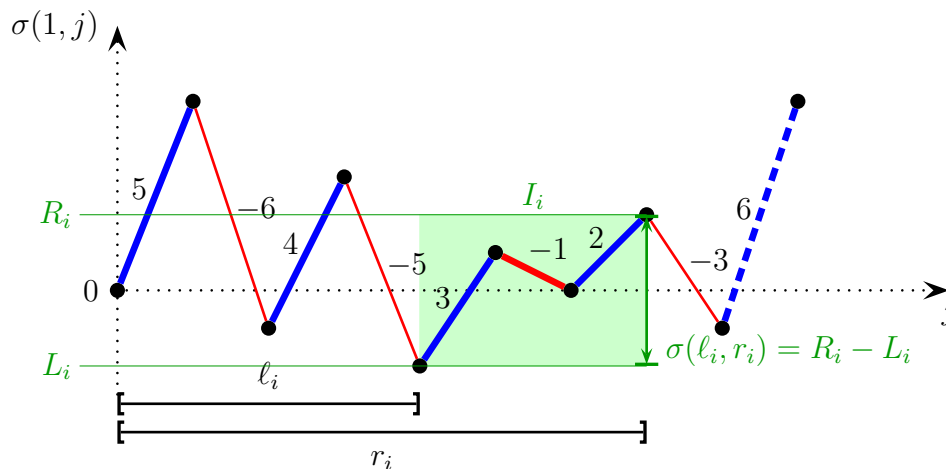


Abbildung 1.20: Skizze: Die Werte ℓ_i , r_i , L_i , R_i und I_i

Wir gehen wie folgt vor: Bearbeite die Folge von links nach rechts und sammle eine Liste von maximal bewerteten Teilfolgen bezüglich des betrachteten Präfixes.

Sei a_m das aktuell betrachtete Element der Folge a :

- Ist $a_m \leq 0$, betrachte die nächste Position $m + 1$, da keine maximal bewertete Teilfolge an Position m beginnen oder enden kann.

- Ist $a_m > 0$, erzeuge eine neue einelementige Liste $I_k = (m, m)$ und bestimme L_k sowie R_k . Die Folge (a_m) erfüllt E1, aber nicht notwendigerweise E2 im betrachteten Präfix $a' = (a_1, \dots, a_m)$ von a .
- Für ein neues I_k wiederhole das Folgende: Durchsuche $I_{k-1}, \dots, I_1$ (von rechts nach links!) bis ein maximales j gefunden wird, so dass $L_j < L_k$ und betrachte die folgenden drei Fälle genauer. Hierbei hierbei gilt immer, dass $I_1, \dots, I_{k-1}$ maximal bewertete Teilfolgen in $(a_1, \dots, a_{m-1})$ sind und I_k zumindest die Eigenschaft E1 erfüllt.

26.10.06

Fall 1: Es existiert gar kein solches j , siehe Abbildung 1.21.

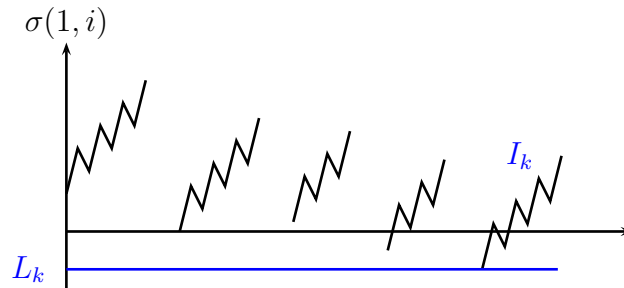


Abbildung 1.21: Skizze: Nichtexistenz von j

Nach Lemma 1.11 ist I_k der Anfang einer maximal bewerteten Teilfolge in der Teilfolge $a' = (a_1, \dots, a_m)$. Nach Lemma 1.15 sind $I_1, \dots, I_{k-1}$ dann auch maximal bewerteten Teilfolgen in a' .

Nach Lemma 1.8 ist I_k auch Teilfolge einer a -MSS a'' . Die Teilfolge I_k ist sogar Anfang von a'' , denn sonst hätten einige Präfixe von a'' einen nichtpositiven Score, was ein Widerspruch zu Lemma 1.6 ist. Somit sind $I_1, \dots, I_{k-1}$ nach Lemma 1.15 auch jeweils eine a -MSS.

Daher geben wir jetzt die Teilfolgen $I_1, \dots, I_{k-1}$ als maximal bewerteten Teilfolgen aus und setzen $I_1 := I_k$ sowie $k = 1$.

Fall 2: Sei j maximal mit $L_j < L_k$ und $R_j < R_k$, siehe Abbildung 1.22. Offensichtlich gilt nach Wahl von j , dass $L_i \geq L_k$ für $i \in [j + 1 : k - 1]$.

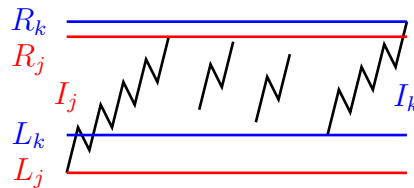


Abbildung 1.22: j maximal mit $L_j < L_k$ und $R_j < R_k$

Weiterhin gilt $R_i \leq R_j$ für alle $i \in [j + 1 : k - 1]$. Andernfalls gäbe es ein $i \in [r_j + 1 : \ell_k - 1]$ mit $\sigma(1, i) > R_j$. Sei i ein minimaler solcher Index. Dann besäße I_j die Oberfolge $(a_{\ell_j}, \dots, a_{r_i})$, die offensichtlich die Eigenschaft E1 erfüllt, was der Eigenschaft E2 von I_j in $(a_1, \dots, a_{m-1})$ widerspricht.

Somit erfüllt die Teilfolge $(a_{\ell_j}, \dots, a_{r_k})$ die Eigenschaft E1, aber nicht notwendigerweise E2. Somit muss ein neues Intervall generiert und angefügt werden. Wir bilden aus den Teilfolgen $I_j, \dots, I_k$ ein neues Intervall, d.h. wir setzen $I_j := (\ell_j, r_k)$ sowie $k := j$ und wiederholen die Prozedur für $I_k = I_j$ (also mit dem neuen k). Alle 3 Fälle sind jetzt wieder neu zu betrachten.

Fall 3: Sei j maximal mit $L_j < L_k$ und $R_j \geq R_k$, siehe Abbildung 1.23.

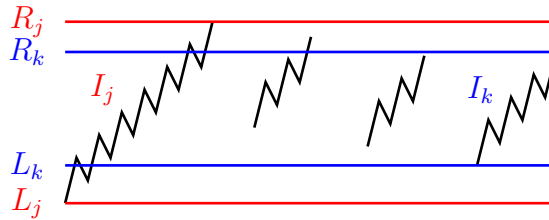


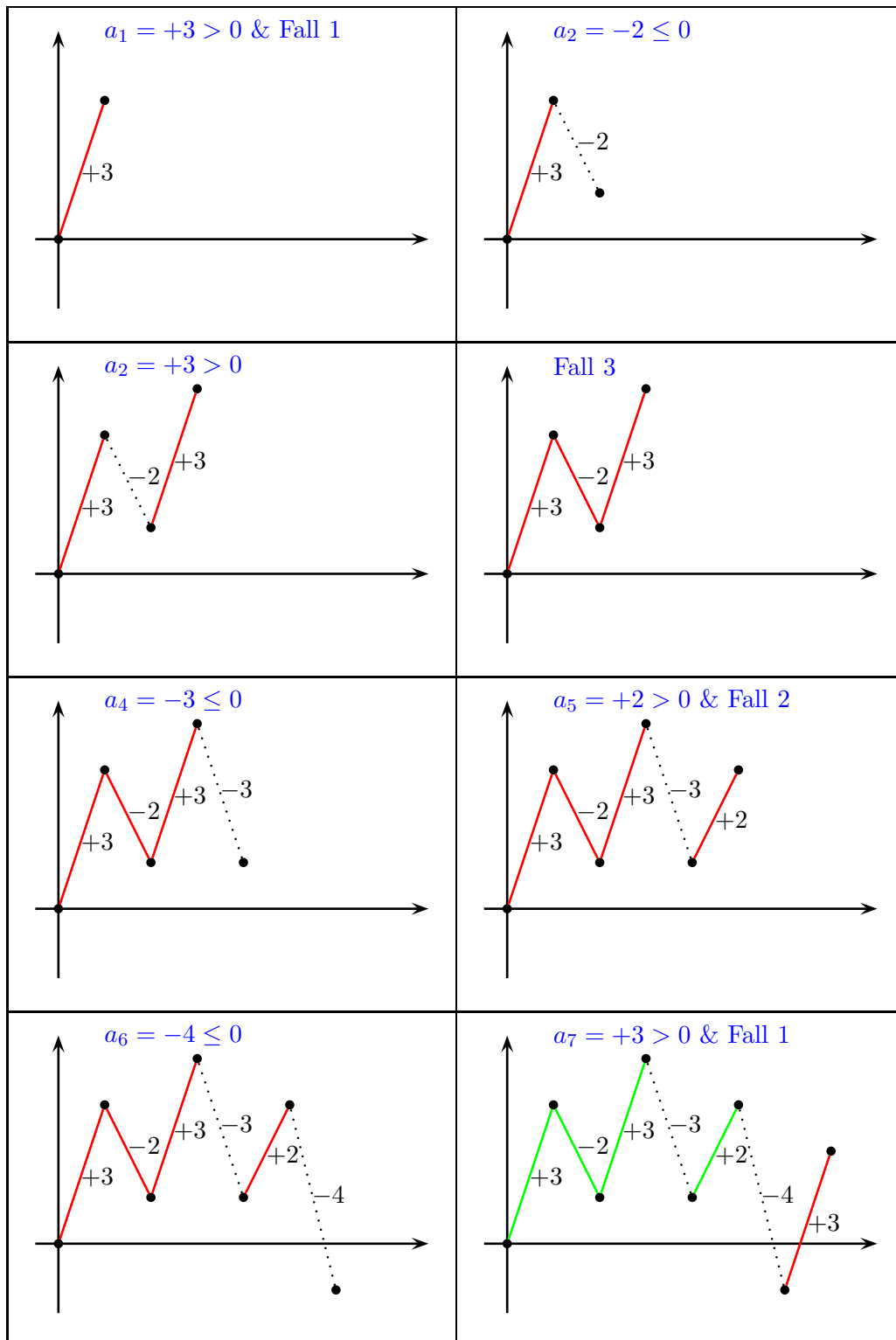
Abbildung 1.23: Skizze: j maximal mit $L_j < L_k$ und $R_j \geq R_k$

Wir zeigen zunächst, dass dann I_j eine a' -MSS mit $a' = (a_1, \dots, a_m)$ ist.

- E1) Dies gilt offensichtlich nach Konstruktion.
- E2) Für einen Widerspruchsbeweis nehmen wir an, es gäbe eine Oberfolge $a'' = (a_i, \dots, a_s)$ von I_j , die Eigenschaft E1 erfüllt.
Gilt $r_j \leq s < \ell_k$, dann muss nach Charakterisierung 2 von Lemma 1.6 $\sigma(1, s) > j$ gelten. Somit wäre a'' auch eine Oberfolge von I_j , die in $(a_1, \dots, a_{\ell_k-1})$ liegt und somit wäre I_j keine maximal bewertete Teilfolge von $(a_1, \dots, a_{m-1})$ gewesen, was nach Konstruktion ein Widerspruch ist. Also muss $s \in [\ell_k : r_k]$ gelten. Dann gibt es aber offensichtlich Suffixe von a'' mit nichtpositiven Score, was nach Lemma 1.6 Charakterisierung 3) nicht sein kann.

Jetzt müssen wir noch zeigen, dass $I_k = (a_{\ell_k}, \dots, a_{r_k})$ eine maximal bewertete Teilfolge von $a' = (a_1, \dots, a_m)$ ist. Da I_j eine maximal bewertete Teilfolge von a' ist genügt es nach Lemma 1.15 zu zeigen, dass I_k eine maximal bewertete Teilfolge von $a'' = (a_{r_j+1}, \dots, a_m)$ ist. Aus Lemma 1.11 folgt aber unmittelbar, dass I_k eine a'' -MSS ist und somit die Behauptung. In diesem Fall ist algorithmisch also gar nichts zu tun.

In jedem Fall wird das Durchsuchen der Listen $I_{k-1}, \dots, I_1$ entweder erneut aufgerufen oder aber alle aufgesammelten Teilfolgen $I_1, \dots, I_k$ sind maximal bewertete Teilfolgen des Präfixes $a' = (a_1, \dots, a_m)$ von a .

Abbildung 1.24: Beispiel: $a = (+3, -2, +3, -3, +2, -4, +3)$

In Abbildung 1.24 auf Seite 21 ist ein Ablauf des gerade beschriebenen AMSS-Algorithmus für die Sequenz $a = (+3, -2, +3, -3, +2, -4, +3)$ schematisch dargestellt. Hierbei entsprechen die roten Linienzüge den gespeicherten disjunkten Teilfolgen $I_1, \dots, I_k$, während die grünen Linienzüge die nach Fall 1 ausgegeben maximal bewerteten Teilfolgen darstellen.

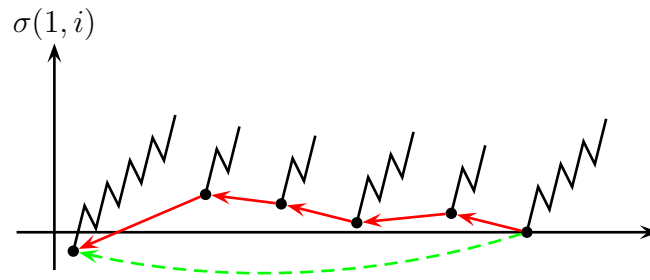
1.2.4 Zeitkomplexität

Wir analysieren jetzt die Zeitkomplexität des soeben vorgestellten Algorithmus:

- 1.) Die Betrachtung von a_m und der eventuell Generierung von I_{k+1} kann pro Folgeelement in konstanter Zeit durchgeführt werden. Daraus ergibt sich ein Gesamtaufwand von $O(n)$.
- 2.) Das Durchsuchen der bislang gesammelten Teilfolgen $I_{k-1}, \dots, I_1$ wird anschließend analysiert.
- 3.) Gesamtaufwand von Fall 1: Es existiert kein j , dann geben wir $I_1, \dots, I_{k-1}$ aus. Dies lässt sich für jedes I_ℓ in konstanter Zeit erledigen, also erhalten wir insgesamt für alle Ausgaben einen Zeitbedarf von $O(n)$ (da es nur $O(n)$ maximal bewertete Teilfolgen geben kann).
- 4.) Gesamtaufwand von Fall 3: Es gilt $L_j < L_k \wedge R_j \geq R_k$. Die Ermittlung des zweiten Falles geschieht in konstanter Zeit. Da anschließend das Durchsuchen beendet ist, tritt dieser Fall maximal n -mal auf, also ist der Aufwand für alle diese Fälle höchstens $O(n)$.
- 5.) Gesamtaufwand von Fall 2: Es gilt $L_j < L_k \wedge R_j < R_k$. Die Verschmelzung von $I_j, \dots, I_k$ lässt sich in Zeit $O(k - j + 1)$ erledigen (Explizites Löschen der Teilfolgen). Da insgesamt höchstens $n - 1$ Verschmelzungen von einzelnen Intervallen möglich sind, beträgt der Gesamtzeitbedarf $O(n)$.

Dabei stellen wir uns eine Verschmelzung von ℓ Intervallen als $\ell - 1$ Verschmelzungen von je zwei Intervallen vor. Da es insgesamt maximal $n - 1$ Verschmelzungen disjunkter Intervalle geben kann, folgt obige Behauptung.

Durchsuchen der Listen: Beim Durchsuchen der Listen von Intervallen gehen wir etwas geschickter vor. Wir merken uns für jedes Intervall I_k dabei das Intervall I_j , für das $L_j < L_k$ und j maximal ist. Beim nächsten Durchlaufen können wir dann einige Intervalle beim Suchen einfach überspringen. Dies ist in der folgenden Abbildung 1.25 schematisch dargestellt.

Abbildung 1.25: Skizze: Übersprungene Folge von L_j -Werten beim Durchsuchen

Jetzt müssen wir uns nur noch überlegen, welchen Aufwand alle Durchsuche-Operationen insgesamt haben. Dabei werden die Kosten zum einen auf die Aufrufe (der Durchsucheprozedur) und zum anderen auf die ausgeschlossenen Intervalle verteilt.

Bei einem Durchsuchen der Liste werden beispielsweise ℓ Intervalle übersprungen. Dann wurden $\ell + 2$ Listen inspiziert, was Kosten in Höhe von $O(\ell + 1)$ verursacht (man beachte, dass $\ell \geq 0$ gilt). Die Kosten der Intervalle, die dabei ausgeschlossen werden, werden einfach anteilig auf die ausgeschlossenen Intervalle umgelegt. Somit erhält jedes ausgeschlossene Intervall Kosten von $O(1)$. Die Kosten für den Rest werden auf den Aufruf eines Durchsuch-Durchlaufs umgelegt, die dann ebenfalls konstant sind.

Somit fallen jeweils Kosten $O(1)$ auf ausgeschlossene Intervalle und Aufrufe an.

Anzahl Aufrufe: Es kann maximal so viele Aufrufe geben, wie der Zeitbedarf in der Analyse der Punkte 3, 4 und 5 verbraucht wird, also $O(n)$.

Anzahl ausgeschlossener Intervalle: Zum einen können diese im Fall 1 ausgegeben werden, von diesen kann es daher ebenfalls maximal $O(n)$ viele geben.

Wenn diese Intervalle nicht selbst ausgegeben werden, müssen diese irgendwann mit anderen Intervallen verschmolzen worden sein (sie können ja nicht einfach verschwinden). Da, wie wir bereits gesehen haben, maximal $O(n)$ Intervalle verschmolzen werden, können auch hierdurch höchstens $O(n)$ Intervalle ausgeschlossen werden.

Mit unserem kleinen Trick kann also auch das Durchsuchen der Intervall-Listen mit einem Aufwand von insgesamt $O(n)$ bewerkstelligt werden. Halten wir das Resultat noch im folgenden Satz fest.

Theorem 1.17 *Das All Maximal Scoring Subsequences Problem für eine gegebene reelle Folge kann in Linearzeit gelöst werden.*

1.3 Bounded All Maximum Scoring Subsequences

Im vorherigen Abschnitt haben wir quasi über eine lokale Definition die Menge aller maximal bewerteten Teilfolgen bestimmt. Wir können auch einen globalen Ansatz wählen und eine Menge von Teilfolgen einer gegebenen Folge fordern, deren aufaddierte Scores maximal ist. Dies ist jedoch nicht sinnvoll, da dann die Menge aller einelementigen Teilfolgen, die positive Elemente beschreiben, eine Lösung dieses Problems ist. Man überlegt sich leicht, dass dies wirklich ein globales Optimum ist.

Daher wollen wir in diesem Abschnitt das Problem ein wenig abwandeln, in dem wir die Länge der Teilfolgen in der Lösungsmenge sowohl von oben als auch von unten beschränken.

1.3.1 Problemstellung

Es wird zusätzlich eine untere $\underline{\lambda}$ und eine obere Schranke $\overline{\lambda}$ vorgegeben, um die Länge der zu betrachtenden Teilfolgen zu beschränken, wobei natürlich $\underline{\lambda} \leq \overline{\lambda} \in \mathbb{N}$ gilt. Mit $Seq(n, k)$ bezeichnen wir die Menge aller 0-1-Zeichenreihen mit genau k konsekutiven 1-Runs, deren Länge durch $\underline{\lambda}$ nach unten und mit $\overline{\lambda}$ nach oben beschränkt ist.

Notation 1.18 Sei $k \leq n \in \mathbb{N}$ und $\underline{\lambda} \leq \overline{\lambda} \in \mathbb{N}$, dann bezeichne

$$Seq(n, k) := \{0^* s_1 0^+ s_2 0^+ \cdots 0^+ s_k 0^* \mid s_i \in 1^*, \underline{\lambda} \leq |s_i| \leq \overline{\lambda}\} \cap \{0, 1\}^n \subseteq \{0, 1\}^n.$$

BOUNDED ALL MAXIMUM SCORING SUBSEQUENCES (BAMSS)

Eingabe: Eine Folge $(a_1, \dots, a_n) \in \mathbb{R}^n$, $\underline{\lambda} \leq \overline{\lambda} \in \mathbb{N}$.

Gesucht: Eine Sequenz $s \in \bigcup_{k=0}^n Seq(n, k)$, die $\sum_{i=1}^n s_i \cdot a_i$ maximiert.

Bemerkung: Durch die Längenbeschränkung können nach unserer alten Definition von allen Maximal Scoring Subsequences die einzelnen Teilfolgen der Lösung überlappen, siehe Abbildung 1.26. Damit ist ein Greedy-Ansatz nicht mehr effizient

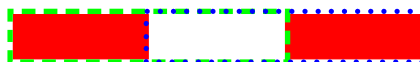


Abbildung 1.26: Skizze: Erlaubte überlappende Sequenzen

möglich. Aus diesem Grund wurde das Problem anders gefasst, nämlich wie oben. Mit Hilfe der Menge $Seq(n, k)$ werden aus der Gesamtfolge Teilstücke ausgewählt (nämlich 1-Runs), die wir als Lösungsteilfolgen zulassen wollen.

1.3.2 Lösung mittels Dynamischer Programmierung

Wir definieren wieder eine Tabelle S für $i, k \in [0 : n]$ wie folgt:

$$S(i, k) := \max \left\{ \sum_{j=1}^i s_j \cdot a_j : s \in Seq(i, k) \right\}$$

Wir müssen also $S(i, k)$ für alle Werte $i \in [0 : n]$, $k \in [0 : n]$ berechnen, wobei $\max \emptyset = \max \{ \} = -\infty$ (entspricht dem neutralen Element) gilt.

Für diese Tabelle S lässt sich die folgende Rekursionsgleichung aufstellen:

$$\begin{aligned} S(i, 0) &= 0 \quad \text{für } i \in [0 : n], \\ S(i, k) &= -\infty \quad \text{für } i < k \cdot \underline{\lambda} + (k - 1), \\ S(i, k) &= \max \left\{ \begin{array}{l} S(i - 1, k), \\ S(i - \lambda - 1, k - 1) + \sum_{j=i-\lambda+1}^i a_j \end{array} : \lambda \in [\underline{\lambda} : \min(\overline{\lambda}, i - 1)] \right\}. \end{aligned}$$

Die Korrektheit der Rekursionsgleichung ergibt sich aus der Tatsache, dass man sich eine optimale Menge von Teilfolgen anschaut und unterscheidet, ob diese mit einer 0 oder einem 1-Run endet, wie in Abbildung 1.27 illustriert.

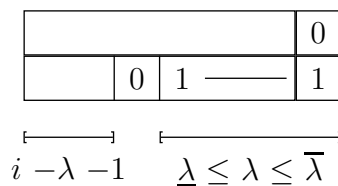


Abbildung 1.27: Skizze: Aufstellen der Rekursionsgleichung

Wie man leicht sieht, gilt:

$$\sum_{j=i-\lambda+1}^i a_j = \sum_{j=1}^i a_j - \sum_{j=1}^{i-\lambda} a_j$$

Somit kann man die einzelnen Summen effizient berechnen, wenn für alle $i \in [0 : n]$ die folgenden Summen bekannt sind: $A(i) := \sum_{j=1}^i a_j$. Berechnet man die Werte $A(i)$

für $i \in [0 : n]$ vorab, so lässt sich jeder Eintrag der Tabelle S in Zeit $O(\bar{\lambda} - \underline{\lambda}) = O(n)$ berechnen.

Die Gesamtlaufzeit beträgt somit $O(n^2(\bar{\lambda} - \underline{\lambda})) \leq O(n^3)$. In der Praxis sind die Schranken $\underline{\lambda}$ und $\bar{\lambda}$ allerdings so nah beieinander, so dass $\bar{\lambda} - \underline{\lambda} = O(1)$ gilt und somit die Laufzeit bei $O((\bar{\lambda} - \underline{\lambda})n^2) = O(n^2)$ bleibt.

Theorem 1.19 *Das BAMSS-Problem lässt sich in Zeit $O(n^2(\bar{\lambda} - \underline{\lambda}))$ lösen.*

Wie wir im nächsten Abschnitt sehen werden, können wir das Problem auch effizienter lösen. Die vorhergehende Lösung hat den Vorteil dass man auch die Maximalanzahl der Teilfolgen einer Lösung beschränken kann, was durchaus biologisch sinnvoll sein kann.

1.3.3 Effiziente Lösung mittels Dynamischer Programmierung

Wir definieren jetzt folgende Tabelle S für $i \in [0 : n]$ wie folgt:

$$S(i) := \max \left\{ \sum_{j=1}^i s_j \cdot a_j : s \in \text{Seq}(i, k) \wedge k \in [0 : n] \right\}$$

Wir müssen also in diesem Fall $S(i)$ für alle Werte $i \in [0 : n]$ berechnen, wobei $\max \emptyset = \max \{ \} = -\infty$ (entspricht dem neutralen Element) gilt.

Für diese Tabelle S lässt sich die folgende Rekursionsgleichung aufstellen:

$$\begin{aligned} S(i) &= 0 \quad \text{für } i \in [0 : \underline{\lambda} - 1], \\ S(i) &= \max \left\{ \begin{array}{l} S(i-1), \\ S(i - \lambda - 1) + \sum_{j=i-\lambda+1}^i a_j \end{array} : \lambda \in [\underline{\lambda} : \bar{\lambda}] \right\} \cup \left\{ \sum_{j=1}^i : i \leq \bar{\lambda} \right\}. \end{aligned}$$

Hierbei gilt die zweite Rekursionsgleichung nur für $i \geq \underline{\lambda}$. Die Korrektheit der Rekursionsgleichung ergibt sich wie vorher.

Die Gesamtlaufzeit beträgt somit $O(n(\bar{\lambda} - \underline{\lambda})) \leq O(n^2)$. In der Praxis sind die Schranken $\underline{\lambda}$ und $\bar{\lambda}$ allerdings so nah beieinander, so dass $\bar{\lambda} - \underline{\lambda} = O(1)$ gilt und somit die Laufzeit bei $O((\bar{\lambda} - \underline{\lambda})n) = O(n)$ bleibt.

Theorem 1.20 *Das BAMSS-Problem lässt sich in Zeit $O(n(\bar{\lambda} - \underline{\lambda}))$ lösen.*

1.4 Bounded Maximal Scoring Subsequence

Jetzt wollen wir nur *eine* längenbeschränkte Maximal Scoring Subsequence finden. Man beachte, dass diese kein Teil der Lösung aus dem Problem des vorherigen Abschnittes sein muss! Wir werden zunächst nur eine obere Längenbeschränkung für die gesuchte Folge betrachten. Eine Hinzunahme einer unteren Längenbeschränkung ist nicht weiter schwierig und wird in den Übungen behandelt.

1.4.1 Problemstellung

Wir formalisieren zunächst die Problemstellung.

BOUNDED MAXIMAL SCORING SUBSEQUENCE (BMSS)

Eingabe: Eine Folge $(a_1, \dots, a_n) \in \mathbb{R}^n$ und $\lambda \in \mathbb{N}$.

Gesucht: Eine (zusammenhängende) Teilfolge $(a_i, \dots, a_j)$ mit $j - i + 1 \leq \lambda$, die $\sigma(i, j)$ maximiert, wobei $\sigma(i, j) = \sum_{\ell=i}^j a_\ell$.

1.4.2 Links-Negativität

Um einer effizienten Lösung des Problems näher zu kommen, benötigen wir zuerst den Begriff der Links-Negativität und einer minimalen links-negativen Partition einer reellen Folge.

Definition 1.21 Eine Folge $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ reeller Zahlen heißt links-negativ, wenn $\sum_{\ell=1}^k a_\ell \leq 0$ für alle $k \in [1 : n - 1]$.

Eine Partition $a = A_1 \cdots A_k$ der Folge a heißt minimal links-negativ, wenn für alle $i \in [1 : k]$ A_i links-negativ ist und $\sigma(A_i) > 0$ für alle $i \in [1 : k - 1]$ ist.

Beispiele:

- 1.) $(-1, 1, -3, 1, 1, 3)$ ist links-negativ,
- 2.) $(2, -3, -4, 5, 3, -3)$ ist **nicht** links-negativ.

Die Partition $(2)(-3, -4, 5, 3)(-3)$ ist eine minimal links-negative.

Lemma 1.22 Jede Folge $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ lässt sich eindeutig in eine minimal links-negative Partition zerlegen.

Beweis: Wir führen den Beweis durch Induktion über n .

Induktionsanfang ($n = 1$): Dies folgt unmittelbar aus der Definition.

Induktionsschritt ($n \rightarrow n + 1$): Betrachte $a' = (a_0, a_1, \dots, a_n)$ und sei dann $a = (a_1, \dots, a_n)$. Nach Induktionsvoraussetzung sei $a = A_1 \cdots A_k$ die(!) minimal links-negative Partition von a .

Fall 1 ($a_0 > 0$): Indem wir a_0 einfach als Segment an die minimal links-negative Partition von a voranstellen, erhalten wir eine minimal links-negative Partition für $a' = (a_0) \cdot A_1 \cdots A_k$.

Fall 2 ($a_0 \leq 0$): Wähle ein minimales i mit

$$a_0 + \sum_{j=1}^i \sigma(A_j) > 0.$$

Dann ist $((a_0) \cdot A_1 \cdots A_i) \cdot A_{i+1} \cdots A_k$ eine minimal links-negative Partition.

Der Beweis der Eindeutigkeit sei dem Leser zur Übung überlassen. ■

Im Folgenden ist die Notation für die minimale links-negative Partition der Suffixe der untersuchten Folge hilfreich.

Notation 1.23 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Jedes Suffix $a^{(i)} = (a_i, \dots, a_n)$ von a besitzt dann eine eindeutige minimal links-negative Partition $a^{(i)} = A_1^{(i)} \cdots A_{k_i}^{(i)}$.

Basierend auf dieser Notation lassen sich links-negative Zeiger definieren.

Definition 1.24 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine reelle Folge und für jedes Suffix ab Position i sei $a^{(i)} = A_1^{(i)} \cdots A_{k_i}^{(i)}$ die zugehörige minimale links-negative Partition. Sei weiter $A_1^{(i)} = (a_i, \dots, a_{p(i)})$, dann heißt $p(i)$ der links-negative Zeiger von i .

Bemerkungen:

- Ist $a_i > 0$, dann ist $p(i) = i$.
- Ist $a_i \leq 0$, dann ist $p(i) > i$ für $i < n$ und $p(n) = n$.

1.4.3 Algorithmus zur Lösung des BMSS-Problems

Wie helfen uns minimale links-negative Partitionen bei der Lösung des Problems? Wenn wir eine Teilfolge $a' = (a_i, \dots, a_j)$ als Kandidaten für eine längenbeschränkte MSS gefunden haben, dann erhalten wir den nächstlängeren Kandidaten durch das Anhängen der Teilfolge $(a_{j+1}, \dots, a_{p(j+1)})$. Nach Definition hätte jedes kürzere Stück einen nichtpositiven Score und würde uns nicht weiter helfen. Wir müssen dann nur noch prüfen, ob die obere Längenbeschränkung eingehalten wird. Dies ist in der folgenden Abbildung 1.28 veranschaulicht.

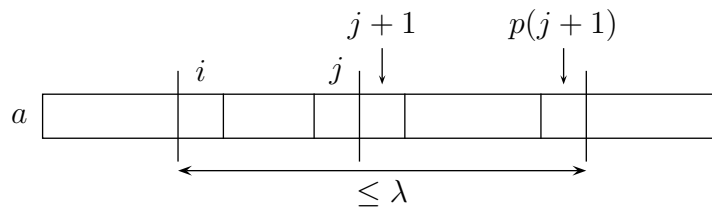


Abbildung 1.28: Skizze: Verlängerung einer maximal Scoring Subsequence

Mit Hilfe der links-negativen Zeiger können wir das Problem mit dem in Abbildung 1.29 angegebenen Algorithmus leicht lösen. Wir versuchen hier ab jeder Position $i \in [1 : n]$ eine Maximal Scoring Subsequence der Länge höchstens λ zu finden.

BMSS (real $a[]$, int n , λ)

```
{
    int ms = 0;    mi = 1;    mj = 0;    /* score and position of current mss */
    int j = 1;    /* (a_i, ..., a_j) is current candidate subsequence */
    for (i = 1; i ≤ n; i++)
    {
        while ((a_i ≤ 0) && (i < n))
            i++;
        j := max(i, j)
        while ((j < n) && (p[j + 1] < i + λ) && (σ(j + 1, p(j + 1)) > 0))
            j = p(j + 1)
        if (σ(i, j) > ms)
        {
            mi = i;    mj = j;    ms = σ(i, j);
        }
    }
}
```

Abbildung 1.29: Algorithmus: Lösung des BMSS Problems

Für die Laufzeit halten wir das Folgende fest. Nach maximal einer konstanten Anzahl von Operationen wird entweder i oder j erhöht, somit ist die Laufzeit $O(n)$.

Wir haben jedoch die links-negativen Zeiger noch nicht berechnet. Wir durchlaufen dazu die Folge vom Ende zum Anfang hin. Treffen wir auf ein Element $a_i > 0$, dann sind wir fertig. Andernfalls verschmelzen wir das Element mit den folgenden Segmenten, bis das neu entstehende Segment einen positiven Score erhält. Dies ist in der folgenden Abbildung 1.30 illustriert, wobei nach dem Verschmelzen mit $A_j^{(i)}$ das Segment ab Position i einen positiven Score erhält.



Abbildung 1.30: Skizze: Verschmelzen eines Elements mit Segmenten

Wir können diese Idee in den in Abbildung 1.31 angegebenen Algorithmus umsetzen, wobei $s(i) := \sigma(i, p(i))$ bezeichnet.

```

COMPUTE_LEFT_NEGATIVE_POINTER (real  $a[]$ , int  $n$ )
{
    for ( $i = n$ ;  $i > 0$ ;  $i--$ )
    {
         $p(i) = i$ ;
         $s(i) = a_i$ ;
        while ( $(p(i) < n) \ \&\& \ (s(i) \leq 0)$ )
        {
             $s(i) = s(i) + s(p(i) + 1)$ ;
             $p(i) = p(p(i) + 1)$ ;
        }
    }
}

```

Abbildung 1.31: Algorithmus: Berechnung links-negativer Zeiger durch Verschmelzen von Segmenten

Eine simple Laufzeit-Abschätzung ergibt, dass die Laufzeit im schlimmsten Fall $O(n^2)$ beträgt. Mit einer geschickteren Analyse können wir jedoch wiederum eine lineare Laufzeit zeigen. Es wird pro Iteration maximal ein neues Segment generiert. In jeder inneren 'while-Schleife' wird ein Segment eliminiert. Es müssen mindestens so viele Segmente generiert wie gelöscht (= verschmolzen) werden. Da nur n Segmente

generiert werden, können auch nur n Segmente gelöscht werden und somit ist die Laufzeit $O(n)$. Halten wir das Resultat im folgenden Satz fest.

Theorem 1.25 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine Folge reeller Zahlen und $\lambda \in \mathbb{N}$. Dann lässt sich eine Maximal Scoring Subsequence der Länge höchstens λ in Zeit $O(n)$ finden.

Es gilt sogar der folgende Satz.

Theorem 1.26 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine Folge reeller Zahlen und $\underline{\lambda} \leq \bar{\lambda} \in \mathbb{N}$. Dann lässt sich eine Maximal Scoring Subsequence der Länge mindestens $\underline{\lambda}$ und höchstens $\bar{\lambda}$ in Zeit $O(n)$ finden.

Beweis: Der Beweis sei dem Leser zur Übung überlassen. ■

02.11.06

1.5 Maximal Average Scoring Subsequence (*)

Jetzt wollen wir eine Teilfolge finden, deren Mittelwert (gemittelt über die Länge) maximal ist.

1.5.1 Problemstellung

Zunächst führen wir noch einige abkürzende Notationen ein.

Notation 1.27 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$, dann ist

$$\sigma(i, j) := \sum_{k=i}^j a_k, \quad \ell(i, j) := j - i + 1, \quad \mu(i, j) := \frac{\sigma(i, j)}{\ell(i, j)}.$$

Damit lässt sich nun das Problem formalisieren.

MAXIMAL AVERAGE SCORING SUBSEQUENCE (MASS)

Eingabe: Eine Folge $a = (a_1, \dots, a_n) \in \mathbb{R}^n$, $\lambda \in \mathbb{N}$.

Gesucht: Eine (zusammenhängende) Teilfolge $(a_i, \dots, a_j)$ mit $\ell(i, j) \geq \lambda$, die $\mu(i, j)$ maximiert, d.h. $\mu(i, j) = \max\{\mu(i', j') \mid j' \geq i' + \lambda - 1\}$.

Wozu haben wir hier noch die untere Schranke λ eingeführt? Ansonsten wird das Problem trivial, denn dann ist das maximale Element, interpretiert als eine elementige Folge, die gesuchte Lösung!

1.5.2 Rechtsschiefe Folgen und fallend rechtsschiefe Partitionen

Zur Lösung des Problems benötigen wir den Begriff einer rechtsschiefen Folge sowie einer fallend rechtsschiefen Partition.

Definition 1.28 Eine Folge $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ heißt rechtsschief, wenn für alle $i \in [1 : n - 1]$ gilt: $\mu(1, i) \leq \mu(i + 1, n)$.

Eine Partition $a = A_1 \cdots A_k$ von a heißt fallend rechtsschief, wenn jedes Segment A_i rechtsschief ist und $\mu(A_i) > \mu(A_j)$ für $i < j$ gilt.

Anschaulich heißt eine Folge $a = (a_1, \dots, a_n)$ also rechtsschief, wenn der Durchschnittswert jedes echten Präfix $(a_1, \dots, a_i)$ kleiner gleich dem Durchschnittswert des korrespondierenden Suffixes $(a_{i+1}, \dots, a_n)$ ist.

Lemma 1.29 Seien $a \in \mathbb{R}^n$ und $b \in \mathbb{R}^m$ zwei reelle Folgen mit $\mu(a) < \mu(b)$. Dann gilt $\mu(a) < \mu(ab) < \mu(b)$ (ab sei die konkatenierte Folge).

Beweis: Es gilt:

$$\begin{aligned} \mu(ab) &= \frac{\sigma(ab)}{\ell(ab)} \\ &= \frac{\sigma(a) + \sigma(b)}{\ell(ab)} \\ &= \frac{\mu(a) \cdot \ell(a) + \mu(b) \cdot \ell(b)}{\ell(ab)}. \end{aligned}$$

Somit ist zum einen (da $\mu(b) > \mu(a)$ ist)

$$\begin{aligned} \mu(ab) &= \frac{\mu(a) \cdot \ell(a) + \mu(b) \cdot \ell(b)}{\ell(ab)} \\ &> \frac{\mu(a) \cdot \ell(a) + \mu(a) \cdot \ell(b)}{\ell(ab)} \\ &= \frac{\mu(a) \cdot (\ell(a) + \ell(b))}{\ell(ab)} \\ &= \mu(a) \end{aligned}$$

und zum anderen (da $\mu(a) < \mu(b)$ ist)

$$\begin{aligned}
 \mu(ab) &= \frac{\mu(a) \cdot \ell(a) + \mu(b) \cdot \ell(b)}{\ell(ab)} \\
 &< \frac{\mu(b) \cdot \ell(a) + \mu(b) \cdot \ell(b)}{\ell(ab)} \\
 &= \frac{\mu(b) \cdot (\ell(a) + \ell(b))}{\ell(ab)} \\
 &= \mu(b).
 \end{aligned}$$

Damit ist die Behauptung gezeigt. ■

Korollar 1.30 Seien $a \in \mathbb{R}^n$ und $b \in \mathbb{R}^m$ zwei reelle Folgen mit $\mu(a) \leq \mu(b)$. Dann gilt $\mu(a) \leq \mu(ab) \leq \mu(b)$.

Beweis: Gilt $\mu(a) < \mu(b)$, so ist es gerade die abgeschwächte Aussage des vorherigen Lemmas. Gilt $\mu(a) = \mu(b)$, so gilt trivialerweise $\mu(a) = \mu(ab) = \mu(b)$. ■

Aus dem Beweis des vorherigen Lemmas folgt sofort auch das folgende Korollar.

Korollar 1.31 Seien $a \in \mathbb{R}^n$ und $b \in \mathbb{R}^m$ zwei reelle Folgen mit $\mu(a) > \mu(b)$. Dann gilt $\mu(a) > \mu(ab) > \mu(b)$.

Lemma 1.32 Seien $a \in \mathbb{R}^n$ und $b \in \mathbb{R}^m$ zwei rechtsschiefe Folgen mit $\mu(a) \leq \mu(b)$. Dann ist auch die Folge ab rechtsschief.

Beweis: Sei p ein beliebiges Präfix von ab . Es ist zu zeigen, dass $\mu(p) \leq \mu(q)$ ist, wobei q das zu p korrespondierende Suffix in ab ist.

Fall 1: Sei $p = a$. Dann gilt $\mu(a) \leq \mu(b)$ nach Voraussetzung, und die Behauptung ist erfüllt, da $q = b$ ist.

Fall 2: Sei jetzt p ein echtes Präfix von a , d.h. $a = pa'$. Da a eine rechtsschiefe Folge ist, gilt $\mu(p) \leq \mu(a')$. Mit dem vorherigen Korollar gilt dann:

$$\mu(p) \leq \mu(pa') \leq \mu(a').$$

Somit ist $\mu(p) \leq \mu(a) \leq \mu(b)$.

Mit $\mu(p) \leq \mu(a')$ und $\mu(p) \leq \mu(b)$, folgt

$$\begin{aligned}
 \mu(p) &= \frac{\ell(a')\mu(p) + \ell(b)\mu(p)}{\ell(a'b)} \\
 &\leq \frac{\ell(a')\mu(a') + \ell(b)\mu(b)}{\ell(a'b)} \\
 &= \frac{\sigma(a') + \sigma(b)}{\ell(a'b)} \\
 &= \mu(a'b) \\
 &= \mu(q).
 \end{aligned}$$

Fall 3: Sei $p = ab'$ mit $b = b'b''$. Mit dem vorherigen Korollar folgt, da $b = b'b''$ rechtsschief ist:

$$\mu(b') \leq \underbrace{\mu(b'b'')}_{=b} \leq \mu(b'').$$

Somit gilt $\mu(a) \leq \mu(b) \leq \mu(b'')$. Also gilt mit $\mu(a) \leq \mu(b'')$ und $\mu(b') \leq \mu(b'')$:

$$\begin{aligned}
 \mu(p) &= \mu(ab') \\
 &= \frac{\sigma(a) + \sigma(b')}{\ell(ab')} \\
 &= \frac{\ell(a)\mu(a) + \ell(b')\mu(b')}{\ell(ab')} \\
 &\leq \frac{\ell(a)\mu(b'') + \ell(b')\mu(b'')}{\ell(ab')} \\
 &= \mu(b'') \\
 &= \mu(q).
 \end{aligned}$$

Damit ist das Lemma bewiesen. ■

Lemma 1.33 Jede Folge $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ besitzt eine eindeutig fallend rechtsschiefe Partition.

Beweis: Wir führen den Beweis durch Induktion über n .

Induktionsanfang ($n = 1$): Klar, nach Definition.

Induktionsschritt ($n \rightarrow n + 1$): Wir betrachten eine Folge $(a_1, \dots, a_n, a_{n+1})$. Sei weiter $A_1 \cdots A_k$ die fallend rechtsschiefe Partition von $(a_1, \dots, a_n)$.

Gilt $\mu(A_k) > a_{n+1} = \mu(a_{n+1})$, dann ist $A_1 \cdots A_k \cdot (a_{n+1})$ eine neue rechtsschiefe Partition.

Andernfalls bestimmen wir ein maximales i mit $\mu(A_{i-1}) > \mu(A_i \cdots A_k \cdot (a_{n+1}))$. Dann behaupten wir, dass die neue Partition $A_1 \cdots A_{i-1} \cdot (A_i \cdots A_k \cdot (a_{n+1}))$ eine fallend rechtsschiefe ist.

Nach Definition ist (a_{n+1}) rechtsschief. Nach dem vorherigen Lemma 1.32 und der Wahl von i ist dann auch $A_k \cdot (a_{n+1})$ rechtsschief. Weiter gilt allgemein für $j \geq i$, dass $A_j \cdots A_k \cdot (a_{n+1})$ rechtsschief ist. Somit ist auch $A_i \cdots A_k \cdot (a_{n+1})$ rechtsschief.

Nach Konstruktion ist also die jeweils konstruierte Partition eine fallend rechtsschiefe Partition.

Es bleibt noch die Eindeutigkeit zu zeigen. Nehmen wir an, es gäbe zwei verschiedene fallend rechtsschiefe Partitionen. Betrachten wir, wie in der folgenden Abbildung 1.32 skizziert, die jeweils linken Teilfolgen in ihren Partition, in denen sich die beiden Partitionen unterscheiden.

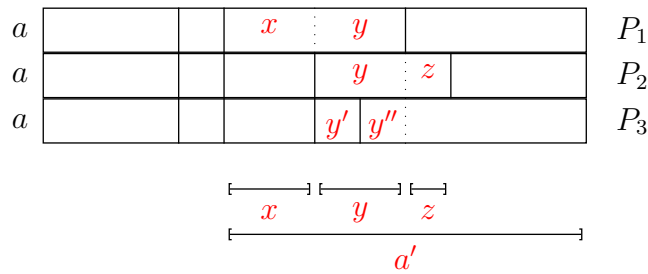


Abbildung 1.32: Skizze: Beweis der Eindeutigkeit der rechtsschiefen Partition

Betrachten wir zuerst den Fall P_1 gegen P_2 . Sei zuerst $z = \varepsilon$. Dann gilt $\mu(x) > \mu(y)$, da P_2 eine fallend rechtsschiefe Partition ist. Da xy in P_1 rechtsschief ist, gilt $\mu(x) \leq \mu(y)$ und wir erhalten den gewünschten Widerspruch.

Sei nun $z \neq \varepsilon$. Da yz nach der Partition P_2 rechtsschief ist, gilt $\mu(y) \leq \mu(z)$. Mit Korollar 1.30 folgt, dass $\mu(y) \leq \mu(yz) \leq \mu(z)$.

Nach Wahl der fallend rechtsschiefen Partition P_2 gilt $\mu(x) > \mu(yz)$. Wie wir eben gezeigt haben, gilt auch $\mu(yz) \geq \mu(y)$. Damit ist $\mu(x) > \mu(y)$ und somit ist xy nicht rechtsschief. Dies ist ein Widerspruch zur Annahme, dass P_1 eine fallend rechtsschiefe Partition ist.

Es bleibt noch der Fall P_1 gegen P_3 zu betrachten. Nach P_1 gilt $\mu(x) \leq \mu(y)$ und nach P_3 gilt $\mu(x) > \mu(y')$. Also gilt $\mu(y) > \mu(y')$. Sei y'' so gewählt, dass $y = y'y''$

und sei weiter $y = y' \cdot Y_2 \cdots Y_k$ eine fallende rechtsschiefe Partition von y , die sich aus P_3 bzw. dem ersten Teil des Beweises ergibt. Dann gilt $\mu(y') > \mu(Y_2) > \cdots > \mu(Y_k)$. Dann muss aber $\mu(y') > \mu(y'')$ sein. Somit gilt $\mu(y) > \mu(y') > \mu(y'')$, was ein offensichtlicher Widerspruch ist. ■

1.5.3 Algorithmus zur Konstruktion rechtsschiefer Zeiger

Kommen wir nun zu einem Algorithmus, der die fallend rechtsschiefe Partition zu einer gegebenen Folge konstruiert. Zuerst benötigen wir noch einige Notationen.

Notation 1.34 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Jedes Suffix $a^{(i)} = (a_i, \dots, a_n)$ von a besitzt eine eindeutige fallend rechtsschiefe Partition: $a^{(i)} = A_1^{(i)} \cdots A_{k_i}^{(i)}$.

Damit können wir die so genannten rechtsschiefen Zeiger definieren.

Definition 1.35 Sei $A_1^{(i)} = (a_i, \dots, a_{p(i)})$, dann heißt $p(i)$ der rechtsschiefe Zeiger von i .

Für den folgenden Algorithmus benötigen wir noch die folgenden vereinfachenden Notationen.

Notation 1.36 Im Folgenden verwenden wir der Einfachheit halber die beiden folgenden Abkürzungen:

$$\begin{aligned} s(i) &:= \sigma(i, p(i)), \\ \ell(i) &:= \ell(i, p(i)) = p(i) - i + 1. \end{aligned}$$

Für die Konstruktion der rechtsschiefen Zeiger arbeiten wir uns wieder vom Ende zum Anfang durch die gegebene Folge durch. Für ein neu betrachtetes Element setzen wir zunächst die rechtsschiefe Folge auf diese einelementige Folge. Ist nun der

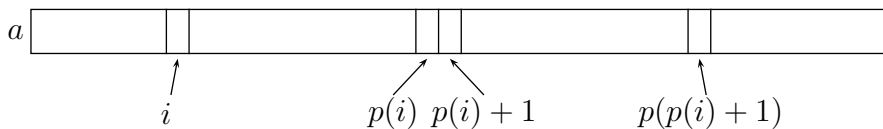


Abbildung 1.33: Skizze: Verschmelzen von Segmenten

Mittelwert des aktuell betrachteten Segments kleiner gleich dem Mittelwert des folgenden Segments, so verschmelzen wir diese beiden und betrachten dieses neue Segment als das aktuelle. Andernfalls haben wir die Eigenschaft einer fallend rechtsschiefen Partition sichergestellt. Dies ist in der Abbildung 1.33 schematisch dargestellt.

Da diese eindeutig ist, erhalten wir zum Schluss die gewünschte fallend rechtsschiefe Partition samt aller rechtsschiefer Zeiger. Somit können wir den in Abbildung 1.34 angegebenen Algorithmus zur Konstruktion rechtsschiefer Zeiger formalisieren.

```

COMPUTE_RIGHTSKEW_POINTER (real a[], int n)
{
    for (i = n; i > 0; i--)
    {
        p(i) = i;
        s(i) = a_i;
        l(i) = 1;
        while ((s(i)/l(i) ≤ s(p(i) + 1)/l(p(i) + 1)) && (p(i) < n))
        {
            s(i) = s(i) + s(p(i) + 1);
            l(i) = l(i) + l(p(i) + 1);
            p(i) = p(p(i) + 1);
        }
    }
}

```

Abbildung 1.34: Algorithmus: Rechtsschiefe Zeiger

Lemma 1.37 *Die rechtsschiefen Zeiger lassen sich in Zeit $O(n)$ berechnen.*

Beweis: Analog zum Beweis der Laufzeit des Algorithmus zur Konstruktion linksnegativer Zeiger. ■

1.5.4 Elementare Eigenschaften von MASS

Bevor wir zur Bestimmung von Teilfolgen vorgegebener Mindestlänge mit einem maximalem Mittelwert kommen, werden wir erst noch ein paar fundamentale Eigenschaften festhalten.

Lemma 1.38 *Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ und $(a_i, \dots, a_j)$ eine kürzeste Teilfolge von a der Länge mindestens λ , die deren Average Score $\mu(i, j)$ maximiert. Dann gilt $\ell(i, j) = j - i + 1 \leq 2\lambda - 1$.*

Beweis: Der Beweis sei dem Leser zur Übung überlassen. ■

Lemma 1.39 Seien $a \in \mathbb{R}^n$, $b \in \mathbb{R}^m$ und $c \in \mathbb{R}^k$ drei reelle Folgen, die die Beziehung $\mu(a) < \mu(b) < \mu(c)$ erfüllen. Dann gilt $\mu(ab) < \mu(abc)$.

Zunächst geben wir ein Gegenbeispiel an, um zu zeigen, dass im vorherigen Lemma die Bedingung $\mu(a) < \mu(b)$ notwendig ist.

Gegenbeispiel:

$a = 11$	
$b = 1$	$\mu(ab) = \frac{12}{2} = 6$
$c = 3$	$\mu(abc) = \frac{15}{3} = 5$

Beweis: Nach Korollar 1.31 gilt $\mu(a) < \mu(ab) < \mu(b) < \mu(c)$. Somit gilt insbesondere $\mu(ab) < \mu(c)$. Nach Korollar 1.31 gilt dann $\mu(ab) < \mu(abc) < \mu(c)$ und das Lemma ist bewiesen. ■

Lemma 1.40 Seien $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ und $p \in \mathbb{R}^\ell$ und sei $A_1 \cdots A_k$ die fallend rechtsschiefe Partition von a . Sei weiter m maximal gewählt, so dass

$$\mu(p \cdot A_1 \cdots A_m) = \max\{\mu(p \cdot A_1 \cdots A_i) \mid i \in [0 : k]\}.$$

Es gilt genau dann $\mu(p \cdot A_1 \cdots A_i) > \mu(A_{i+1})$, wenn $i \geq m$ gilt.

Beweis:

$\Rightarrow$: Sei i so gewählt, dass $\mu(p \cdot A_1 \cdots A_i) > \mu(A_{i+1})$ gilt. Da $A_1 \cdots A_k$ die fallend rechtsschiefe Partition von a ist, gilt $\mu(A_1) > \mu(A_2) > \cdots > \mu(A_k)$. Dann gilt auch $\mu(p \cdot A_1 \cdots A_i) > \mu(A_{i+1}) > \mu(A_{i+2}) > \cdots > \mu(A_k)$. Nach Korollar 1.31 gilt dann $\mu(p \cdot A_1 \cdots A_i) > \mu(p \cdot A_1 \cdots A_{i+1}) > \mu(A_{i+1}) > \cdots > \mu(A_k)$. Mit Korollar 1.31 gilt also auch $\mu(p \cdot A_1 \cdots A_i) > \mu(p \cdot A_1 \cdots A_j) > \mu(A_j) > \cdots > \mu(A_k)$ für alle $j \in [i : k]$. Aufgrund der Definition von m muss also $i \geq m$ gelten.

$\Leftarrow$: Zum Beweis der Behauptung unterscheiden wir die folgenden beiden Fälle, je nachdem, ob $\mu(p \cdot A_1 \cdots A_m) > \mu(A_{m+1})$ oder ob $\mu(p \cdot A_1 \cdots A_m) \leq \mu(A_{m+1})$ gilt.

Fall 1: $\mu(p \cdot A_1 \cdots A_m) > \mu(A_{m+1}) > \mu(A_{m+2}) > \cdots > \mu(A_k)$.

Korollar 1.31 liefert $\mu(p \cdot A_1 \cdots A_{m+1}) > \mu(A_{m+1}) > \mu(A_{m+2}) > \cdots > \mu(A_k)$. Nach wiederholter Anwendung von Korollar 1.31 gilt dann für $j \in [1 : k - m]$ natürlich auch $\mu(p \cdot A_1 \cdots A_{m+j}) > \mu(A_{m+j}) > \mu(A_{m+j+1}) > \cdots > \mu(A_k)$. Also gilt für $i \geq m$ $\mu(p \cdot A_1 \cdots A_i) > \mu(A_i) > \mu(A_{i+1}) > \cdots > \mu(A_k)$ und somit die Behauptung.

Fall 2: $\mu(p \cdot A_1 \cdots A_m) \leq \mu(A_{m+1})$.

Nach Korollar 1.30 gilt $\mu(p \cdot A_1 \cdots A_m) \leq \mu(p \cdot A_1 \cdots A_{m+1}) \leq \mu(A_{m+1})$. Aufgrund der Wahl von m (das Maximum von $\mu(p \cdot A_1 \cdots A_i)$ wird für $i = m$ angenommen) gilt dann $\mu(p \cdot A_1 \cdots A_m) = \mu(p \cdot A_1 \cdots A_{m+1})$. Daraus ergibt sich ein Widerspruch zur Maximalität von m .

Damit ist die Behauptung des Lemmas gezeigt. ■

1.5.5 Ein Algorithmus für MASS

Somit können wir einen Algorithmus zur Lösung unseres Problems angeben, der im Wesentlichen auf Lemma 1.40 basiert. Wir laufen von links nach rechts durch die Folge, und versuchen, ab Position i die optimale Sequenz mit maximalem Mittelwert zu finden. Nach Voraussetzung hat diese mindestens die Länge λ und nach Lemma 1.38 auch höchstens die Länge $2\lambda - 1$. Dies ist in Abbildung 1.35 illustriert.

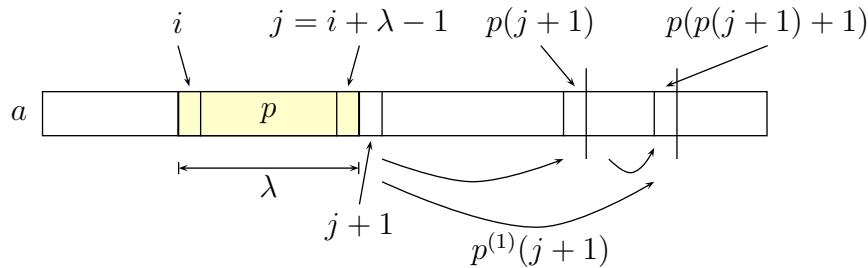


Abbildung 1.35: Skizze: Auffinden der optimalen Teilfolge

Ein einfacher Algorithmus würde jetzt alle möglichen Längen aus dem Intervall $[\lambda : 2\lambda - 1]$ ausprobieren. Der folgende, in Abbildung 1.36 angegebene Algorithmus würde dies tun, wenn die Prozedur `locate` geeignet implementiert wäre, wobei `locate` einfach die Länge einer optimalen Teilfolge zurückgibt. Würden wir `locate` mittels einer linearen Suche implementieren, so würden wir eine Laufzeit von $O(n\lambda) = O(n^2)$ erhalten.

Mit Hilfe von Lemma 1.40 können wir aber auch geschickter vorgehen. Das Lemma besagt nämlich, dass wenn wir rechts vom optimalen Schnitt sind, der Mittelwert größer als der des folgenden Segments ist. Links vom Optimum gilt, dass der Wert kleiner als der Mittelwert vom folgenden Segment ist. Somit könnten wir hier eine binäre Suche ausnutzen.

Allerdings kennen wir die Grenzen der Segmente nicht explizit, sondern nur als lineare Liste über die rechtsschiefen Zeiger. Daher konstruieren wir uns zu den

MASS (real $a[]$, int n , λ)

```

{
  int  $mi = 1$ ;    $mj = \lambda$ ;    $mm = \mu(mi, mj)$ ;
  for ( $i = 1$ ;  $i \leq n$ ;  $i++$ )
  {
     $j = i + \lambda - 1$ ;
    if ( $\mu(i, j) \leq \mu(j + 1, p(j + 1))$ )
       $j = \text{locate}(i, j)$ ;
    if ( $\mu(i, j) > mm$ )
    {
       $mi = i$ ;    $mj = j$ ;    $mm = \mu(mi, mj)$ ;
    }
  }
}

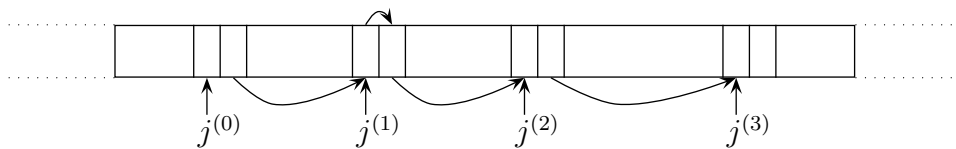
```

Abbildung 1.36: Algorithmus: MASS

rechtsschiefen Zeiger noch solche, die nicht nur auf das nächste Segment, sondern auch auf das 2^k -te folgende Segment angibt. Dafür definieren wir erst einmal formal die Anfänge der nächsten 2^k -ten Segment sowie die darauf basierenden iterierten rechtsschiefen Zeiger wie folgt:

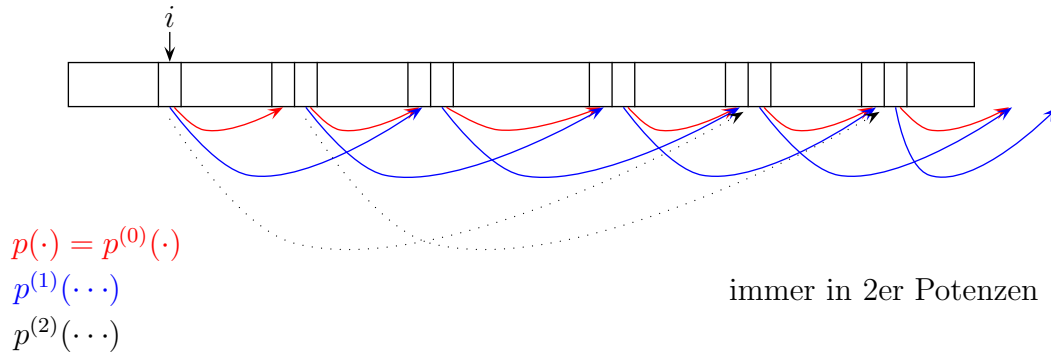
$$\begin{aligned}
 j^{(0)} &:= j, \\
 j^{(k)} &= \min\{p(j^{(k-1)} + 1), n\}, \\
 p^{(0)}(i) &= p(i), \\
 p^{(k)}(i) &= \min\{p^{(k-1)}(p^{(k-1)}(i) + 1), n\}.
 \end{aligned}$$

Hierbei gibt $p^{(k)}(i)$ das Ende nach 2^k Segmenten an. Dies ist in folgenden Abbildungen illustriert, in Abbildung 1.37 die Definition von $j^{(\cdot)}$ und in Abbildung 1.38 die Definition von $p^{(\cdot)}$.



$$j^{(2)} = p(j^{(1)} + 1)$$

Abbildung 1.37: Skizze: Definition der Werte $j^{(\cdot)}$

Abbildung 1.38: Skizze: Definition der Werte $p^{(k)}$

Definition 1.41 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$ eine reelle Folge und $p(\cdot)$ die zugehörigen rechtsschiefen Zeiger. Der k -te iterierte rechtsschiefe Zeiger $p^{(k)}(i)$ ist rekursiv definiert durch $p^{(k)}(i) := \min\{p^{(k-1)}(p^{(k-1)}(i) + 1), n\}$ und $p^{(0)}(i) := p(i)$.

Diese iterierten rechtsschiefen Zeiger lassen sich aus den rechtsschiefen Zeigern in Zeit $O(\log(\lambda))$ berechnen, da wir ja maximal die 2λ -iterierten rechtsschiefen Zeiger aufgrund der oberen Längenbeschränkung von $2\lambda - 1$ benötigen.

Mit Hilfe dieser iterierten rechtsschiefen Zeiger können wir jetzt die binäre Suche in der Prozedur `locate` wie im folgenden Algorithmus implementieren, der in Abbildung 1.39 angegeben ist. Hierbei ist zu beachten, dass die erste zu testende Position bei der binären Suche bei maximal $\log(2\lambda - 1) - 1 \leq \log(\lambda)$ liegt.

```

LOCATE (int  $i, j$ )
{
  for ( $k = \log(\lambda); k \geq 0; k--$ )
  {
    if ( $(j \geq n) \parallel (\mu(i, j) > \mu(j + 1, p(j + 1)))$ )
      return  $j$ ;
    if ( $(p^{(k)}(j + 1) < n) \&\&$ 
        ( $\mu(i, p^{(k)}(j + 1)) \leq \mu(p^{(k)}(j + 1) + 1, p(p^{(k)}(j + 1) + 1))$ )
       $j = p^{(k)}(j + 1)$ ;
  }
  if ( $(j < n) \&\& (\mu(i, j) \leq \mu(j + 1, p(j + 1)))$ )
     $j = p(j + 1)$ ;
  return  $j$ ;
}

```

Abbildung 1.39: Algorithmus: `locate( $i, j$ )`

Die Strategie der Suche ist noch einmal in der folgenden Abbildung 1.40 illustriert.

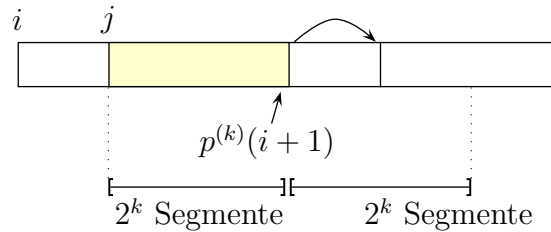


Abbildung 1.40: Skizze: Strategie von `locate`

Zusammenfassend erhalten wir das folgende Theorem.

Theorem 1.42 Sei $a = (a_1, \dots, a_n) \in \mathbb{R}^n$. Eine kürzeste Teilfolge der Länge mindestens λ mit maximalem durchschnittlichem Score kann in Zeit $O(n \log(\lambda))$ gefunden werden.

Wir wollen hier noch erwähnen, dass mittlerweile für das MASS-Problem Algorithmen mit linearer Laufzeit bekannt sind.

2.1 Definition von Suffix-Tries und Suffix-Trees

In diesem Kapitel betrachten wir einige Algorithmen zur Konstruktion von Suffix-Tries bzw. Suffix-Trees. Letztere werden zur effizienten Suche von Wörtern und Teilstrings (z.B. Tandem-Repeats) im nächsten Kapitel benötigt.

2.1.1 Σ -Bäume und (kompakte) Σ^+ -Bäume

Zunächst definieren wir die so genannten Σ -Bäume.

Definition 2.1 Sei Σ ein Alphabet. Ein Σ -Baum ist ein gewurzelter Baum mit Kantenmarkierungen aus Σ , so dass kein Knoten zwei ausgehende Kanten mit derselben Markierung besitzt. Ein Σ -Baum wird oft auch als Trie bezeichnet.

In Abbildung 2.1 ist ein Beispiel eines Σ -Baumes mit $\Sigma = \{a, b\}$ angegeben.

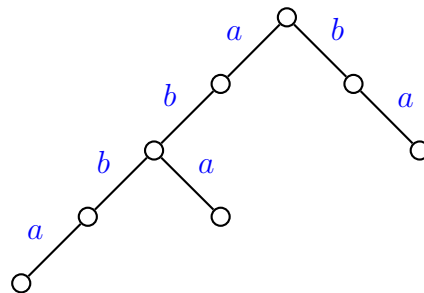


Abbildung 2.1: Beispiel: Ein Σ -Baum T_1 mit $\Sigma = \{a, b\}$

Erlaubt man als Kantenmarkierung Zeichenreihen über Σ , so erhält man die so genannten Σ^+ -Bäume.

Definition 2.2 Sei Σ ein Alphabet. Ein Σ^+ -Baum ist ein gewurzelter Baum mit Kantenmarkierungen aus Σ^+ , so dass kein Knoten zwei ausgehende Kanten besitzt, deren Markierungen mit demselben Zeichen aus Σ beginnen.

In Abbildung 2.2 ist ein Beispiel eines Σ^+ -Baumes mit $\Sigma = \{a, b\}$ angegeben.

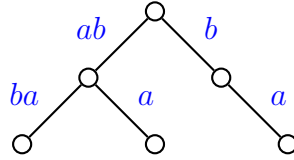


Abbildung 2.2: Beispiel: Ein Σ^+ -Baum T_2 mit $\Sigma = \{a, b\}$

Werden in Σ^+ -Bäumen Knoten mit nur einem Kind verboten, so erhalten wir kompakte Σ^+ -Bäume.

Definition 2.3 Ein Σ^+ -Baum heißt kompakt, wenn es keinen Knoten außer der Wurzel mit nur einem Kind gibt. Ein solcher kompakter Σ^+ -Baum wird auch als kompaktifizierter Trie bezeichnet.

In Abbildung 2.3 ist ein Beispiel eines Σ -Baumes mit $\Sigma = \{a, b\}$ angegeben.

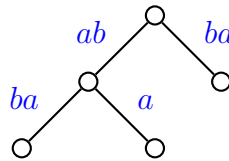


Abbildung 2.3: Beispiel: Ein kompakter Σ^+ -Baum T_3 mit $\Sigma = \{a, b\}$

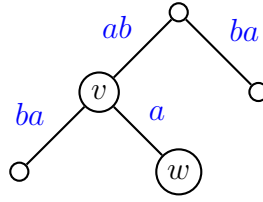
2.1.2 Grundlegende Notationen und elementare Eigenschaften

In diesem Abschnitt wollen wir noch einige grundlegende Notationen und Eigenschaften von Σ^+ -Bäumen einführen.

Sei im Folgenden T ein Σ^+ -Baum.

- Für einen Knoten $v \in V(T)$ bezeichnen wir mit $\text{path}(v)$ die Konkatenation der Kantenmarkierungen auf dem Pfad von der Wurzel zu v , siehe auch Abbildung 2.4, und $|\text{path}(v)|$ wird als *Worttiefe* von v bezeichnet.
- Wir sagen, dass eine Zeichenreihe $w \in \Sigma^*$ von T *dargestellt* wird, wenn es einen Knoten $v \in V(T)$ und ein $x \in \Sigma^*$ gibt, so dass $\text{path}(v) = wx$.

In Abbildung 2.4 stellt der Σ^+ -Baum unter anderen das Wort abb dar.

Abbildung 2.4: Beispiel: $\text{path}(v) = ab$ und $\text{path}(w) = aba$ in T_3

- Gilt $\text{path}(v) = w$ für einen Knoten $v \in V(T)$ und eine Zeichenreihe $w \in \Sigma^*$, dann schreiben wir $\overline{w}$ für v .

In Abbildung 2.4 gilt im angegebenen Σ^+ -Baum beispielsweise $\overline{ab} = v$ und $\overline{\varepsilon}$ bezeichnet die Wurzel dieses Baumes.

- Mit $\text{words}(T)$ bezeichnen wir die Menge der Wörter, die im Σ^+ -Baum T dargestellt werden. Formal lässt sich das wie folgt definieren:

$$\text{words}(T) = \{w : \exists v \in V(T), x \in \Sigma^* : \text{path}(v) = wx\}.$$

In Abbildung 2.4 gilt für den Σ^+ -Baum T beispielsweise

$$\text{words}(T) = \{\varepsilon, a, ab, aba, abb, abba, b, ba\},$$

wobei ε wie üblich das leere Wort bezeichnet.

Für die Bäume aus den Abbildungen 2.1, 2.2 und 2.3 gilt

$$\text{words}(T_1) = \text{words}(T_2) = \text{words}(T_3).$$

2.1.3 Suffix-Tries und Suffix-Trees

Bevor wir zu Suffix-Tries und Suffix-Trees kommen, wiederholen wir erst noch die Definitionen und zugehörigen Notationen zu Präfixen, Suffixen und Teilwörtern.

Notation 2.4 Sei Σ ein Alphabet und $w \in \Sigma^*$. Sei weiter v ein Teilwort von w (d.h. es gibt $x, y \in \Sigma^*$ mit $w = xvy$), dann schreiben wir auch $v \sqsubseteq w$.

Beobachtung 2.5 Seien $x, w \in \Sigma^*$ und $\$, \# \notin \Sigma$, dann ist x genau dann ein Präfix (bzw. Suffix) von w , wenn $\$x \sqsubseteq \w (bzw. $x\$ \sqsubseteq w\$$) gilt.

Damit kommen wir zur Definition eines Suffix-Tries.

Definition 2.6 Sei Σ ein Alphabet. Ein Suffix-Trie für ein Wort $t \in \Sigma^*$ ist ein Σ -Baum T mit der Eigenschaft $\text{words}(T) = \{w \in \Sigma^* : w \sqsubseteq t\}$.

In Abbildung 2.5 ist der Suffix-Trie für das Wort *abbab* angegeben.

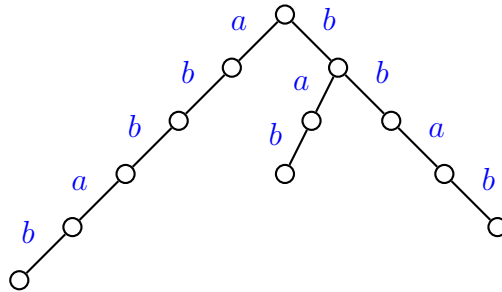


Abbildung 2.5: Beispiel für einen Suffix-Trie mit $t = abbab$

Kommen wir nun zur Definition eines Suffix-Baumes.

Definition 2.7 Sei Σ ein Alphabet und $t \in \Sigma^*$. Ein Suffix-Baum (engl. Suffix-Tree) für t ist ein kompakter Σ^+ -Baum $T = T(t)$ mit $\text{words}(T) = \{w \in \Sigma^* : w \sqsubseteq t\}$.

In Abbildung 2.6 ist links ein Suffix-Baum für das Wort *abbab* angegeben.

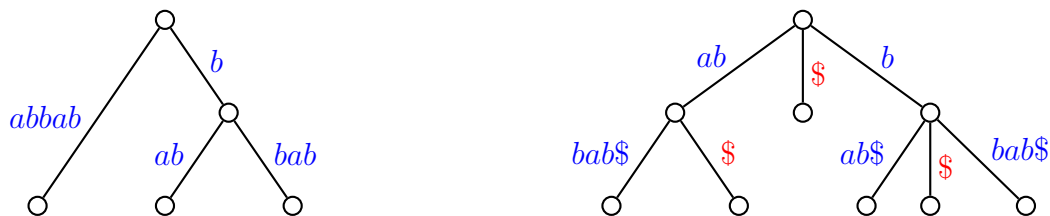


Abbildung 2.6: Beispiel: Suffix-Bäume für $t = abbab$ und $t = abbab\$$

Warum diese Bäume Suffix-Tries bzw. Suffix-Trees statt Infix-Tries bzw. Infix-Trees heißen, wird gleich in den folgenden Bemerkungen klar.

Bemerkungen: Halten wir zunächst ein paar nützliche Eigenschaften von Suffix-Tries bzw. -Trees fest.

- Oft betrachtet man statt des Suffix-Tries bzw. Suffix-Trees für $t \in \Sigma^*$ den Suffix-Trie bzw. Suffix-Tree für $t\$ \in (\Sigma \cup \{\$\})^*$, wobei $\$ \notin \Sigma$ gilt. Der Vorteil dieser Betrachtungsweise ist, dass dann jedes Suffix von t zu einem Blatt des Suffix-Tries bzw. Suffix-Trees für t korrespondiert. Damit wird dann auch die Bezeichnung Suffix-Trie bzw. Suffix-Tree klar.

- Als Kantenmarkierungen werden keine Zeichenreihen verwendet, sondern so genannte *Referenzen* (Start- und Endposition) auf das Originalwort. Für ein Beispiel siehe auch Abbildung 2.7. Statt ab wird dort die Referenz $(1, 2)$ für t_1t_2 im Wort $t = abbab\$$ angegeben. Man beachte, dass ab mehrfach in t vorkommt, nämlich ab Position 1 und ab Position 4. Die Wahl der Referenz wird dabei nicht willkürlich sein. Wir betrachten dazu die Kante mit Kantenmarkierung w , deren Referenz bestimmt werden soll. Im Unterbaum des Suffix-Baumes, der am unteren Ende der betrachteten Kante gewurzelt ist, wird das längste im Unterbaum dargestellte Suffix s betrachtet. Dann muss auch $w \cdot s$ ein Suffix von t sein und als Referenz für die Kantenmarkierung w wird $(|t| - |w \cdot s| + 1, |t| - |s|)$ gewählt.

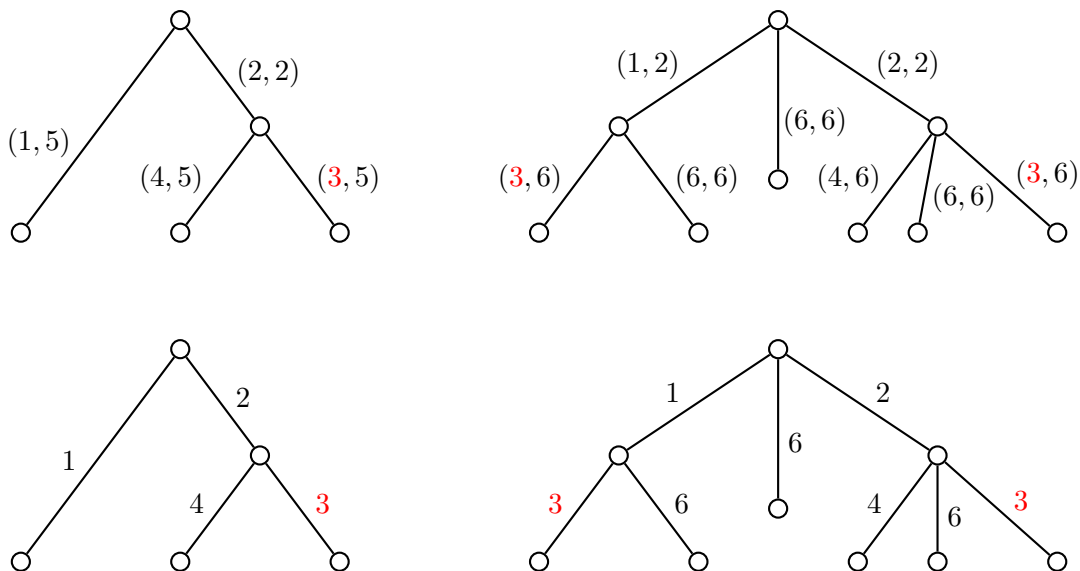


Abbildung 2.7: Beispiel: Suffix-Bäume für $t = abbab$ und $t = abbab\$$ mit normalen (oben) bzw. reduzierten (unten) Referenzen als Kantenmarkierungen

- Der Platzbedarf für einen Suffix-Baum beträgt mit dieser Referenzdarstellung $O(|t|)$. Würde man hingegen die Zeichenreihen an die Kanten schreiben, so könnte der Platzbedarf auch auf $\Theta(|t|^2)$ ansteigen. Der Beweis ist dem Leser als Übungsaufgabe überlassen.
- In der Literatur wird der Suffix-Trie manchmal auch als *atomic suffix tree* und der Suffix-Tree selbst als *compact suffix tree* bezeichnet. Dies führt oft zu Verwirrung, insbesondere, da heutzutage oft unter einem kompakten Suffix-Baum (compact suffix tree) etwas ganz anderes verstanden wird.

2.2 Repräsentationen von Bäumen

Zuerst überlegen wir uns einige Möglichkeiten, wie man Bäume im Allgemeinen und Suffix-Bäume im Speziellen überhaupt speichern kann. In einem Baum muss man im Wesentlichen die Kinder eines Knoten speichern. Dafür gibt es verschiedene Möglichkeiten, die jeweils ihre Vor- und Nachteile haben.

Allgemein erinnern wir uns zuerst daran, dass ein Baum mit m Blättern, der keinen Knoten mit nur einem Kind besitzt, maximal $m-1$ innere Knoten haben kann. Somit besitzt ein Suffix-Baum für t weniger als $|t|$ innere Knoten und somit insgesamt weniger als $2|t|$ Knoten.

Für eine detailliertere Analyse in der Bioinformatik nehmen wir an, dass sowohl Integers als auch Verweise für Ihre Speicherung 4 Bytes benötigen. Für die Kantenmarkierungen werden wir in der Regel nicht die volle Referenz (i, j) speichern, sondern oft nur die Startposition i . Da in einem Suffixbaum alle Suffixe gespeichert sind, muss von dem Knoten, den wir über die Kante (i, j) erreichen, eine Kante mit der Startposition $j + 1$ beginnen. Wenn wir uns diese merken (meist als erste ausgehende Kante in unseren Darstellungen), können wir daraus die Endposition j rekonstruieren. Hierfür ist es wichtig, dass wir für ein Teilwort w von t die Referenz geeignet wählen (d.h. wie im vorherigen Abschnitt beschrieben).

Im Folgenden ist bei der Zeitanalyse der verschiedenen Darstellungen von Bäumen mit Zeit die Zugriffszeit auf ein Kind gemeint, zu dem die Kantenmarkierung mit einem bestimmten Zeichen beginnt.

2.2.1 Darstellung der Kinder mit Feldern

Die Kinder eines Knotens lassen sich sehr einfach mit Hilfe eines Feldes der Größe $|\Sigma|$ darstellen.

- Platz: $O(|t| \cdot |\Sigma|)$.

Dies folgt daraus, dass für jeden Knoten ein Feld mit Platzbedarf $O(|\Sigma|)$ benötigt wird.

- Zeit: $O(1)$.

Die üblichen Realisierungen von Feldern erlauben einen Zugriff in konstanter Zeit auf die einzelnen Feldelemente.

Der Zugriff ist also sehr schnell, wo hingegen der Platzbedarf, insbesondere bei großen Alphabeten, doch sehr groß werden kann.

Für eine DNA-Sequenz sind also pro Knoten maximal 20 Bytes nötig (16 Bytes für die Verweise und 4 Bytes für die Referenz auf die Kantenmarkierung). Da es etwa doppelt so viele Knoten wie Nukleotide gibt, sind als pro Nukleotid maximal 40 Bytes nötig.

Für ein Protein sind also pro Knoten maximal 84 Bytes nötig (80 Bytes für die Verweise und 4 Bytes für die Referenz auf die Kantenmarkierung). Da es etwa doppelt so viele Knoten wie Aminosäuren in einem Protein gibt, sind als pro Aminosäure maximal 168 Bytes nötig.

Für die zweite Referenz müssen wir uns noch explizit das Kind merken, deren Referenz im Wesentlichen das Ende der gesuchten Referenz angibt. Dazu sind im Falle von DNA bzw. Proteinen noch jeweils 2 Bits bzw. 5 Bits je inneren Knoten nötig (wenn man die Kinder in konstanter Zeit finden will).

Insgesamt sind daher pro Nukleotid bzw. pro Aminosäure etwa 42 Bytes bzw. etwa 173 Bytes nötig.

2.2.2 Darstellung der Kinder mit Listen

Eine andere Möglichkeit ist, die Kinder eines Knotens in einer linearen Liste zu verwalten, wobei das erste Kind einer Liste immer dasjenige ist, das für den zweiten Referenzwert der inzidenten Kante zu seinem Elter nötig ist.

- Platz: $O(|t|)$.

Für jeden Knoten ist der Platzbedarf proportional zur Anzahl seiner Kinder. Damit ist Platzbedarf insgesamt proportional zur Anzahl der Knoten des Suffix-Baumes, da jeder Knoten (mit Ausnahme der Wurzel) das Kind eines Knotens ist. Im Suffix-Baum gilt, dass jeder Knoten entweder kein oder mindestens zwei Kinder hat. Für solche Bäume ist bekannt, dass die Anzahl der inneren Knoten kleiner ist als die Anzahl der Blätter. Da ein Suffix-Baum für einen Text der Länge m maximal m Blätter besitzt, folgt daraus die Behauptung für den Platzbedarf.

- Zeit: $O(|\Sigma|)$.

Leider ist hier die Zugriffszeit auf ein Kind sehr groß, da im schlimmsten Fall (aber auch im Mittel) die gesamte Kinderliste eines Knotens durchlaufen werden muss und diese bis zu $|\Sigma|$ Elemente umfassen kann.

In vielen Anwendungen wird jedoch kein direkter Zugriff auf ein bestimmtes Kind benötigt, sondern eine effiziente Traversierung aller Kinder. Dies ist mit der Geschwisterliste aber auch wieder in konstanter Zeit pro Kind möglich, wenn keine besondere Ordnung auf den Kindern berücksichtigt werden soll.

Pro Knoten des Baumes sind also maximal 12 Bytes nötig (8 Bytes für die Verweise und 4 Bytes für die Kantenmarkierung). Für die gespeicherte Sequenz sind also 24 Bytes pro Element nötig. Bei geschickterer Darstellung der Blätter (hier ist ja kein Verweis auf sein ältestes Kind nötig) kommen wir mit 20 Bytes aus.

Wenn man eine Ordnung in den Listen aufrecht erhalten will, so muss man auf die reduzierten Referenzen verzichten und pro Nukleotid nochmals 4–8 Bytes opfern (je nachdem, wie geschickt man die Referenzen an den zu Blättern inzidenten Kanten implementiert).

2.2.3 Darstellung der Kinder mit balancierten Bäumen

Die Kinder lassen sich auch mit Hilfe von balancierten Suchbäumen (AVL-, Rot-Schwarz-, B-Bäume, etc.) verwalten:

- Platz: $O(|t|)$

Da der Platzbedarf für einen Knoten ebenso wie bei linearen Listen proportional zur Anzahl der Kinder ist, folgt die Behauptung für den Platzbedarf unmittelbar.

- Zeit: $O(\log(|\Sigma|))$.

Da die Tiefe von balancierten Suchbäumen logarithmisch in der Anzahl der abzuspeichernden Schlüssel ist, folgt die Behauptung unmittelbar.

Auch hier ist eine Traversierung der Kinder eines Knotens in konstanter Zeit pro Kind möglich, indem man einfach den zugehörigen balancierten Baum traversiert.

Auch hier sind für eine Sequenz also pro Element maximal etwa 24–26 Bytes nötig. Dies folgt daraus, dass balancierte Bäume etwa genauso viele Links besitzen wie Listen (wenn man annimmt, dass die Daten auch in den inneren Knoten gespeichert werden und man auf die unnötigen Referenzen in den Blättern verzichtet, sonst erhöht sich der Platzbedarf jeweils um etwa einen Faktor 2). Es werden jedoch noch zusätzlich 1–2 Bytes für die Verwaltung der Balancierung benötigt.

2.2.4 Darstellung des Baumes mit einer Hash-Tabelle

Eine weitere Möglichkeit ist die Verwaltung der Kinder aller Knoten in einem einzigen großen Feld der Größe $O(|t|)$. Um nun für ein Knoten auf ein spezielles Kind zuzugreifen wird dann eine *Hashfunktion* verwendet:

$$h : V \times \Sigma \rightarrow \mathbb{N} : (v, a) \mapsto h(v, a).$$

Hierbei interpretieren wir die Knotenmenge V in geeigneter Weise als natürliche Zahlen, d.h. $V \subset \mathbb{N}$. Zu jedem Knoten und dem Symbol, die das Kind identifizieren, wird ein Index des globalen Feldes bestimmt, an der die gewünschte Information enthalten ist.

Leider bilden Hashfunktionen ein relativ großes Universum von potentiellen Referenzen (hier Paare von Knoten und Symbolen aus Σ , also mit einer Mächtigkeit von $\Theta(|t| \cdot |\Sigma|)$) auf ein kleines Intervall ab (hier Indizes aus $[1 : \ell]$ mit $\ell = \Theta(|t|)$). Daher sind so genannte *Kollisionen* prinzipiell nicht auszuschließen. Ein Beispiel ist das so genannte *Geburtstagsparadoxon*. Ordnet man jeder Person in einem Raum eine Zahl aus dem Intervall $[1 : 366]$ zu (nämlich ihren Geburtstag), dann ist bereits ab 23 Personen die Wahrscheinlichkeit größer als 50%, dass zwei Personen denselben Wert erhalten. Also muss man beim Hashing mit diesen Kollisionen leben und diese geeignet auflösen. Für die Details der Kollisionsauflösung verweisen wir auf andere Vorlesungen.

Um solche Kollisionen überhaupt festzustellen, enthält jeder Feldeintrag i neben den normalen Informationen noch die Informationen, wessen Kind er ist und über welches Symbol er von seinem Elter erreichbar ist. Somit lassen sich Kollisionen leicht feststellen und die üblichen Operationen zur Kollisionsauflösung anwenden.

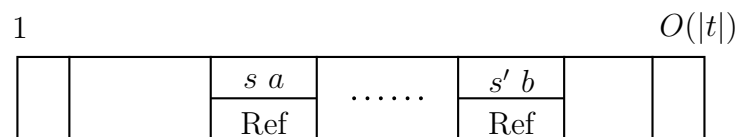


Abbildung 2.8: Skizze: Realisierung mittels eines Feldes und Hashing

- Platz: $O(|t|)$

Das folgt unmittelbar aus der obigen Diskussion.

- Zeit: $O(1)$

Im Wesentlichen erfolgt der Zugriff in konstanter Zeit, wenn man voraussetzt, dass sich die Hashfunktion einfach (d.h. in konstanter Zeit) berechnen lässt und dass sich Kollisionen effizient auflösen lassen.

Hier sind für jeden Eintrag 4 Bytes für die Kantenmarkierung nötig. Für die Verwaltung (ein Verweis auf den Elter und das Zeichen, das das Kind identifiziert) sind 5 Bytes nötig. Um die Endposition zu merken, muss man sich noch das entsprechende Kind merken, was ein weiteres Byte kostet. Da es maximal $2|t|$ Kanten gibt, sind somit maximal 20 Bytes pro Element nötig. Da beim Hashing für eine effiziente

Kollisionsauflösung das Feld nur bis zu etwa 80% gefüllt sein darf, muss das Feld entsprechend größer gewählt werden, so dass der Speicherbedarf insgesamt maximal 25 Bytes pro Element der Sequenz beträgt.

2.2.5 Speicherplatzeffiziente Feld-Darstellung

Man kann die Methode der linearen Listen auch selbst mit Hilfe eines Feldes implementieren. Dabei werden die Geschwisterlisten als konsequente Feldelemente abgespeichert. Begonnen wird dabei mit der Geschwisterliste des ältesten Kindes der Wurzel, also mit den Kindern der Wurzel, die dann ab Position 1 des Feldes stehen. Die Wurzel selbst wird dabei nicht explizit dargestellt.

Auch hier unterscheidet man wieder zwischen internen Knoten und Blättern. Ein interner Knoten belegt zwei Feldelemente, ein Blatt hingegen nur ein Feldelement. Ein innerer Knoten hat einen Verweis auf das Feldelement, in dem die Geschwisterliste des ältesten Kindes konsequentiv abgespeichert wird. Im zweiten Feldelement steht die Startposition der Kantenmarkierung innerhalb von t für die Kante, die in den Knoten hineinreicht. Für ein Blatt steht nur die Startposition der Kantenmarkierung innerhalb von t für die Kante, die in das Blatt hineinreicht.

Ein Beispiel ist in Abbildung 2.9 für $t = abbab\$$ angegeben. Zur Unterscheidung der beiden Knotentypen wird ein Bit spendiert. Dies ist in der Abbildung 2.9 zu können wir als hochgestelltes B zu erkennen. Des Weiteren müssen mit Hilfe eines weiteren Bits auch noch die Enden der Geschwisterlisten markiert werden. Dies ist in Abbildung 2.9 mit einem höhergestellten $*$ markiert. Verweise sind dort als rote Zahlen geschrieben.

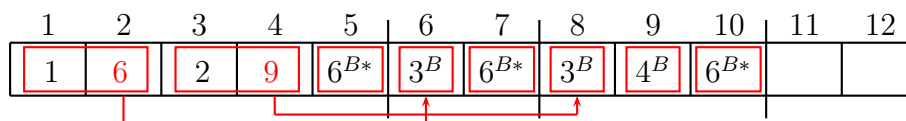


Abbildung 2.9: Beispiel: Feld-Darstellung des Suffix-Baumes aus Abbildung 2.6 für $t = abbab\$$

Es bleibt nur die Frage, wie man das Ende einer Kantenmarkierung findet. Das ist jedoch sehr einfach: Hat man den Beginn der Kantenmarkierung, so folgt man dem Verweis dieses Knotens. Im folgenden Knoten steht dann der Beginn der nächsten Kantenmarkierung; dieser muss nur noch um es erniedrigt werden. Handelt es sich um ein Blatt, so gibt es keinen Folgeverweis, aber die Endposition ist dann das Ende, also $|t|$. Hierfür ist aber wichtig, dass die Knoten der Geschwisterliste in der

richtigen Reihenfolge abgespeichert werden, nämlich das Kind, dessen zugehörige Kantenmarkierung am frühestens in t auftritt, als erstes.

- Platz: $O(|t|)$.

Wie bei Listen

- Zeit: $O(|\Sigma|)$.

Leider ist hier die Zugriffszeit auf ein Kind sehr groß, da im schlimmsten Fall (aber auch im Mittel) die gesamte Kinderliste eines Knotens durchlaufen werden muss und diese bis zu $|\Sigma|$ Elemente umfassen kann.

Auch hier ist eine Traversierung der Kinder eines Knotens in konstanter Zeit pro Kind möglich, indem man einfach im Feld weiterläuft bis man auf eine Markierung stößt, die das Ende der Geschwisterliste markiert.

Pro inneren Knoten wird nun eine Zahl, ein Verweis und zwei Bits gespeichert, pro Blatt eine Zahl und zwei Bits. Spendiert man von den 4 Bytes jeweils zwei Bits um die Bits dort zu Markierung von Blättern und dem Ende von Geschwisterlisten zu speichern, kommt man mit maximal 12 Bytes pro Sequenzelement aus.

07.11.06

2.3 WOTD-Algorithmus

Nun wollen wir einen einfachen Algorithmus zur Konstruktion von Suffix-Bäumen angeben. Der so genannte *WOTD-Algorithmus* (für write-only-top-down) fügt die Kanten in den Suffix-Baum von der Wurzel beginnend (daher top-down) ein.

2.3.1 Die Konstruktion

Zuerst werden alle möglichen Suffixe der Zeichenkette $t \in \Sigma^*$ bestimmt. Im nächsten Schritt werden diese nach dem ersten Buchstaben sortiert. Hierzu wird vorzugsweise ein Bucket-Sort verwendet. Rekursiv wird dann daraus der Suffix-Baum aufgebaut. Siehe hierzu den Algorithmus in Abbildung 2.10.

Für den eigentlichen Aufbau traversiert man den zu konstruierenden Suffix-Baum top-down. Für jeden Knoten gibt es eine Menge S von Zeichenreihen, die noch zu verarbeiten sind. Diese wird dann nach dem ersten Zeichen (mittels eines Bucket-Sorts) sortiert und liefert eine Partition der Zeichenreihe in Mengen S_c mit demselben ersten Zeichen für $c \in \Sigma$. Für jede nichtleere Menge S_c wird ein neues Kind generiert, wobei die inzidente Kante als Kantenmarkierung das längste gemeinsame Präfix p

aus S_c erhält. Dann werden von jeder Zeichenreihe aus S_c dieses Präfix p entfernt und in einer neuen Menge S'_c gesammelt. An dem neu konstruierten Kind wird nun die Prozedur rekursiv mit der Menge S'_c aufgerufen. Da Werte immer nur geschrieben und nie modifiziert werden, erklärt sich nun der Teil write-only aus WOTD.

In der Implementierung wird man die Menge S (bzw. S_c) nicht wirklich als Menge von Wörtern implementieren, da sonst die Menge S bereits quadratischen Platzbedarf hätte. Da S jeweils eine Menge von Suffixen beschreibt, wird man diese Menge als Menge der Anfangspositionen der Suffixe in S realisieren.

Ein Beispiel für den Ablauf dieses Algorithmus mit dem Wort $t = ababb$ ist in Abbildung 2.11 angegeben.

2.3.2 Zeitbedarf

Wir wollen nun den Zeitbedarf des Algorithmus analysieren. Für jeden Knoten wird ein Bucket-Sort nach dem ersten Zeichen ausgeführt. Dies geht in Zeit $O(|S|)$ wir verteilen diese Kosten auf das jeweils erste Zeichen in den Mengen S_c . Somit erhält jedes erste Zeichen in S_c konstante Kosten. Zur Ermittlung des längsten gemeinsamen Präfixes p in S_c werden in jeder Zeichenreihe in S_c maximal $|p| + 1$ Vergleiche ausgeführt. Diese verteilen wir nun auf die Zeichen des längsten gemeinsamen Präfixes. Da nach Vorsortierung $|p| \geq 1$ gilt, erhält auch hier wieder jedes Zeichen des längsten gemeinsamen Präfixes konstante Kosten. Da anschließend das längste gemeinsame Präfix entfernt wird, enthält jedes Zeichen aus $S(t)$ maximal konstant viele Einheiten. Da für ein Wort t der Länge n gilt, dass die Anzahl der Zeichen aller Suffixe von t gleich $\binom{n+1}{2}$ ist, beträgt die Laufzeit im worst-case $O(n^2)$.

Theorem 2.8 *Für ein Wort $t \in \Sigma^n$ mit $n \in \mathbb{N}$ kann der zugehörige Suffix-Baum mit dem WOTD-Algorithmus im worst-case in Zeit $O(n^2)$ konstruiert werden.*

In der Praxis bricht der Algorithmus ja ab, wenn $|S_c| = 1$ ist. Für jeden internen Knoten v eines Suffix-Baumes gilt im average-case $|\text{path}(v)| = O(\log_{|\Sigma|}(n))$. Somit werden im average-case maximal die ersten $O(\log_{|\Sigma|}(n))$ Zeichen eines Wortes aus $S(t)$ betrachtet. Daraus ergibt sich dann im average-case eine Laufzeit von $O(n \log_{|\Sigma|}(n))$.

Theorem 2.9 *Für ein Wort $t \in \Sigma^n$ mit $n \in \mathbb{N}$ kann der zugehörige Suffix-Baum mit dem WOTD-Algorithmus im average-case in Zeit $O(n \log_{|\Sigma|}(n))$ konstruiert werden.*

2.4 Der Algorithmus von Ukkonen

In diesem Abschnitt wollen wir einen Algorithmus vorstellen, der im worst-case eine lineare Laufzeit zur Konstruktion eines Suffix-Baumes besitzt.

2.4.1 Suffix-Links

Zunächst einmal benötigen wir für die Darstellung des Algorithmus den Begriff eines Suffix-Links.

Definition 2.10 Sei $t \in \Sigma^*$ und sei T der Suffix-Baum zu t . Zu jedem inneren Knoten $\overline{aw}$ von T mit $a \in \Sigma$, $w \in \Sigma^*$ (also mit Ausnahme der Wurzel) ist der Suffix-Link des Knotens $\overline{aw}$ als $\overline{w}$ definiert (in Zeichen $\text{slink}(\overline{aw}) = \overline{w}$).

In Abbildung 2.12 sind die Suffix-Links für den Suffix-Baum für $t = ababb$ als gestrichelte Linien dargestellt.

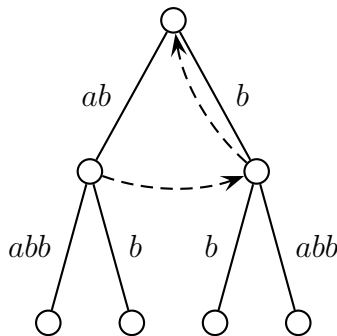


Abbildung 2.12: Beispiel: Die Suffix-Links im Suffix-Baum für $t = ababb$

Als nächstes sollten wir uns noch überlegen, dass Suffix-Links überhaupt wohldefiniert sind. Es könnte ja durchaus sein, dass es einen Knoten $\overline{aw}$ für $a \in \Sigma$ und $w \in \Sigma^*$ gibt, aber es keinen Knoten v mit $\text{path}(v) = w$ gibt.

Wir müssen also Folgendes zeigen: Wenn $\overline{aw}$ ein innerer Knoten des Suffix-Baums ist, dann ist auch $\overline{w}$ ein innerer Knoten.

Ist $\overline{aw}$ ein innerer Knoten des Suffix-Baumes, dann muss es Zeichen $b \neq c \in \Sigma$ geben, so dass sowohl $awb \sqsubseteq t$ als auch $awc \sqsubseteq t$ gilt. Dabei sind bx und cx' mit $x, x' \in \Sigma^*$ die Kantenmarkierungen von (mindestens) zwei ausgehenden Kanten aus dem Knoten $\overline{aw}$. Dies folgt aus der Definition des Suffix-Baumes als kompakter Σ^+ -Baum und ist in Abbildung 2.13 schematisch dargestellt.

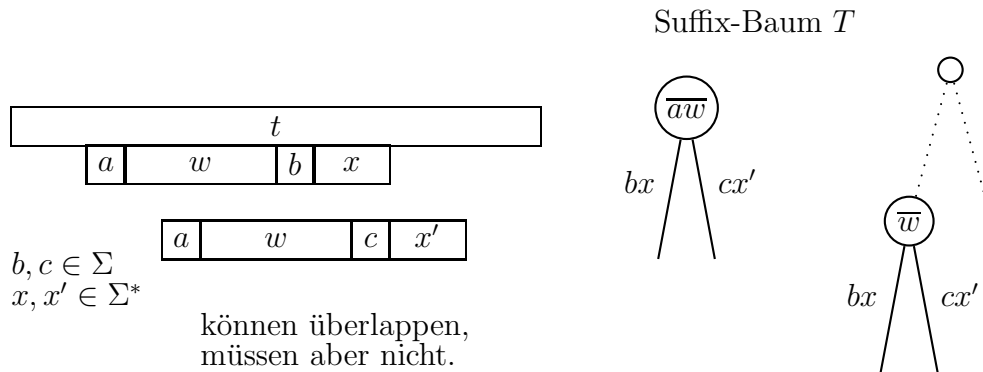


Abbildung 2.13: Schema: Suffix-Links sind wohldefiniert

Gilt jedoch sowohl $awb \sqsubseteq t$ als auch $awc \sqsubseteq t$, dann gilt auch $wb \sqsubseteq t$ und $wc \sqsubseteq t$. Somit muss $\overline{w}$ ein Knoten im Suffix-Baum für t sein.

Man beachte, dass man für Blätter eines Suffix-Baumes auf diese Weise keine Suffix-Links definieren kann. Hier kann es passieren, dass für ein Blatt $\overline{aw}$ kein Knoten $\overline{w}$ im Suffix-Baum existiert.

2.4.2 Verschachtelte Suffixe und verzweigende Teilwörter

Für die folgende Beschreibung des Algorithmus von Ukkonen zur Konstruktion von Suffix-Bäumen benötigen wir noch einige Definitionen und Notationen, die in diesem Abschnitt zur Verfügung gestellt werden.

Definition 2.11 Sei $t \in \Sigma^*$. Ein Suffix $s \sqsubseteq t$ heißt verschachtelt (engl. nested), wenn sowohl $s\$ \sqsubseteq t\$$ als auch $sa \sqsubseteq t$ für ein $a \in \Sigma$ gilt.

Ein verschachteltes Suffix ist in Abbildung 2.14 schematisch dargestellt. Man beachte, dass anders als in der Skizze das verschachtelte Vorkommen mit dem Suffix auch überlappen kann.



Abbildung 2.14: Skizze: verschachteltes Suffix

Definition 2.12 Sei $t \in \Sigma^*$. Ein Teilwort $s \sqsubseteq t$ heißt rechtsverzweigend (engl. rightbranching), wenn es $a \neq b \in \Sigma$ mit $sa \sqsubseteq t$ und $sb \sqsubseteq t$ gibt.

2.4.3 Idee von Ukkonens Algorithmus

Sei $t = t_1 \cdots t_n \in \Sigma^*$. Der Algorithmus konstruiert dann sukzessive Suffix-Bäume $T^1, \dots, T^n$ mit $T^i = T(t_1 \cdots t_i)$. Nach Definition stellt T^i alle Teilwörter von $t_1 \cdots t_i$ und T^{i+1} alle Teilwörter von $t_1 \cdots t_i \cdot t_{i+1}$ dar.

T^1 ist offensichtlich leicht zu erstellen. Da jedes Teilwort von $t_1 \cdots t_i$ auch ein Teilwort von $t_1 \cdots t_i \cdot t_{i+1}$ ist, ist T^{i+1} quasi eine Erweiterung von T^i . Wir müssen uns nur noch überlegen, wie wir T^{i+1} aus T^i konstruieren können.

Im Folgenden verwenden wir die folgenden Abkürzungen: $x := t_1 \cdots t_i$ und $a := t_{i+1}$. Also müssen wir in T^i alle Teilwörter aus xa einfügen, die keine Teilwörter von x sind, um T^{i+1} zu erhalten. Dazu halten wir erst noch einige Notationen fest.

Notation 2.13 $I := \{w \in \Sigma^* : w \sqsubseteq xa \wedge w \not\sqsubseteq x\}$.

Halten wir zuerst die folgenden offensichtlichen Lemmata fest.

Lemma 2.14 *Alle Wörter aus I sind Suffixe von xa .*

Beweis: Dies folgt unmittelbar aus der Definition von I . ■

Lemma 2.15 *Sei $sa \in I$, dann ist $\overline{sa}$ in T^{i+1} ein Blatt.*

Beweis: Da sa nach dem vorhergehenden Lemma ein Suffix von xa ist und sa kein Teilwort von x ist, folgt die Behauptung. ■

Damit können wir noch folgende wichtige Notation einführen.

Notation 2.16 $I^* := \{sa \in I : \overline{s} \text{ ist kein Blatt in } T^i\} \subseteq I$.

Für alle Wörter in $sa \in I \setminus I^*$ ist nicht sonderlich viel zu tun, da es sich bei $\overline{s}$ um ein Blatt in T^i handelt. Für die Konstruktion von T^{i+1} wird an die Kantenmarkierung des zu $\overline{s}$ inzidenten Blattes nur ein a angehängt und die Bezeichnung dieses Blattes wird zu $\overline{sa}$. Somit sind alle Wörter aus $I \setminus I^*$ nun ebenfalls dargestellt. Die eigentliche Arbeit besteht also nur bei der Darstellung der Wörter aus I^* .

Lemma 2.17 *Für ein Suffix sa von xa gilt genau dann $sa \in I^*$, wenn s ein verschachteltes Suffix von x ist und sa kein Teilwort von x ist.*

Beweis: $\Rightarrow$: Da $sa \in I^* \subseteq I$ gilt, ist sa nach Definition von I kein Teilwort von x .

Da $sa \in I^*$ ist, ist $\bar{s}$ kein Blatt in $T^i = T(x)$. Somit muss s sowohl ein Suffix von x als auch ein echtes Teilwort von x sein, d.h. s ist ein verschachteltes Suffix von x .

$\Leftarrow$: Da s ein verschachteltes Suffix von x ist, ist sa ein Suffix von xa . Da außerdem $sa \not\sqsubseteq x$, folgt $s \in I$.

Da s ein verschachteltes Suffix von x ist, gibt es ein $b \in \Sigma$, so dass auch sb ein Teilwort von x ist. Somit kann $\bar{s}$ kein Blatt in $T^i = T(x)$ sein und es folgt $sa \in I^*$. ■

Definition 2.18 Das längste verschachtelte Suffix von x heißt aktives Suffix und wird mit $\alpha(x)$ bezeichnet.

Als Beispiel hierfür ist das aktive Suffix von $x = ababb$ gleich $\alpha(x) = b$ und von $x = abbab$ gleich $\alpha(x) = ab$. Beachte, dass für das leere Wort kein aktives Suffix definiert ist.

Lemma 2.19 Es gelten die folgenden Aussagen:

1. Ein Suffix s von x ist genau dann verschachtelt, wenn $|s| \leq |\alpha(x)|$.
2. Für alle Suffixe s von x gilt: $sa \in I^* \Leftrightarrow |\alpha(x)a| \geq |sa| > |\alpha(xa)|$
3. $\alpha(xa)$ ist ein Suffix von $\alpha(x) \cdot a$.
4. Ist $sa = \alpha(xa)$ und $s \neq \alpha(x)$, dann ist s ein verschachteltes Suffix, sogar ein rechtsverzweigendes Suffix von x .

Beweis: Zu 1.: $\Rightarrow$: Da $\alpha(x)$ das längste verschachtelte Suffix ist, folgt die Behauptung unmittelbar.

$\Leftarrow$: Ist s ein Suffix von x und gilt $|s| \leq |\alpha(x)|$, so ist s ein Suffix von $\alpha(x)$. Da $\alpha(x)$ ein verschachteltes Suffix von x ist, kommt es in x ein weiteres Mal als Teilwort vor und somit auch s . Also ist s ein verschachteltes Suffix. Dies ist in Abbildung 2.15 illustriert.

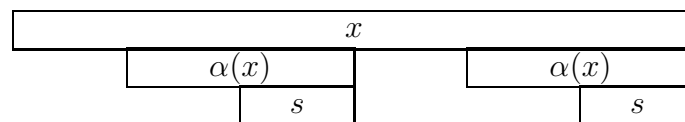


Abbildung 2.15: Skizze: Das Suffix s von x mit $|s| \leq |\alpha(x)|$ ist verschachtelt

Zu 2.: Es gilt offensichtlich mit Lemma 2.17 und mit Teil 1.:

$$\begin{aligned}
 sa \in I^* &\Leftrightarrow s \text{ ist ein verschachteltes Suffix von } x \text{ und } sa \text{ ist kein Teilwort von } x \\
 &\Leftrightarrow |s| \leq |\alpha(x)| \text{ und } sa \text{ ist kein verschachteltes Suffix von } xa \\
 &\Leftrightarrow |s| \leq |\alpha(x)| \wedge |sa| > |\alpha(xa)| \\
 &\Leftrightarrow |\alpha(x)a| \geq |sa| > |\alpha(xa)|
 \end{aligned}$$

Zu 3.: Offensichtlich sind $\alpha(xa)$ und $\alpha(x)a$ Suffixe von xa . Nehmen wir an, dass $\alpha(xa)$ kein Suffix von $\alpha(x) \cdot a$ wäre. Da $\alpha(xa)$ das längste verschachtelte Suffix von xa ist, muss es nochmals als Teilwort in x auftreten. Dann kann aber nicht $\alpha(x)$ das längste verschachtelte Suffix von x sein. Diese Situation ist in Abbildung 2.16 illustriert.

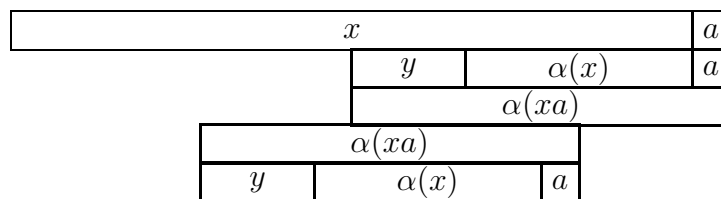


Abbildung 2.16: Skizze: Die Suffixe $\alpha(x) \cdot a$ und $\alpha(xa)$ von xa und ihre Wiederholungen in xa

Zu 4.: Sei also $\alpha(xa) = sa$ und $\alpha(x) \neq s$. Da $\alpha(xa) = sa$ ein verschachteltes Suffix von xa ist, muss auch s ein verschachteltes Suffix von x sein. Somit ist $|\alpha(x)| > |s|$ (da ja nach Voraussetzung $s \neq \alpha(x)$) und s ist ein echtes Suffix von $\alpha(x)$. Nach Definition taucht also $\alpha(x)$ (und damit auch s) ein weiteres Mal in x auf. Da sa ein längstes verschachteltes Suffix von xa ist, muss bei dem weiteren Vorkommen von $\alpha(x)$ direkt dahinter ein Zeichen $b \in \Sigma$ mit $b \neq a$ auftauchen (siehe auch Abbildung 2.17). Also ist s ein rechtsverzweigendes Suffix von x , da in x sowohl sa als auch sb auftritt. Diese Situation ist in Abbildung 2.17 illustriert. ■

14.11.06

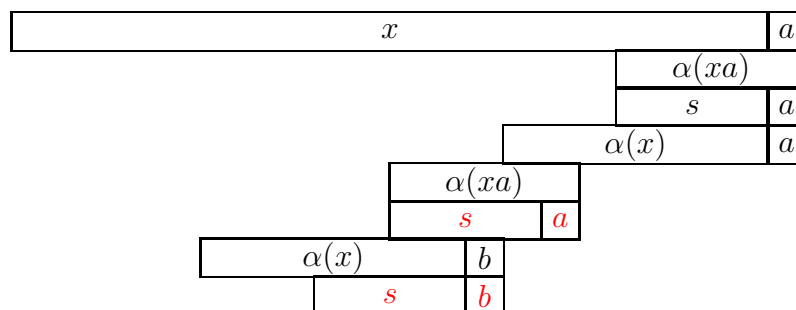


Abbildung 2.17: Skizze: Die Suffixe $\alpha(x) \cdot a$ und $\alpha(xa)$ von xa

2.4.4 Ukkonens Online Algorithmus

Die Aussage 2 des Lemmas 2.19 gibt uns eine Charakterisierung derjenigen Suffixe, die in T^i noch dargestellt werden müssen, um T^{i+1} zu erhalten. Mit Hilfe von Aussage 3 des Lemmas 2.19 wissen wir weiter, dass die noch darzustellenden Suffixe Suffixe des alten aktiven Suffixes verlängert um a sind.

Wie können wir nun alle neu darzustellenden Suffixe von xa finden? Wir beginnen mit dem Suffix $\alpha(x)$ in T^i und versuchen dort das Zeichen a anzuhängen. Anschließend durchlaufen wir alle Suffixe von $\alpha(x)$ in abnehmender Länge (was sich elegant durch die Suffix-Links erledigen lässt) und hängen jeweils ein a an. Wir enden mit dieser Prozedur, wenn für ein Suffix s von $\alpha(x)$ das Wort sa bereits im Suffix-Baum dargestellt ist, d.h. die einzufügende Kante mit Kantenmarkierung a bereits existiert (bzw. eine Kante, deren Kantenmarkierung nach (einem Suffix von) s bereits ein a enthält). Nach dem vorhergehenden Lemma erhalten wir durch Ablaufen der Kantenmarkierung a den aktiven Suffix $\alpha(xa)$ von T^{i+1} . Diese Position wird dann als Startposition für die Erweiterung auf den Suffix-Baum T^{i+1} verwenden.

Wir müssen jedoch noch berücksichtigen, dass wir auf einen Suffix s von $\alpha(x)$ treffen, wobei $\bar{s}$ kein Knoten in T^i ist. In diesem Fall müssen wir in die Kante, an der die Darstellung von s endet, Aufbrechen und an der geeigneten Stelle einen neuen Knoten einfügen.

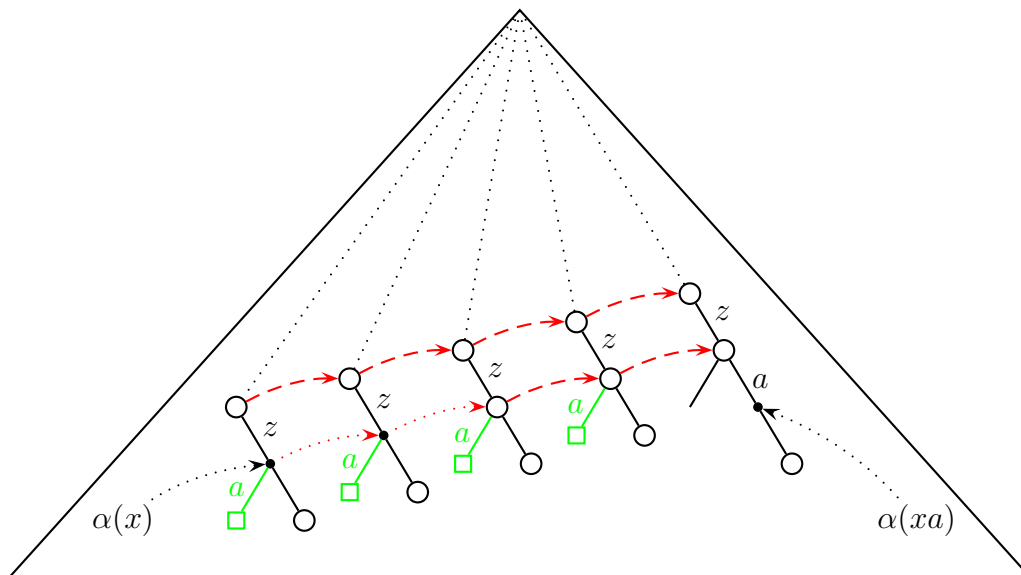
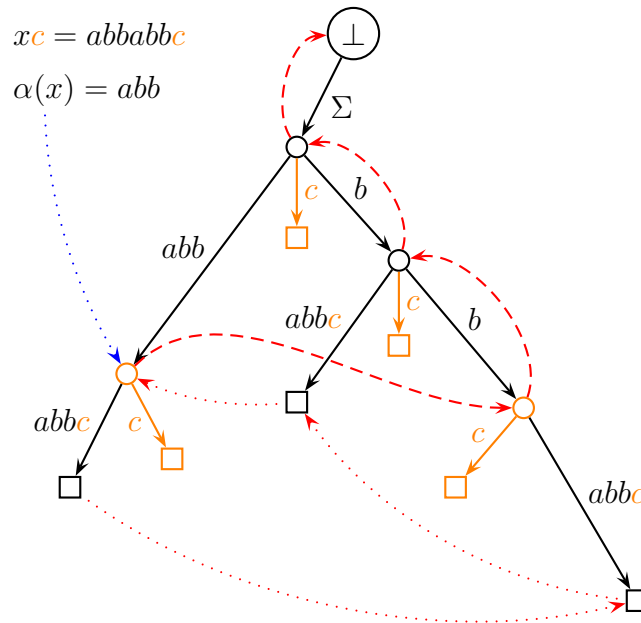


Abbildung 2.18: Skizze: Übergang von T^i zu T^{i+1} in Ukkonens Algorithmus

Die Erweiterung eines Suffix-Baumes für x zu einem Suffix-Baum für xa ist in Abbildung 2.18 schematisch dargestellt. Dort sind die rot gepunkteten Linien die virtuellen Suffix-Links, denen wir gerne folgen möchten, die aber gar nicht existieren, da

Abbildung 2.19: Beispiel: Konstruktion von $T(abbabbc)$ aus $T(abbabb)$

deren Anfangspunkt in T^i gar nicht zu einem Knoten im Baum T^i gehören. Die rot gestrichelten Linien sind wirklich vorhandene Suffix-Links in T^i .

In Abbildung 2.19 ist ein Beispiel für eine Erweiterung eines Suffix-Baumes angegeben. Dort wird auch das mehrfache Aufbrechen von Kanten illustriert. Die rot gepunkteten Linien entsprechen virtuelle Suffix-Links, die rot gestrichelten Linien echten Suffix-Links. Die Erweiterung ist orangefarben dargestellt.

In diesem Beispiel haben wir noch eine virtuelle Superwurzel $\perp$ eingeführt. Mit $\text{slink}(\bar{\epsilon}) = \perp$ können wir sicherstellen, dass jeder innere Knoten einen Suffix-Link besitzt. Dazu nehmen wir noch eine Baum-Kante von $\perp$ zu $\bar{\epsilon}$ an, die mit allen Zeichen aus Σ markiert ist. Damit stellen wir sicher, dass wir wieder zur Wurzel zurückkehren können, wenn wir eigentlich fälschlicherweise dem Suffix-Link von $\bar{\epsilon}$ zu $\perp$ gefolgt sind.

Wie wir noch sehen werden, ist der einzige Grund für die Einführung von $\perp$, dass die Beschreibung von Ukkonens Algorithmus einheitlicher wird, da wir beim Folgen der Suffix-Links dann immer irgendwann auf einen Knoten stoßen werden, von dem eine Kante mit dem Zeichen $a \in \Sigma$ ausgeht.

Wie finden wir in diesem Beispiel den Suffix-Link von $\alpha(x) = abb$? Im Suffix-Baum für $abbabb$ gibt es ja keinen Knoten $\overline{abb}$. Wir folgen stattdessen dem Suffix-Link des Knotens, der als erstes oberhalb liegt: $\bar{\epsilon}$. Somit landen wir mit dem Suffix-Link im Knoten $\perp$. Vom Knoten $\bar{\epsilon}$ mussten wir ja noch die Zeichenreihe abb ablesen um bei

```

UKKONEN (string  $t = t_1 \cdots t_n$ )
{
    tree  $T = T(t_1)$ ;
    for ( $i = 2$ ;  $i \leq n$ ;  $i++$ )
    {
        string  $x = t_1 \cdots t_{i-1}$ ;
        char  $a = t_i$ ;
        string  $s = \alpha(x)$ ;
        while ( $sa$  is not represented in  $T$ ) do
        {
            insert  $\overline{sa}$  in  $T$ ;
             $s = s_2 \cdots s_{|s|}$ ;    /* via suffix-links */
                                   /* Note that  $s_2 \cdots s_{|s|} := \perp$  iff  $s = \varepsilon$  */
        }
         $\alpha(xa) = sa$ ;           /* Note that  $\perp \cdot a := \varepsilon$  for any  $a \in \Sigma$  */
    }
}

```

Abbildung 2.20: Algorithmus: Abstrakte Fassung von Ukkonens Algorithmus

der Position von $\alpha(x)$ zu landen. Dies müssen wir jetzt auch vom Knoten $\perp$ aus tun. Damit landen wir dann eigentlich beim Knoten $\overline{bb}$, der aber dummerweise im Suffix-Baum für *abbabb* auch nicht existiert (der aber jetzt eingefügt wird).

Für das weitere Folgen des Suffix-Links von $\overline{bb}$ (wobei es den Suffix-Link ja noch nicht gibt) starten wir wieder von $\overline{b}$ aus. Nach Folgen des Suffix-Links landen wir in $\overline{\varepsilon}$. Nun laufen wir noch das Wort *b* ab, um die eigentlich Position $\overline{b}$ im Suffix-Baum zu finden, an der wir weitermachen. Die folgenden Schritte sind nun einfach und bleiben dem Leser zur Übung überlassen.

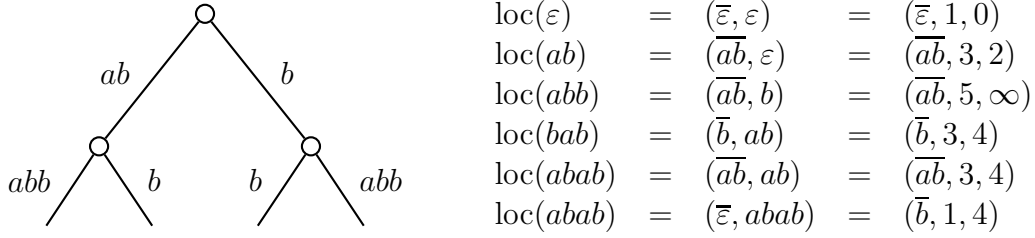
Die sich aus dieser Diskussion ergebende, erste abstrakte Fassung von Ukkonens Algorithmus ist in Abbildung 2.20 angegeben.

Um dies etwas algorithmischer beschreiben zu können, benötigen wir erst noch eine angemessene Darstellung der vom Suffix-Baum T dargestellten Wörter, die wir als Lokation bezeichnen wollen (in Zeichen $\text{loc}(v)$).

Definition 2.20 Sei $t \in \Sigma^*$ und sei $T = T(t)$ der zugehörige Suffix-Baum. Für $s \in \text{words}(T)$ ist $\text{loc}(s) = (\overline{u}, v) \equiv (\overline{u}, j + |u|, j + |s| - 1)$ eine Lokation von s in T , wenn $\overline{u} \in V(T)$, $s = uv$ für ein $v \in \Sigma^*$ und $j = \min \{k : t_k \cdots t_{k+|s|-1} = s\}$ gilt. Ist $\overline{s}$ ein Blatt in T , dann muss darüber hinaus $v \neq s$ gelten

Für eine Lokation $(\bar{v}, i, j)$ eines Blattes schreiben wir darüber hinaus im Folgenden $(\bar{v}, i, \infty)$, da für ein Blatt immer $j = |t|$ gilt. Diese Schreibweise bezeichnet man als *offene Lokation* bzw. auch als *offene Referenz*, da der Endpunkt keine Rolle spielt. Das ist der Grund, warum Lokationen keine Blätter verwenden dürfen.

Die Definition wollen wir uns in Abbildung 2.21 noch einmal an einem Beispiel genauer anschauen.



Abbildungung 2.21: Beispiel: Einige Lokationen im Suffix-Baum für *ababb*

Für unseren Algorithmus benötigen wir insbesondere so genannte kanonische Lokationen.

Definition 2.21 Eine Lokation $\text{loc}(s) = (\bar{v}, w) \equiv (\bar{v}, i, j)$ heißt kanonisch, wenn für jede andere Lokation $\text{loc}(v) = (\bar{v}', w') \equiv (\bar{v}', i', j')$ gilt, dass $|v| > |v'|$.

Für die in Abbildung 2.21 angegebenen Lokationen sind die ersten fünf kanonische, die sechste jedoch nicht.

Für eine kanonische Lokation gilt also das Folgende:

- Wenn $\bar{s}$ ein Verzweigungsknoten in T ist, dann gilt

$$\text{loc}(s) := (\bar{s}, \varepsilon) \equiv (\bar{s}, j + |s|, j + |s| - 1),$$

wobei $j = \min \{k : t_k \cdots t_{k+|s|-1} = s\}$.

- Wenn $\bar{s}$ ein Blatt ist, d.h. $\bar{u} \xrightarrow{v} \bar{s}$ und $s = uv$, dann gilt

$$\text{loc}(s) := (\bar{u}, v) \equiv (\bar{u}, j + |u|, \infty),$$

wobei $j = \min \{k : t_k \cdots t_{k+|uv|-1} = uv = s\}$.

- Wenn kein Knoten $\bar{s}$ in T existiert, dann existiert eine Kante $\bar{u} \xrightarrow{vw} \overline{uvw}$ mit $s = uv$, $v \neq \varepsilon$, $w \neq \varepsilon$ und es gilt

$$\text{loc}(s) := (\bar{u}, v) \equiv (\bar{u}, j + |u|, j + |s| - 1),$$

wobei $j = \min \{k : t_k \cdots t_{k+|uv|+1} = uv = s\}$.

```

UKKONEN (string  $t = t_1 \cdots t_n$ )
{
    tree  $T = T(t_1)$ ;
    ref  $(\bar{v}, w) = (r(T), \varepsilon)$ ;    /*  $r(T)$  is the root of  $T$  */
    for ( $i = 2$ ;  $i \leq n$ ;  $i++$ )
    {
        node  $x = y = \text{NIL}$ ;
        while (not  $T.\text{lookup}((\bar{v}, w), t_i)$ )    /* while  $((\bar{v}, w \cdot t_i) \notin T)$  */
        {
             $y = T.\text{insert}((\bar{v}, w), t_i)$ ;    /* returns node adjacent to new leaf */
            if ( $x$ )  $\text{slink}(x) = y$ ;
             $x = y$ ;
             $(\bar{v}, w) = \text{canonize}((\text{slink}(\bar{v}), w))$ ;
        }
        if ( $x$ )  $\text{slink}(x) = T.\text{lookup}((\bar{v}, w), t_i)$ ;
         $(\bar{v}, w) = \text{canonize}(\bar{v}, w \cdot t_i)$ ;
    }
}

```

Abbildung 2.22: Algorithmus: Ukkonens Algorithmus

Aus unserer Diskussion ergibt sich nun Ukkonens Algorithmus, der im Pseudo-Code in Abbildung 2.22 angegeben ist.

Hierbei schaut die Prozedur $\text{lookup}((\bar{v}, w), a)$ nach, ob im Baum T das Wort vwa dargestellt wird, d.h. ob man ab der kanonischen Lokation $(\bar{v}, w)$ in T den Buchstaben a weiterverfolgen kann. Ist $w = \varepsilon$, dann wird nur überprüft, ob $\bar{v}$ eine ausgehende Kante hat deren Kantenmarkierung mit a beginnt. Ist $w \neq \varepsilon$, dann wird die ausgehende Kante von $\bar{v}$ betrachtet, deren Kantenmarkierung mit dem ersten Buchstaben von w beginnt. In dieser Kantenmarkierung wird jetzt nur der Buchstabe an Position $|w| + 1$ mit a verglichen. Dabei liefert lookup einen Null-Zeiger, wenn $(\bar{v}, w \cdot a)$ keine Lokation in T ist, andernfalls wird der Knoten $(\bar{v}, w) = \overline{vw}$ geliefert. Der Leser möge sich selbst davon überzeugen, dass in dem im Algorithmus benötigten Fall (nach dem Ende der while-Schleife) der Knoten $\overline{vw}$ in T wirklich existieren muss.

Man beachte, dass die Buchstaben an den Positionen 2 mit $|w|$ nicht verglichen werden müssen, da diese nach Konstruktion übereinstimmen müssen. Dies folgt aus der Tatsache, dass das zur Lokation $(\bar{v}, w)$ gehörige Wort vw nach Konstruktion von Ukkonens Algorithmus in T dargestellt sein muss.

Die Prozedur canonize macht aus einer gegebenen Lokation $(\bar{v}, w)$ eine kanonische, indem sie vom Knoten $\bar{v}$ soweit im Suffix-Baum die Zeichenreihe w verfolgt, bis die kanonische Lokation gefunden wird. Man beachte hierbei, dass der Aufwand

proportional zur Anzahl überlaufener Kanten plus 1 ist, da jeweils nur das erste Zeichen einer ausgehenden Kante betrachtet werden muss (die folgenden Zeichen dieser Kantenmarkierung müssen dann, wie vorher schon diskutiert, auch wieder übereinstimmen).

Wie man in Abbildung 2.19 sehen kann, kann es auch wirklich passieren, dass man bei `canonize` mehrere Kanten überlaufen muss. Dort ist die kanonische Lokation des Suffixes $\alpha(x) = abb$ gerade $(\bar{\epsilon}, abb)$. Folgt man dem zugehörigen Suffix-Link, so gelangt man zur Lokation $(\perp, abb)$. Um die zugehörige kanonische Lokation $(\bar{b}, b)$ zu erhalten, muss man zwei Knoten im Baum hinabsteigen.

Die Prozedur `insert(( $\bar{v}, w$ ),  $a$ )` fügt an der Lokation $(\bar{v}, w)$ eine Kante mit Kantenmarkierung a zu einem neuen Blatt ein. War $(\bar{v}, w)$ bereits ein Knoten im Suffix-Baum, so wird die neue Kante dort nur angehängt. Beschreibt die Lokation $(\bar{v}, w)$ eine Position innerhalb einer Kante, so wird in diese Kante an der zur Lokation $(\bar{v}, w)$ entsprechenden Stelle ein neuer Knoten eingefügt, an den die Kante zu dem neuen Blatt angehängt wird. Man überlegt sich leicht, dass dies in konstanter Zeit erledigt werden kann. Der zu dem neuen Blatt adjazente innere Knoten kann dann leicht wie gefordert zurückgeliefert werden.

Dabei muss auch noch ein Suffix-Link vom vorhergehenden Knoten auf diesen neu eingefügten Knoten gesetzt werden. Mit einer geeigneten Buchhaltung kann auch dies in konstanter Zeit erledigt werden, da der Knoten, der den neuen Suffix-Link auf den neu eingefügten bzw. aktuell betrachteten Knoten erhält, gerade eben vorher betrachtet wurde.

2.4.5 Zeitanalyse

Es bleibt noch die Analyse der Zeitkomplexität von Ukkonens Algorithmus. Das Einfügen eines neuen Knotens kann, wie bereits diskutiert, in konstanter Zeit geschehen. Da ein Suffix-Baum für t maximal $O(|t|)$ Knoten besitzt, ist der Gesamtzeitbedarf hierfür $O(|t|)$.

Der Hauptaufwand liegt in der Kanonisierung der Lokationen. Für einen Knoten kann eine Kanonisierung nämlich mehr als konstante Zeit kosten. Bei der Kanonisierung wird jedoch die Lokation $(\bar{v}, w)$ zu $(\overline{vv'}, w')$ für $v' \in \Sigma^+$ und $w' \in \Sigma^*$ mit $v'w' = w$. Damit wandert das Ende der Zeichenreihe v des in der Lokation verwendeten internen Knotens bei einer Kanonisierung immer weiter zum Wortende von t hin. Zwar wird $\bar{v}$ beim Ablaufen der Suffix-Links verkürzt, aber dies geschieht nur am vorderen Ende. Somit kann eine Verlängerung am Wortende von v bei der verwendeten Lokation $(\bar{v}, w)$ maximal $|t|$ Mal auftreten. Also haben alle Aufrufe zur Kanonisierung eine lineare Laufzeit.

Theorem 2.22 *Ein Suffix-Baum kann mit Ukkonens Algorithmus in Zeit $O(n)$ mit Platzbedarf $O(n)$ konstruiert werden.*

Wir erwähnen hier noch, dass der Algorithmus folgende Eigenschaft einer Darstellung eines Suffix-Baumes benötigt: Man muss von einem Knoten das Kind finden können, das sich über diejenige Kanten erreichen lässt, dessen erstes Zeichen der Kantenmarkierung vorgegeben ist. Verwendet man also für die Repräsentation aus Platzgründen keine Felder, so erhöht sich der Aufwand in der Regel um den Faktor $|\Sigma|$, was bei kleinen Alphabeten noch tolerierbar ist.

Des Weiteren wollen wir noch anmerken, dass in manchen Lehrbüchern eine Laufzeit für Ukkonens Algorithmus von $O(n \log(n))$ angegeben wird. Dabei wird jedoch statt des uniformen Kostenmaßes das logarithmische Kostenmaß verwendet (also die Bit-Komplexität). In diesem Fall hat jedoch der Suffix-Baum selbst schon die Größe $O(n \log(n))$, da dort ja Zahlen (Referenzen und Zeiger) aus dem Intervall $[0 : n]$ vorkommen, deren Darstellung ja $\Theta(\log(n))$ Bits benötigt. Daher handelt es sich auch in diesem Modell um einen asymptotisch optimalen Algorithmus zur Erstellung von Suffix-Bäumen, da die Ausgabegröße bereits $O(n \log(n))$ betragen kann. Man beachte jedoch, dass die Eingabegröße nur $\Theta(n \log(|\Sigma|))$ beträgt. Im logarithmischen Kostenmaß ist also der Suffix-Baum im Allgemeinen größer als die Eingabe!

16.11.06

Repeats

3.1 Exakte und maximale Repeats

In diesem Kapitel wollen wir uns mit Wiederholungen (so genannten) Repeats in Zeichenreihen (insbesondere in Genomen) beschäftigen. Dazu müssen wir uns vor allem überlegen, wie man interessante Wiederholungen charakterisiert.

3.1.1 Erkennung exakter Repeats

Zunächst einmal müssen wir formal definieren, was wir unter einer Wiederholung verstehen wollen. Dazu benötigen wir erst noch die Begriffe einer Hamming- und Alignment-Distanz.

Definition 3.1 Seien $s, t \in \Sigma^n$, dann ist die Hamming-Distanz von s und t definiert als $\delta_H(s, t) := |\{i \in [1 : n] : s_i \neq t_i\}|$.

Definition 3.2 Seien $s, t \in \Sigma^*$, dann ist die Alignment-Distanz von s und t definiert als $\delta_A(s, t) := \min \{k \in \mathbb{N}_0 : \sigma(\hat{s}, \hat{t})\}$, wobei $(\hat{s}, \hat{t})$ ein Alignment von s und t ist und $\sigma(\hat{s}, \hat{t})$ eine Alignment-Distanz ist

Als Alignment-Distanz kann beispielsweise die Anzahl von Substitutionen, Insertionen und Deletionen im Alignment $(\hat{s}, \hat{t})$ (also die EDIT-Distanz) verwendet werden. Nun können wir definieren, was eine Wiederholung sein soll.

Definition 3.3 Sei $t = t_1 \cdots t_n \in \Sigma^*$ eine Zeichenreihe. Ein Paar $((i_1, j_1), (i_2, j_2))$ heißt

- exaktes Paar, wenn $t_{i_1} \cdots t_{j_1} = t_{i_2} \cdots t_{j_2}$. Das Wort $t_{i_1} \cdots t_{j_1}$ wird dann auch als exaktes Repeat bezeichnet. Mit $\mathcal{R}(t)$ bezeichnen wir die Menge aller exakten Paare und mit $\overline{\mathcal{R}}(t)$ die Menge aller exakten Repeats.
- k -mismatch Repeat, wenn $\delta_H(t_{i_1} \cdots t_{j_1}, t_{i_2} \cdots t_{j_2}) \leq k$. Mit $\mathcal{R}_k^H(t)$ bezeichnen wir die Menge aller k -mismatch Repeats.
- k -difference Repeat, wenn $\delta_A(t_{i_1} \cdots t_{j_1}, t_{i_2} \cdots t_{j_2}) \leq k$. Mit $\mathcal{R}_k^A(t)$ bezeichnen wir die Menge aller k -difference Repeats.

In Abbildung 3.1 sind schematisch zwei solcher Repeats dargestellt. Man beachte dabei, dass nach Definition ein Repeat sich mit sich selbst überlappen kann (wie bei den roten Teilwörtern) oder es auch mehr als zwei Vorkommen eines Repeats geben kann.

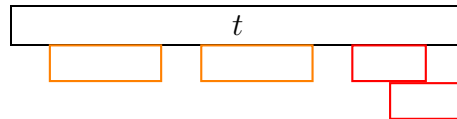


Abbildung 3.1: Skizze: Schematische Darstellung von Repeats

Wir behaupten nun, dass wir alle solchen exakten Repeats in „linearer Zeit“ finden können. Dies folgt aus der Tatsache, dass alle exakten Repeats von t zu Lokationen eines Suffix-Baums $T(t\$)$ korrespondieren, die sich nicht auf einer zu einem Blatt inzidenten Kante befinden. Sobald man sich im Suffix-Baum nicht auf einer zu einem Blatt inzidenten Kanten befindet, gibt es Verlängerungsmöglichkeiten, so dass man in mindestens zwei *verschiedenen* Blättern landen kann. Also ist das betrachtete Wort jeweils ein Präfix von mindestens zwei verschiedenen Suffixen, d.h. das betrachtete Wort beginnt an mindestens zwei verschiedenen Positionen in t .

In Abbildung 3.2 ist für das Wort $abbbab$ der Suffix-Baum für $abbbab\$$ angegeben und einige darin enthaltene exakte Repeats durch Angabe von exakten Paaren.

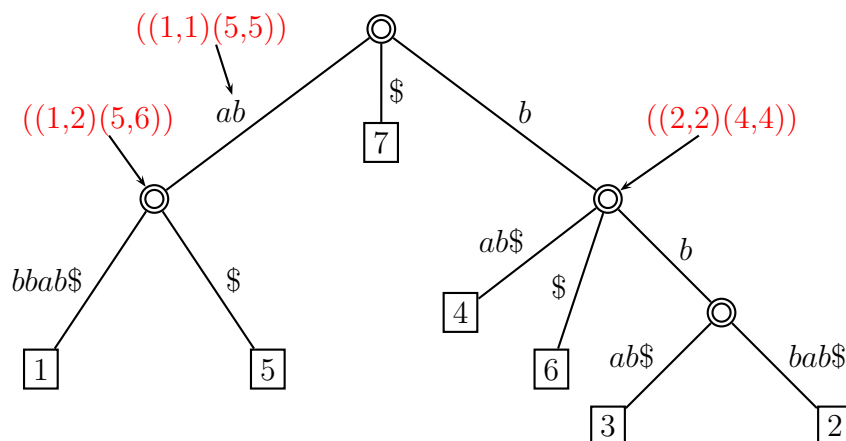


Abbildung 3.2: Beispiel: Suffix-Baum für $t\$ = abbbab\$$ und die exakten Repeats

Durchlaufen wir nun den Suffix-Baum $T(t\$)$ gekürzt um alle Blätter und die dazu inzidenten Kanten mit einer Tiefensuche, so können wir alle dort zu den noch enthaltenen Lokationen korrespondierenden Zeichenreihen ausgeben, und geben damit alle exakten Repeats aus.

Theorem 3.4 *Sei $t \in \Sigma^n$, lassen sich alle exakten Repeats als Wörter bzw. Referenzen in Zeit $O(n + k)$ ermitteln, wobei k die Anzahl aller Zeichen in den exakten Repeats bzw. die Anzahl der exakten Repeats ist.*

Somit ist hier mit linearer Zeit gemeint, dass der Algorithmus linear in der Eingabe- und Ausgabegröße läuft. Man beachte, dass es durchaus Wörter über Σ der Länge n geben kann, die $\Theta(n^2)$ exakte Repeats besitzen. Wir merken noch an, dass dieser Algorithmus optimal ist.

Somit haben wir aber nur die exakten Repeats ausgegeben, wir wissen jedoch nicht wo diese auftreten. Daher ist in der Regel ein Algorithmus zu Auffinden aller exakten Paare interessanter. Auch diesen können wir analog zu denen der exakten Repeats definieren. Der Algorithmus zum Auffinden aller exakten Paare geht dabei wie folgt vor:

1. Konstruktion von $T(t\$)$ in Zeit $O(n)$.
2. Tiefensuche durch $T(t\$)$ in Zeit $O(n)$. Während der Tiefensuche führe folgende Schritte aus:
 - (a) Jedes Blatt liefert die Indexposition zurück, an der das aufgefundene Suffix beginnt. Der Zeitbedarf hierfür ist $O(n)$.
 - (b) Jeder innere Knoten liefert eine Liste der Blätter, die von diesem Knoten erreichbar sind, an seinen Elter zurück. Der Zeitbedarf hierfür ist $O(|n|)$, da zum einen immer nur ein Zeiger auf den Beginn und das Ende der Liste übergeben wird und zum anderen, da das Zusammenhängen der Listen der Kinder insgesamt in Zeit $O(|n|)$ zu realisieren ist (da man auch das Ende der Liste kennt).
 - (c) Für einen inneren Knoten $\bar{v}$ mit der Blattliste L und mit Elter $\bar{w}$ generiere die exakten Paare $((i + \ell - 1), (j, j + \ell - 1))$ mit $i \neq j \in L$ und $\ell \in [|w| + 1 : |v|]$. Dies lässt sich in Zeit $O(n + |\text{output}|)$ erledigen. Nach der vorherigen Diskussion beschreibt jedes ausgegebene exakte Paar einen exakten Repeat. Man überlegt sich leicht, dass alle ausgegeben exakten Paare paarweise verschieden sind.

Theorem 3.5 *Sei $t \in \Sigma^n$, dann lassen sich alle exakten Paare in Zeit $O(n + k)$ ermitteln, wobei k die Anzahl der exakten Paare ist.*

Wir merken noch an, dass dieser Algorithmus optimal ist.

Betrachten wir das Wort $a^n \in \Sigma^*$, dann enthält es nach unserer Definition $\Theta(n^3)$ exakte Paare. Es gilt nämlich, dass für alle $\ell \in [1 : n - 1]$ und $i_1 < i_2 \in [1 : n - \ell + 1]$

das Paar $((i_1, i_1 + \ell - 1), (i_2, i_2 + \ell - 1))$ ein exaktes Paar ist. Andererseits existieren nur $n - 1$ exakte Repeats a^i in a^n mit $i \in [1 : n - 1]$.

Somit ist auch hier mit linearer Zeit gemeint, dass der Algorithmus linear in der Eingabe- und Ausgabegröße läuft.

Man kann diesen Algorithmus auch leicht so modifizieren, dass er nur exakte Paare ausgibt, die einen Repeat mit einer Mindestlänge ausgibt. Man muss dann nur innere Knoten berücksichtigen, deren Worttiefe hinreichend groß ist.

3.1.2 Charakterisierung maximaler Repeats

Wie wir gesehen haben, gibt es für die Bestimmung exakter Repeats bzw. Paare einen optimalen Algorithmus. Dennoch muss dieser nicht sehr effizient sein, insbesondere dann nicht, wenn viele exakte Repeats bzw. Paare vorkommen. Daher werden wir die Problemstellung überarbeiten und versuchen uns auf die interessanten Repeats zu beschränken.

Definition 3.6 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei $t' = t'_0 \cdots t'_{n+1} = \#t_1 \cdots t_n\$$ mit $\# \neq \$ \notin \Sigma$. Ein Tripel (i, j, ℓ) mit $\ell \in [1 : n]$ sowie $i < j \in [1 : n - \ell + 1]$ heißt maximales Paar, wenn $t_i \cdots t_{i+\ell-1} = t_j \cdots t_{j+\ell-1}$ und $t'_{i-1} \neq t'_{j-1}$ sowie $t'_{i+\ell} \neq t'_{j+\ell}$ gilt. Ist (i, j, ℓ) ein maximales Paar, dann ist $t_i \cdots t_{i+\ell-1}$ ein maximales Repeat.

Mit $\mathcal{R}_{\max}(t)$ bezeichnet man die Menge aller maximalen Paare von t und mit $\overline{\mathcal{R}}_{\max}(t) = \{t_i \cdots t_{i+\ell-1} : (i, j, \ell) \in \mathcal{R}_{\max}(t)\}$ die Menge aller maximalen Repeats von t .

Anschaulich bedeutet dies, dass ein Repeat maximal ist, wenn jede Verlängerung dieses Repeats um ein Zeichen kein Repeat mehr ist. Es kann jedoch Verlängerungen geben, die wiederum ein maximales Repeat sind. Für $t = a\boxed{aa}bb\boxed{aa}ab$ ist aa mit dem maximalen Paar $(2, 6, 2)$ ein maximales Repeat, aber auch $aaab$ ist ein maximales Repeat. Des Weiteren ist für aa das Paar $(2, 7, 2)$ nicht maximal.

Lemma 3.7 Sei $t \in \Sigma^*$ und sei $T = T(\#t\$)$ der Suffix-Baum für $\#t\$$ mit $\#, \$ \notin \Sigma$. Wenn $\alpha \in \overline{\mathcal{R}}_{\max}(t)$, dann existiert ein innerer Knoten v in T mit $\text{path}(v) = \alpha$.

Beweis: Betrachten wir das zu α gehörige maximale Paar (i, j, ℓ) . Nach Definition des maximalen Paares gilt insbesondere $t'_{i+\ell} \neq t'_{j+\ell}$. Somit muss also α ein rechtsverzweigendes Teilwort von t' und daher muss $\overline{\alpha}$ ein innerer Knoten in T sein. ■

Wir erhalten damit unmittelbar noch das folgende Korollar.

Korollar 3.8 *Ein Wort $t \in \Sigma^*$ besitzt höchstens $|t|$ viele maximale Repeats.*

Beweis: Jedem maximalen Repeat entspricht nach dem vorhergehenden Lemma ein innerer Knoten und in einem Suffix-Baum für $\phi t \$$ mit $|t| + 2$ Blättern kann es maximal $|t| + 1$ viele innere Knoten geben. In der obigen Formel kann nur dann das Maximum angenommen werden, wenn jeder innere Knoten genau zwei Kinder hat. Da die Wurzel jedoch mindestens drei Kinder hat (über die Kanten, die mit ϕ , $\$$ und einem $a \in \Sigma$ beginnen), kann es maximal $|t|$ innere Knoten geben. ■

Wir merken noch an, dass die Anzahl der maximalen Paare keineswegs linear in der Länge des Textes t begrenzt sein muss. Dies sei dem Leser zur Übung überlassen.

Die Beweisidee des Lemmas ist noch einmal in Abbildung 3.3 illustriert. Man sollte sich auch klar machen, dass es für ein maximales Repeat α zwei Vorkommen in t existieren müssen, an denen sich die angrenzenden Zeichen unterscheiden (d.h. $t'_{i-1} \neq t'_{j-1}$ und $t'_{i+\ell} \neq t'_{j+\ell}$, wenn das maximale Paar (i, j, ℓ) das maximale Repeat α beschreibt). Es kann durchaus zwei Vorkommen von α in t geben, so dass sich die Wörter nach vorne oder hinten zu längeren Repeats verlängern lassen.

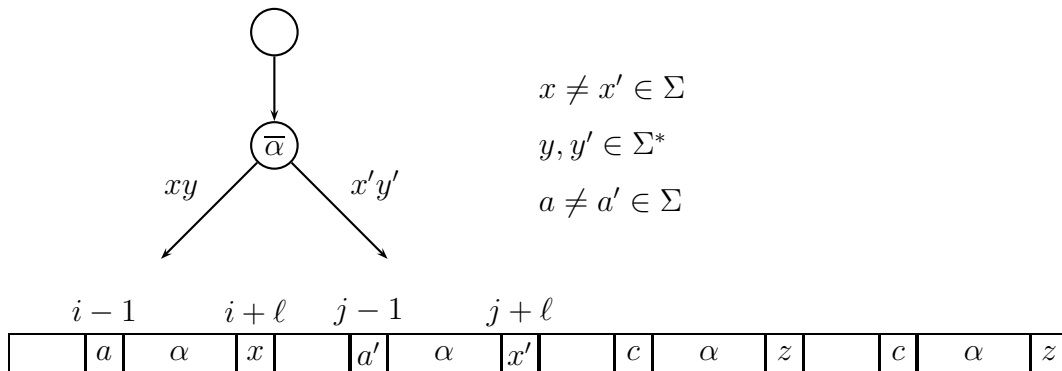


Abbildung 3.3: Skizze: Maximale Repeats enden an inneren Knoten

Wie man dort bemerkt, nutzen wir hier nur aus, dass nach einem maximalen Repeat eines maximalen Paares (i, j, ℓ) die Zeichen unterschiedlich sind (d.h. $t'_{i+\ell} \neq t'_{j+\ell}$), aber nicht die Zeichen unmittelbar davor (d.h. $t'_{i-1} \neq t'_{j-1}$). Um diese bei der Bestimmung maximaler Repeats auch noch berücksichtigen zu können, benötigen wir noch die folgende Definition.

Definition 3.9 Sei $t \in \Sigma^n$ und sei $t' = t'_0 \cdots t'_{n+1} = \mathfrak{t}t\$$.

- Das Linkszeichen von Position $i \in [1 : n]$ ist definiert als t'_{i-1} .
- Das Linkszeichen eines Blattes $\bar{s} \neq \bar{\mathfrak{t}}$ von $T(\mathfrak{t}t\$)$ ist das Zeichen $t'_{n-|s|+1}$.
- Ein innerer Knoten von $T(\mathfrak{t}t\$)$ heißt linksdivers, wenn in seinem Teilbaum zwei Blätter mit einem verschiedenen Linkszeichen existieren.

In Abbildung 3.4 ist noch einmal die Definition der Linkszeichen von Blättern illustriert.

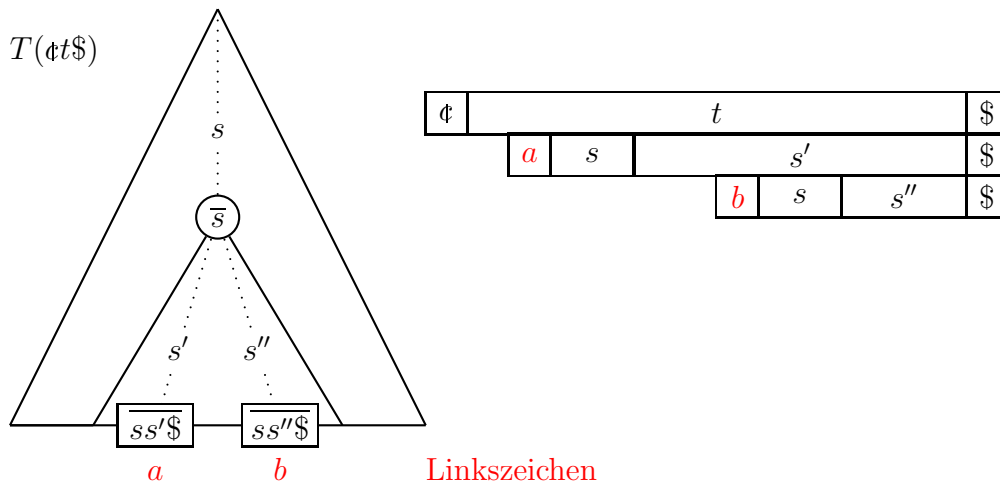


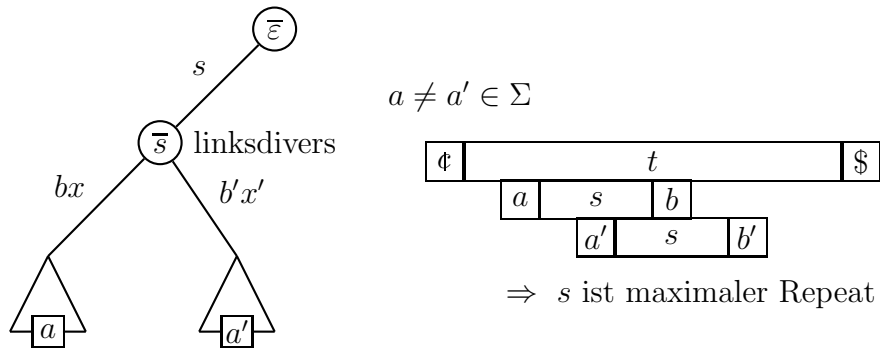
Abbildung 3.4: Skizze: Definition von Linkszeichen

In der Definition ist für einen linksdiversen Knoten v nur gefordert, dass im Teilbaum von v zwei verschiedenen Linkszeichen auftreten müssen. Man kann sich aber leicht überlegen, dass es dann sogar zwei Blätter in den Teilbäumen von zwei verschiedenen Kindern von v mit unterschiedlichem Linkszeichen geben muss.

Theorem 3.10 Sei $t \in \Sigma^*$, $t' = \mathfrak{t}t\$$ und $T = T(t')$ der Suffix-Baum für t' . Die Zeichenreihe $s \in \Sigma^*$ ist genau dann ein maximaler Repeat von t , wenn der Knoten $\bar{s}$ in T existiert und linksdivers ist.

Beweis: $\Leftarrow$: Da $\bar{s}$ linksdivers ist, muss es zwei Blätter $\bar{v}$ und $\bar{w}$ im Teilbaum von $\bar{s}$ geben, die unterschiedliche Linkszeichen besitzen.

Wir halten zunächst fest, dass $\bar{s}$ zwei verschiedene Kinder besitzen muss, in deren Teilbäumen die Blätter mit den unterschiedlichen Linkszeichen auftreten müssen. Dieser Fall ist in Abbildung 3.5 illustriert.

Abbildung 3.5: Skizze: $\bar{s}$ ist linksdivers

Seien bx und $b'x'$ mit $b \neq b' \in \Sigma$ und $x, x' \in \Sigma^*$ die beiden Kantemarkierungen zu den zwei Kindern, die die verschiedenen Linkszeichen im Unterbaum besitzen. Weiter seien $a \neq a' \in \Sigma$ die beiden verschiedenen Linkszeichen, wobei a das Linkszeichen des Blattes ist, der im über die Kante bx erreichbaren Teilbaum liegt. Dann muss sowohl asb als auch $a'sb'$ ein Teilwort von $t' = \phi t \$$ sein. Somit gibt es ein maximales Paar, das genau diese Vorkommen in t' beschreibt und somit ist s ein maximaler Repeat.

$\Rightarrow$: Sei nun s ein maximaler Repeat in t , dann muss es nach Definition ein maximales Paar (i, j, ℓ) geben mit $t_i \cdots t_{i+\ell-1} = t_j \cdots t_{j+\ell-1} = s$, $t'_{i-1} \neq t'_{j-1}$ und $t'_{i+\ell} \neq t'_{j+\ell}$. Somit ist s ein rechtsverzweigendes Teilwort von $\phi t \$$ und $\bar{s}$ muss dann ein interner Knoten sein. Da ja $t'_{i-1} \neq t'_{j-1}$ ist, ist das Linkszeichen vom Blatt $\overline{t'_i \cdots t'_{n+1}}$ ungleich dem Linkszeichen von Blatt $\overline{t'_j \cdots t'_{n+1}}$. Da weiter $t'_{i+\ell} \neq t'_{j+\ell}$ ist, tauchen die beiden unterschiedlichen Linkszeichen in zwei verschiedenen Teilbäumen der Kinder von $\bar{s}$ auf und $\bar{s}$ ist somit linksdivers. Dies ist auch in Abbildung 3.6 illustriert ■

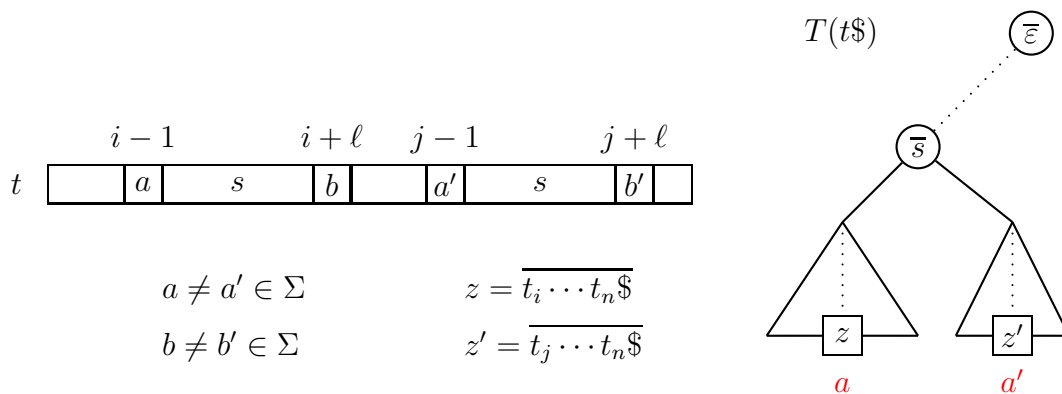


Abbildung 3.6: Skizze: Ein maximales Repeat, das durch sein zugehöriges maximales Paar einen internen Knoten im zugehörigen Suffix-Baum induziert

Damit erhalten auch eine kompakte Darstellung von maximalen Repeats in t . Wir müssen nur den Suffix-Baum für $\#t\$$ konstruieren und dort alle nicht linksdiversen Knoten (inklusive der Blätter) und die dazu inzidenten Kanten entfernen. Man sieht leicht, dass der restliche Baum zusammenhängend ist und nach dem vorherigen Satz alle maximalen Repeats charakterisiert.

3.1.3 Erkennung maximaler Repeats

Wir müssen nun nur noch einen Algorithmus zur Erkennung linksdiverser Knoten entwickeln.

Wir werden das Problem auch hier wieder mit einer Tiefensuche durch den Suffix-Baum $T = T(\#t\$)$ erledigen. Dazu bestimmen wir an den Blättern das jeweilige Linkszeichen. Hierfür müssen wir für ein Blatt nur wissen, ab welcher Position i das zugehörige Suffix beginnt, dann ist t'_{i-1} das zugehörige Linkszeichen.

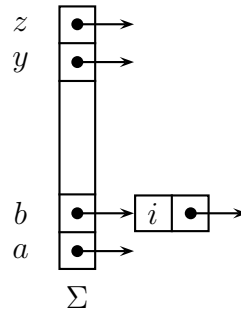
An den inneren Knoten bekommen wir von den Kindern entweder die Information zurück, ob das Kind linksdivers ist, oder das Zeichen das an allen Blättern des Teilbaums des Kindes als Linkszeichen notiert ist. Liefern alle Kinder dasselbe Zeichen zurück, so ist der Knoten nicht linksdivers und wir geben dieses Zeichen an seinen Elter zurück. Sind die von den Kindern zurückgelieferten Zeichen unterschiedlich oder ist eines der Kinder linksdivers, so geben wir an den Elter die Information linksdivers zurück.

Die Laufzeit an den Blättern ist konstant und an den inneren Knoten proportional zum Grad des Knotens. Da die Summe der Anzahl der Kinder aller Knoten eines Baumes mit n Knoten gerade $n - 1$ ist, ist die Gesamtlaufzeit zur Bestimmung linksdiverser Knoten linear in der Größe des Suffix-Baumes T und somit $O(|t|)$.

Jetzt müssen wir nur alle maximalen Paare (Position und Länge des Repeats in der Zeichenreihe) generieren. Hierfür konstruieren wir für jeden Knoten $\bar{v}$ ein Feld $L_{\bar{v}}[\cdot]$ von Listen, wobei das Feld über die Zeichen des Alphabets Σ indiziert ist und die Listenelemente Positionen innerhalb von t (also Werte aus $[1 : |t|]$) sind. Intuitiv bedeutet dies für einen Knoten $\bar{v}$, dass in seiner Liste für das Zeichen $a \in \Sigma$ die Positionen gespeichert sind, für die a ein Linkszeichen im zugehörigen Teilbaum ist.

An den Blättern erzeugen wir das Feld von Listen wie folgt. Sei b das Linkszeichen des Blattes, dessen zugehöriges Suffix an Position i beginnt. Dann enthalten alle Feldelemente die leere Liste mit Ausnahme für das Zeichen $b \in \Sigma$, die eine einelementige Liste mit Listenelement i enthält. Dies ist in Abbildung 3.7 illustriert.

An den inneren Knoten erzeugen wir dieses Feld, indem wir für jedes Feldelement $a \in \Sigma$ die entsprechenden Listen der Kinder aneinanderhängen. Bevor wir für ein

Abbildung 3.7: Skizze: Das Feld der Listen $L_{\bar{v}}$ für ein Blatt $\bar{v}$ mit Linkszeichen $b \in \Sigma$

Kind die Listen an die Listen des Knotens $\bar{v}$ anhängen, geben wir erst noch die maximalen Paare wie folgt aus: Für alle $a \neq b \in \Sigma$ und für jedes Paar von Kindern $\bar{v}' \neq \bar{v}''$ von $\bar{v}$ erzeuge für alle $i \in L_{\bar{v}'}[a]$ und $j \in L_{\bar{v}''}[b]$ die Paare $(i, j, |v|)$. Wenn wir das Kind $\bar{v}''$ neu aufnehmen, dann können wir statt der Liste vom Kind $\bar{v}'$ auch die bislang bereits konstruierte Listen von $\bar{v}$ hernehmen.

Für die Laufzeit ist wichtig, dass die Listen konkateniert und nicht kopiert werden, da die Listen ja sehr lang werden können. Prinzipiell muss dazu im Feld jeweils auch noch ein Zeiger auf das Letzte Feldelement gespeichert werden. Die Idee der Konkatenation der Felderlisten der Kinder ist (ohne die Zeiger auf das das letzte Listenelement) in Abbildung 3.8 illustriert.

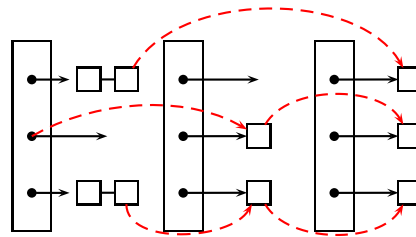


Abbildung 3.8: Skizze: Konkatenation der Felder von Listen zu einem neuen Feld von Listen

Der Aufwand der Konkatenation der Felder von Listen ist wiederum proportional zur Summe der Grade über alle Baumknoten und ist damit wieder linear, d.h. $O(|t|)$. Allerdings kann die Ausgabe der maximalen Paare wiederum mehr Zeit benötigen, jedoch nur konstante Zeit pro maximales Paar. Somit ist die Laufzeit wiederum linear in der Eingabe- und Ausgabegröße.

Theorem 3.11 *Sei $t \in \Sigma^n$, dann lassen sich alle maximalen Paare in Zeit $O(n+k)$ ermitteln, wobei k die Anzahl der maximalen Paare ist.*

Wie im vorherigen Abschnitt lassen sich auch die maximalen Repeats sehr leicht ausgeben. Hierbei ist nur zu beachten, dass die Felder von Listen nicht benötigt werden. Für einen Knoten ist es ausreichend zu wissen, ob in seinem Teilbaum a als Linkszeichen vorkommt oder nicht. Die Details seien dem Leser zur Übung überlassen.

Theorem 3.12 *Sei $t \in \Sigma^n$, dann lassen sich alle maximalen Repeats als Wörter bzw. Referenzen in Zeit $O(n+k)$ ermitteln, wobei k die Anzahl aller Zeichen in den maximalen Repeats bzw. die Anzahl der maximalen Repeats ist.*

Auch hier kann man den Algorithmus leicht so modifizieren, dass er nur exakte Paare ausgibt, deren maximalen Repeats eine Mindestlänge aufweisen.

3.1.4 Revers-komplementäre Repeats

In diesem Abschnitt wollen wir uns mit den so genannten revers-komplementären Repeats beschäftigen, die in Genomen auftreten können, wenn das Repeat eigentlich auf dem anderen Strang der DNA-Doppelhelix liegt. Dazu müssen wir erst einmal formal definieren, was wir unter revers-komplementären Repeats verstehen wollen. Wir beginnen mit der Definition von revers-komplementären Sequenzen.

Definition 3.13 *Sei $\Sigma = \{A, C, G, T\}$ ein Alphabet. Für $\sigma \in \Sigma$ ist das über π zugehörige komplementäre Zeichen $\tilde{\sigma} = \pi(\sigma)$, wobei π eine Permutation auf Σ mit $\pi^2(\sigma) = \sigma$ ist.*

Für $w \in \Sigma^$ ist das über π zugehörige revers-komplementäre Wort $\tilde{w}$ wie folgt definiert:*

$$\tilde{w} = \begin{cases} \varepsilon & \text{falls } w = \varepsilon, \\ \tilde{v} \cdot \tilde{\sigma} & \text{falls } w = \sigma \cdot v \text{ mit } \sigma \in \Sigma, v \in \Sigma^*. \end{cases}$$

Für das biologisch relevante Alphabet $\Sigma = \{A, C, G, T\}$ ist das komplementäre Zeichen wie folgt definiert:

$$\tilde{\sigma} = \pi(\sigma) := \begin{cases} A & \text{falls } \sigma = T, \\ C & \text{falls } \sigma = G, \\ G & \text{falls } \sigma = C, \\ T & \text{falls } \sigma = A. \end{cases}$$

Hierauf basierend können wir nun revers-komplementäre Repeats definieren.

Definition 3.14 *Sei $t \in \Sigma^*$. Ein Teilwort s von t heißt revers-komplementäres Repeat, wenn es ein Tripel (i, j, ℓ) mit $t_i \cdots t_{i+\ell-1} = s$ und $t_j \cdots t_{j+\ell-1} = \tilde{s}$ gibt.*

Die Lösung zum Auffinden exakter bzw. maximaler revers-komplementärer Repeats ist nun einfach, da wir das Problem auf exakte bzw. maximale Repeats zurückführen können. Sei $t \in \Sigma^*$ die Sequenz in der wir revers-komplementäre Repeats ermitteln wollen. Wir konstruieren zuerst das Wort $t' = \#t\#\tilde{t}\#$, wobei $\#$, $\tilde{t} \notin \Sigma$ neue Zeichen sind. Ein revers-komplementäres Repeat in t entspricht nun einem normalen Repeat in t' . Dies ist in Abbildung 3.9 schematisch dargestellt.

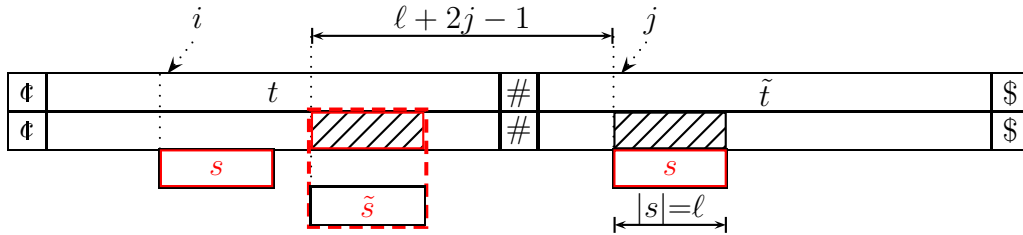


Abbildung 3.9: Skizze: revers-komplementäre Repeats

Wir müssen dabei nur beachten, dass wir auch normale Repeats in t finden (bzw. dann natürlich und nur dann auch in $\tilde{t}$). Wir müssen bei der Ermittlung der Repeats also darauf achten, dass ein Paar (i, j, ℓ) nur dann interessant ist, wenn $i \in [1 : n]$ und $j \in [n + 2 : 2n + 1]$ ist (wobei wir annehmen, dass $t' = t'_0 \cdots t'_{2n+2}$ ist). Als Ausgabe generieren wir dann $(i, j - (\ell + 2j -), \ell)$.

Um nicht mit der Überprüfung von Repeats innerhalb von t (bzw. $\tilde{t}$) aufgehalten zu werden, führen wir in den entsprechenden Algorithmen immer zwei Listen mit. Eine mit den Einträge von Positionen aus $[1 : n]$ und eine mit den Einträgen von Positionen aus $[n + 2 : 2n + 1]$. Bei der Ausgabe müssen wir dann immer nur Paare betrachten, bei denen die Positionen aus den jeweils entgegengesetzten Listen stammen.

Theorem 3.15 *Sei $t \in \Sigma^*$, dann lassen sich alle Paare, die exakte bzw. maximale revers-komplementäre Repeats darstellen, in Zeit $O(n + k)$ ermitteln, wobei k die Anzahl der ausgegeben Paare ist.*

21.11.06

3.2 Tandem-Repeats

In diesem Abschnitt wollen wir uns mit so genannten Tandem-Repeats beschäftigen, das sind kurz gesprochen exakte Repeats, die unmittelbar hintereinander in t vorkommen.

3.2.1 Was sind Tandem-Repeats

Zunächst einmal definieren wir formal, was Tandem-Repeats sind.

Definition 3.16 Für $t \in \Sigma^*$ heißt ein Paar (i, ℓ) Tandem-Repeat-Paar in t , wenn $t_i \cdots t_{i+\ell-1} = t_{i+\ell} \cdots t_{i+2\ell-1}$ gilt. Mit $\mathcal{T}(t)$ bezeichnen wir die Menge aller Tandem-Repeat-Paare von t . Das Wort $t_i \cdots t_{i+2\ell-1}$ heißt dann auch Tandem-Repeat. Mit $\overline{\mathcal{T}}(t)$ bezeichnen wir die Menge aller Tandem-Repeats von t . Die Länge eines Tandem-Repeat-Paares (i, ℓ) bzw. eines Tandem-Repeats $\alpha\alpha$ ist 2ℓ bzw. $2|\alpha|$.

In Abbildung 3.10 ist schematisch ein Tandem-Repeat dargestellt.

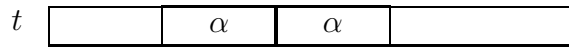


Abbildung 3.10: Skizze: Tandem-Repeat

Wir reden auch manchmal etwas locker von einem Tandem-Repeat anstelle eines Tandem-Repeat-Paares, wenn klar ist, welches zugehörige Tandem-Repeat-Paar hier gemeint ist.

Definition 3.17 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei $t' = t'_0 \cdots t'_{n+1} = \#t_1 \cdots t_n\$$. Ein Tandem-Repeat-Paar (i, ℓ) von t heißt rechtsverzweigend bzw. linksverzweigend, wenn $t'_{i+\ell} \neq t'_{i+2\ell}$ bzw. $t'_{i-1} \neq t'_{i+\ell-1}$ gilt.

Wir sagen auch manchmal etwas locker, dass ein Tandem-Repeat rechtsverzweigend ist, wenn klar ist, welches das zugehörige Tandem-Repeat-Paar ist. Man beachte, dass es Tandem-Repeats geben kann, für die es sowohl ein Tandem-Repeat-Paar geben kann, das rechtsverzweigend ist, als auch ein anderes, das nicht rechtsverzweigend ist. In Abbildung 3.11 ist schematisch ein rechtsverzweigendes Tandem-Repeat dargestellt.

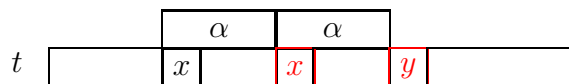


Abbildung 3.11: Skizze: Ein rechtsverzweigendes Tandem-Repeat mit $x \neq y \in \Sigma$

Definition 3.18 Sei $t \in \Sigma^*$ und sei (i, ℓ) ein Tandem-Repeat-Paar in t . Ist auch $(i+1, \ell)$ bzw. $(i-1, \ell)$ ein Tandem-Repeat-Paar, so nennen wir das Tandem-Repeat-Paar (i, ℓ) eine Linksrotation von $(i+1, \ell)$ bzw. Rechtsrotation von $(i-1, \ell)$.

Zunächst einmal wiederholen wir die Definition des im Folgenden wichtigen Begriffs der Worttiefe eines Knotens im Suffix-Baum.

Definition 3.21 Sei $t \in \Sigma^*$ und sei $T = T(t\$)$ der Suffix-Baum für $t\$$. Für einen Knoten $v \in V(T)$ definieren wir seine Worttiefe als $|path(v)|$.

Für das Folgende ist auch der Begriff der Blattlisten (engl. leaf lists) von fundamentaler Bedeutung.

Definition 3.22 Sei $t \in \Sigma^*$ und sei $T = T(t\$)$ der Suffix-Baum zu $t\$$. Für ein Blatt $\bar{v} \in V(T)$ ist die Blattliste $LL(\bar{v})$ als die Menge $\{|t\$| - |v| + 1\}$ definiert (d.h. die Indexposition, an der der zugehörige Suffix in t auftritt). Die Blattliste $LL(\bar{v})$ eines inneren Knotens $\bar{v} \in V(T)$ ist definiert durch

$$LL(\bar{v}) := \bigcup_{\substack{\bar{v}' \in V(T) \\ (\bar{v}, \bar{v}') \in E(T)}} LL(\bar{v}')$$

(d.h. die Menge aller Indexposition von Suffixen, deren zugehörigen Blätter sich im Teilbaum von $\bar{v}$ befinden).

Mit Hilfe dieser Begriffe können wir nun eine Charakterisierung von Tandem-Repeats durch Knoten im zugehörigen Suffix-Baum angeben.

Lemma 3.23 Sei $t \in \Sigma^n$ und sei $i < j \in [1 : n]$ sowie $\ell := j - i$. Dann sind folgende Aussagen äquivalent:

1. (i, ℓ) ist ein Tandem-Repeat-Paar in t .
2. Die Positionen i und j erscheinen in der Blattliste eines Knotens $v \in V(T(t\$))$ mit Worttiefe mindestens ℓ .

Beweis: 1 $\Rightarrow$ 2 : Nach Voraussetzung ist also (i, ℓ) ein Tandem-Repeat-Paar in t . Dies bedeutet, dass $t_i \cdots t_{i+\ell-1} = t_{i+\ell} \cdots t_{i+2\ell-1}$ gilt. Dies ist in Abbildung 3.13 illustriert.

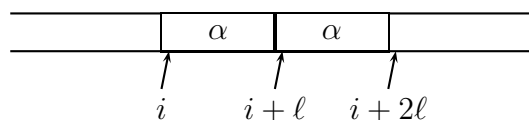


Abbildung 3.13: Skizze: (i, ℓ) ist ein Tandem-Repeat-Paar in t

Sei also $\alpha\alpha$ das zu (i, ℓ) gehörige Tandem-Repeat. Wir unterscheiden jetzt zwei Fälle, je nachdem, ob der Suffix-Baum $T = T(t\$)$ den Knoten $\overline{\alpha}$ enthält oder nicht.

Fall 1: Wir nehmen an, dass $\overline{\alpha}$ existiert. Dann schauen wir uns die Situation im Suffix-Baum T , wie in Abbildung 3.14 angegeben, genauer an. Dann erscheinen i und $j = i + \ell$ in der Blattliste $LL(\overline{\alpha})$. Weiterhin ist offensichtlich die Worttiefe von $\overline{\alpha}$ genau $|\alpha| = \ell$

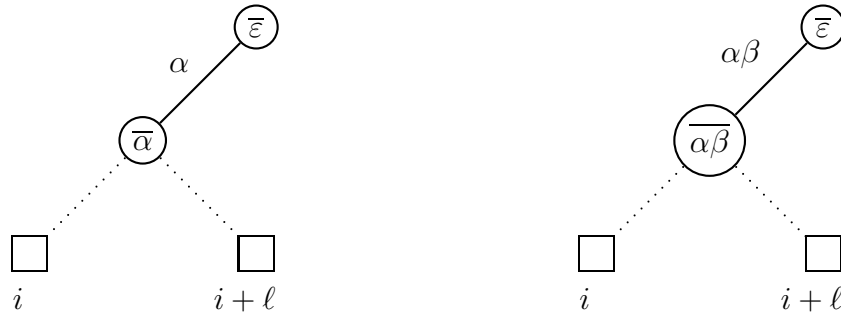


Abbildung 3.14: Skizze: Tandem-Repeat $\alpha\alpha$ und der Suffix-Baum $T(t\$)$

Fall 2: Wir nehmen jetzt an, dass $\overline{\alpha}$ nicht existiert. Sei dazu β so gewählt, dass $\overline{\alpha\beta}$ der Knoten in $T(t\$)$ ist, an dem sich die Pfade zu den Blättern von $t_i \cdots t_n\$$ und $t_{i+\ell} \cdots t_n\$$ trennen. Dann befinden sich jedoch i und $j = i + \ell$ in der Blattliste $LL(\overline{\alpha\beta})$ und die Worttiefe von $\overline{\alpha\beta}$ ist gleich $|\alpha\beta| > \ell$.

2 $\Rightarrow$ 1 : Sei also v ein Knoten von $T(t\$)$ mit Worttiefe mindestens ℓ so gewählt, dass $i, j \in LL(v)$ gilt. Weiterhin sei $\text{path}(v) = \alpha\beta$ mit $|\alpha| = \ell$ und $\beta \in \Sigma^*$. Dieser Fall ist in Abbildung 3.15 illustriert.

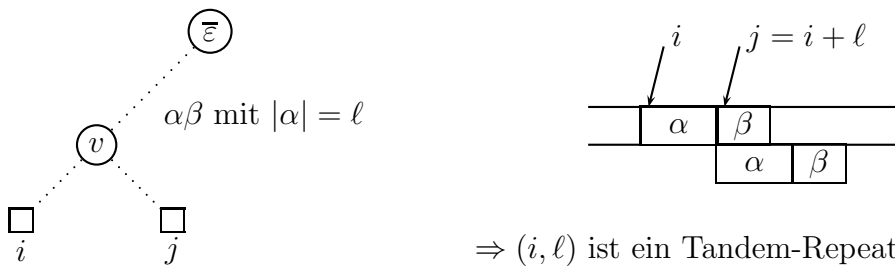


Abbildung 3.15: Skizze: Knoten v mit Worttiefe mindestens ℓ und $i, j \in LL(v)$

Daraus folgt nun sofort, dass $t_i \cdots t_{i+|\alpha\beta|-1} = \alpha\beta$ und $t_j \cdots t_{j+|\alpha\beta|-1} = \alpha\beta$. Da $|\alpha| = \ell$ und $\ell = j - i$ ist, folgt, dass sich die beiden Vorkommen in t überlappen müssen und $\alpha\alpha$ somit ein Tandem-Repeat zum Tandem-Repeat-Paar (i, ℓ) sein muss. Dies ist im

linken Teil der Abbildung 3.15 besonders gut zu erkennen. Damit ist die Behauptung bewiesen. ■

Wir können auch eine analoge Charakterisierung für rechtsverzweigende Tandem-Repeats angeben.

Lemma 3.24 *Sei $t \in \Sigma^n$ sei $i < j \in [1 : n]$ sowie $\ell := j - i > 0$. Dann sind folgende Aussagen äquivalent:*

1. *Das Paar (i, ℓ) ist ein rechtsverzweigendes Tandem-Repeat-Paar.*
2. *Die Positionen i und j erscheinen in der Blattliste eines Knotens $v \in V(T(t\$))$ mit der Worttiefe genau ℓ , aber i und j erscheinen nicht gemeinsam in einer anderen Blattliste eines Knotens mit Worttiefe größer als ℓ .*

Beweis: Der Beweis ist analog zum Beweis in Lemma 3.23 und bleibt dem Leser zur Übung überlassen. ■

3.2.3 Algorithmus von Stoye und Gusfield

Aus dem vorherigen Lemma 3.24 lässt sich sofort der folgende, in Abbildung 3.16 angegebene Algorithmus von Gusfield und Stoye für das Auffinden rechtsverzweigender Tandem-Repeats herleiten. Dabei werden die rechtsverzweigenden Tandem-Repeats explizit durch die Abfrage $t_{i+|\text{path}(v)|} \neq t_{i+2|\text{path}(v)|}$ herausgefiltert. Aufgrund der Tatsache, dass die Anfangsposition eines rechtsverzweigenden Tandem-Repeat-Paares nur an einem Knoten aufgefunden werden können (siehe Lemma 3.24), wird jedes rechtsverzweigende Tandem-Repeat-Paar genau einmal ausgegeben.

```

TANDEMREPEATS (string  $t$ )
{
    tree  $T = \text{SuffixTree}(t\$)$ ;
    for each ( $v \in V(T)$ ) do /* using DFS */
    {
        create  $LL(v) = LL(v_1) \cup \dots \cup LL(v_k)$ ; /* creating leaf lists */
        for each ( $i \in LL(v)$ ) do
            if ( $(i + |\text{path}(v)| \in LL(v)) \ \&\& \ (t_{i+|\text{path}(v)|} \neq t_{i+2|\text{path}(v)|})$ )
                output ( $i, |\text{path}(v)|$ );
    }
}

```

Abbildung 3.16: Algorithmus von Stoye und Gusfield

Der Test, ob $i + |\text{path}(v)| \in LL(v)$ gilt, wird effizient mit Hilfe einer DFS-Nummerierung der Blätter durchgeführt. Haben wir zu jedem Blatt j seine DFS-Nummer $\text{DFS}(j)$ und zu jedem inneren Knoten v das DFS-Intervall $\text{DFS_Int}(v)$, das die DFS-Nummer der Blätter in seinem Teilbaum beschreibt, so gilt:

$$j \in LL(v) \Leftrightarrow \text{DFS}(j) \in \text{DFS_Int}(v)$$

Es ist nämlich offensichtlich effizienter, zwei Intervallgrenzen abzufragen als eine ganze Liste (als Darstellung der Menge) zu durchlaufen. Ein Beispiel ist hierfür in Abbildung 3.17 angegeben

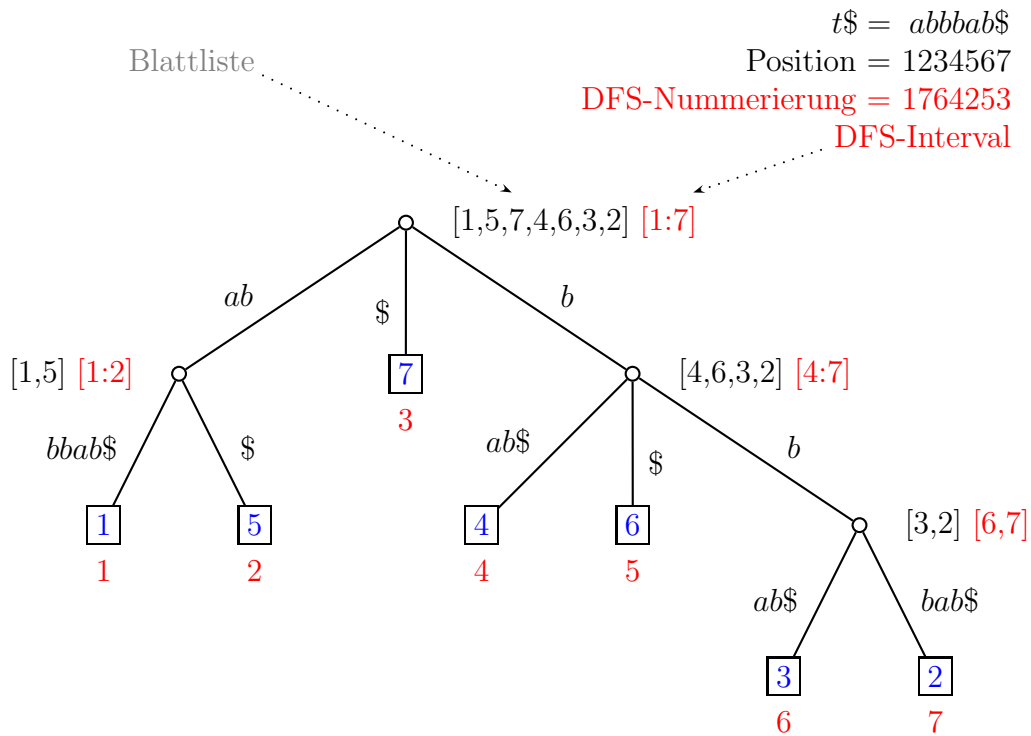


Abbildung 3.17: Beispiel: Auffinden von Tandem-Repeat-Paaren mit dem Algorithmus von Gusfield und Stoye

Wir geben noch ein paar Beispiele zur Bestimmung von $i + |\text{path}|(v) \in LL(v)$ an.

Beispiel 1: Wir betrachten zuerst den Knoten $v = \bar{b}$ und die Abfrage $i = 4 \stackrel{?}{\in} LL(\bar{b})$.

Der ausgeführte Test lautet dann: $j = i + |\text{path}(\bar{b})| = 4 + 1 = 5 \stackrel{?}{\in} LL(\bar{b})$.

Da $\text{DFS}(5) = 2 \notin [4 : 7] = \text{DFS_Int}(\bar{b})$, fällt die Antwort negativ aus.

Beispiel 2: Wir betrachten jetzt den Knoten $v = \bar{b}$ und die Abfrage $i = 2 \stackrel{?}{\in} LL(\bar{b})$.

Der ausgeführte Test lautet dann: $j = i + |\text{path}(\bar{b})| = 2 + 1 = 3 \stackrel{?}{\in} LL(\bar{b})$.

Da $\text{DFS}(3) = 6 \in [4 : 7] = \text{DFS_Int}(\bar{b})$, fällt die Antwort positiv aus.

Da weiter $t_{i+|\text{path}(v)|} = t_{2+1} = b = t_{2+2} = t_{i+2|\text{path}(v)|}$, handelt es sich um keinen rechtsverzweigenden Tandem-Repeat.

Beispiel 3: Wir betrachten erneut den Knoten $v = \bar{b}$ und die Abfrage $i = 3 \stackrel{?}{\in} LL(\bar{b})$.

Der ausgeführte Test lautet dann: $j = i + |\text{path}(\bar{b})| = 3 + 1 = 4 \stackrel{?}{\in} LL(\bar{b})$.
Da $\text{DFS}(4) = 4 \in [4 : 7] = \text{DFS_Int}(\bar{b})$, fällt die Antwort positiv aus.

Da weiter $t_{i+|\text{path}(v)|} = t_{3+1} = b \neq a = t_{3+2} = t_{i+2|\text{path}(v)|}$, handelt es sich um einen rechtsverzweigenden Tandem-Repeat.

3.2.4 Laufzeitanalyse und Beschleunigung

Wir wollen nun die Laufzeit für diesen Algorithmus analysieren. Die Konstruktion des Suffix-Baumes und das Erstellen der Blattlisten (mittels einer Tiefensuche und dem Konkatenieren der Blattlisten der Kinder an den inneren Knoten) benötigt lineare Zeit. Die reine Tiefensuche im zweiten Teil durchläuft $O(n)$ Knoten und für jeden Knoten werden seine Blattlisten durchlaufen. Da die Blattliste maximal $O(n)$ Blätter enthalten kann, beträgt der Zeitaufwand $O(n^2)$. Hierbei haben wir schon ausgenutzt, dass sich der Test $i + |\text{path}(v)| \in LL(v)$ in konstanter Zeit mit Hilfe der DFS-Nummern ausführen lässt.

Theorem 3.25 *Sei $t \in \Sigma^n$. Alle rechtsverzweigenden Tandem-Repeats von t können in Zeit $O(n^2)$ ermittelt werden.*

Berücksichtigen wir die Beobachtung zu Beginn, dass man alle Tandem-Repeats aus rechtsverzweigenden Tandem-Repeats rekonstruieren kann und dass es maximal $O(n^2)$ Tandem-Repeats in einem Wort der Länge n geben kann, erhalten wir das folgende Korollar.

Korollar 3.26 *Sei $t \in \Sigma^n$. Alle Tandem-Repeats von t können in Zeit $O(n^2)$ ermittelt werden.*

Wir wollen jetzt noch einen einfachen Trick vorstellen, mit dem sich die Laufzeit beschleunigen lässt. Wir werden dazu zunächst nicht direkt die Blattliste $LL(v)$ erzeugen, sondern eine disjunkte Vereinigung davon.

Betrachten wir dazu einen Knoten v mit seinen Kinder $v_1, \dots, v_k$. Sei v' ein Kind von v mit einer längsten Blattliste, d.h. $|LL(v')| \geq |LL(v_i)|$ für $i \in [1 : k]$. Dann definieren wir:

$$LL'(v) := \bigcup_{\substack{i=1 \\ v_i \neq v'}}^k LL(v_i).$$

23.11.06

Es gilt dann $LL(v) = LL'(v) \cup LL(v')$ mit $LL'(v) \cap LL(v') = \emptyset$. Dies ist in Abbildung 3.18 illustriert.

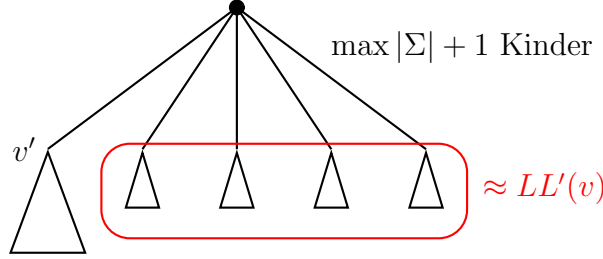


Abbildung 3.18: Skizze: Partitionierung von $LL(v)$ in $LL'(v) \cup LL(v')$

Weiter wissen wir aufgrund von Lemma 3.24 für unseren Test $i + \ell \stackrel{?}{\in} LL(v)$ mit $i \in LL(v)$ und $\ell := |\text{path}(v)|$: Gilt $i \notin LL'(v)$ (d.h. $i \in LL(v')$), dann muss $j = i + \ell \in LL'(v)$ gelten.

Der im Algorithmus durchgeführte Test, ob $i + \ell \stackrel{?}{\in} LL(v)$ für $i \in LL(v)$ und $\ell := |\text{path}(v)|$ wird dann wie folgt ausgeführt.

Gilt $i \in LL'(v)$: In diesem Fall testen wir ganz normal

$$j = i + \ell \stackrel{?}{\in} LL(v) = LL'(v) \cup LL(v').$$

Gilt $i \notin LL'(v)$: Dann gilt $i \in LL(v')$ und es muss $j = i + \ell \in LL'(v)$ gelten. Wir testen daher umgekehrt für $j \in LL(v')$, ob gilt:

$$i = j - \ell \stackrel{?}{\in} LL(v) = LL'(v) \cup LL(v').$$

Dadurch müssen wir in der **for**-Schleife nicht mehr über alle Listenelemente der Blattliste von v gehen, sondern es genügt die Liste $LL'(v)$ zu durchlaufen. Der so modifizierte Algorithmus ist in Abbildung 3.19 angegeben. Auch hier werden in der realen Implementierung die Tests auf Enthaltensein in Listen durch Enthaltensein in DFS-Intervallen ersetzt.

Wie verhalten sich nun die Längen der Blattlisten zueinander? Wir halten zunächst fest, dass der Knoten v maximal $|\Sigma| + 1$ Kinder besitzt und dass die längste Blattliste mindestens so lang sein muss wie eine durchschnittliche. Nach Wahl von v' gilt also:

$$|LL(v')| \geq \frac{|LL(v)|}{|\Sigma| + 1}$$

Mit $|LL(v)| = |LL(v')| + |LL'(v)|$ folgt sofort:

$$|LL'(v)| = |LL(v)| - |LL(v')| \leq |LL(v)| - \frac{|LL(v)|}{|\Sigma| + 1} = \frac{|\Sigma|}{|\Sigma| + 1} \cdot |LL(v)|.$$

```

FASTTANDEMREPEATS (string  $t$ )
{
    tree  $T = \text{SuffixTree}(t\$)$ ;
    for each ( $v \in V(T)$ ) do /* using DFS */
    {
        list  $LL(v) = \emptyset, LL'(v) = \emptyset, LL(v') = \emptyset$ ;
        for each ( $(v, w) \in T$ ) do /*  $w$  is a child of  $v^*$  */
        {
            if ( $|LL(w)| > |LL(v')|$ )
                 $v' = w$ ;
            else
                 $LL'(v) = LL'(v) \cup LL(w)$ ;
        }
        for each ( $i \in LL'(v)$ ) do
        {
            if ( $(i + |\text{path}(v)| \in LL'(v) \cup LL(v')) \ \&\& \ (t_{i+|\text{path}(v)|} \neq t_{i+2|\text{path}(v)|})$ )
                output ( $i, |\text{path}(v)|$ );
            if ( $(i - |\text{path}(v)| \in LL'(v) \cup LL(v')) \ \&\& \ (t_i \neq t_{i+|\text{path}(v)|})$ )
                output ( $i - |\text{path}(v)|, |\text{path}(v)|$ );
        }
         $LL(v) = LL'(v) \cup LL(v')$ ;
    }
}

```

Abbildung 3.19: Beschleunigter Algorithmus von Stoye und Gusfield

Wenn also ein Blatt aus $LL'(v)$ an einem Vergleich beteiligt war, befindet es sich anschließend in einer Blattliste, die mindestens um den Faktor $\frac{|\Sigma|+1}{|\Sigma|}$ größer ist. Dies kann nur

$$\log_{\frac{|\Sigma|+1}{|\Sigma|}}(n) = \frac{\log(n)}{\log\left(1 + \frac{1}{|\Sigma|}\right)} \approx |\Sigma| \cdot \log(n)$$

oft passieren. Somit kann jedes Blatt nur $O(\log(n))$ Vergleiche initiieren. Die Laufzeit aller Tests beträgt also insgesamt: $O(n \log(n))$.

Theorem 3.27 *Alle rechtsverzweigenden Tandem-Repeat-Paare (Tandem-Repeats) eines Wortes $t \in \Sigma^n$ können in Zeit $O(n \log(n))$ ermittelt werden.*

Durch die obere Schranke der Laufzeit erhalten wir auch sofort eine Beschränkung der Anzahl rechtsverzweigender Tandem-Repeats eines Wortes.

Korollar 3.28 *Ein Wort $t \in \Sigma^n$ besitzt maximal $O(n \log(n))$ rechtsverzweigende Tandem-Repeat-Paare (bzw. Tandem-Repeats).*

Korollar 3.29 *Alle Tandem-Repeat-Paare (bzw. Tandem-Repeats) von $t \in \Sigma^n$ können in Zeit $O(n \log(n) + k)$ ermittelt werden, wobei k die Anzahl der Tandem-Repeat-Paare (bzw. Tandem-Repeats) in t ist.*

3.2.5 Algorithmus von Main und Lorentz

In diesem Abschnitt wollen wir noch eine alternative Methode zur Erkennung von Tandem-Repeats angeben. Dieses Verfahren von Main und Lorentz basiert auf einem Divide-and-Conquer-Ansatz. Wir nehmen wieder an, dass $t \in \Sigma^*$ mit $t = t_1 \cdots t_n$. Der Algorithmus geht wie folgt vor.

- 0. Schritt:** Ist $n \leq 1$, dann enthält t keine Tandem-Repeats und wir sind fertig. Andernfalls berechne $h := \lfloor \frac{n}{2} \rfloor$.
- 1. Schritt:** Finde rekursiv alle Tandem-Repeats in $t_1 \cdots t_h$.
- 2. Schritt:** Finde rekursiv alle Tandem-Repeats in $t_{h+1} \cdots t_n$.
- 3. Schritt:** Finde alle Tandem-Repeats $\alpha\alpha$ der Länge 2ℓ mit $\ell \in [1 : \lfloor n/2 \rfloor]$, die an Position $i \in [h - \ell + 2, h]$ beginnen.
- 4. Schritt:** Finde alle Tandem-Repeats $\alpha\alpha$ der Länge 2ℓ mit $\ell \in [1 : \lfloor n/2 \rfloor]$, die an Position $i \in [h - 2\ell + 2, h - \ell + 1]$ beginnen.

Das Prinzip des Algorithmus ist in der folgenden Abbildung 3.20 noch einmal schematisch dargestellt.

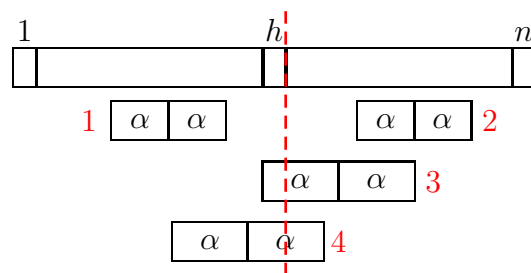


Abbildung 3.20: Skizze: Die vier Schritte des Algorithmus von Main und Lorentz

Die Implementierung der Schritte 0, 1, und 2 sind trivial. Wir müssen uns also nur noch Gedanken um die Implementierung von Schritt 3 und 4 machen. Wie man leicht sieht, sind die Schritte 3 und 4 symmetrisch, so dass wir uns nur einen der beiden Schritte genauer anschauen müssen. Wie wir noch sehen werden, kann

der Schritt 3 bzw. 4 in Zeit $O(n)$ realisiert werden. Damit erhalten wir folgende Rekursionsgleichung für die Laufzeit des Algorithmus von Main und Lorentz:

$$T(n) = \begin{cases} O(1) & \text{falls } n \leq 1, \\ T(\lfloor \frac{n}{2} \rfloor) + T(\lceil \frac{n}{2} \rceil) + O(n) & \text{falls } n \geq 2. \end{cases}$$

Das ist dieselbe Laufzeitanalyse wie für Merge-Sort und die beträgt bekanntlich $O(n \log(n))$. Somit können wir schon das folgende Theorem festhalten.

Theorem 3.30 *Sei $t \in \Sigma^*$ mit $n = |t|$. Mit dem Algorithmus von Main und Lorentz lassen sich alle rechtsverzweigenden Tandem-Repeat-Paare in Zeit $O(n \log(n))$ finden.*

Für den Beweis dieses Satzes müssen wir nur noch zeigen, wie sich Schritt 3 und 4 in linearer Zeit implementieren lassen.

3.2.6 Longest Common Extensions

Für Schritt 3 bzw. 4 im Algorithmus von Main und Lorentz benötigen wir das Konzept der so genannten *longest common extension*.

Definition 3.31 *Sei $t \in \Sigma^*$ mit $n = |t|$ und sei $t' = \#t\$$. Für $i < j \in [1 : n]$ ist eine longest common forward extension bzw. longest common backward extension (oder kurz longest common extension) von i und j in t definiert als eine Zahl $\ell \in \mathbb{N}_0$ mit $t_i \cdots t_{i+\ell-1} = t_j \cdots t_{j+\ell-1}$ und $t'_{i+\ell} \neq t'_{j+\ell}$ (in Zeichen $\text{lce}_f(i, j) = \ell$) bzw. $t_{i-\ell+1} \cdots t_i = t_{j-\ell+1} \cdots t_j$ und $t'_{i-\ell} \neq t'_{j-\ell}$ (in Zeichen $\text{lce}_b(i, j) = \ell$).*

Eine solche longest common (forward) extension ist in der Abbildung 3.21 noch einmal schematisch dargestellt.

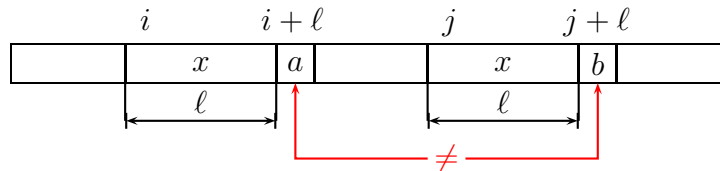


Abbildung 3.21: Skizze: Longest Common Extension

Bevor uns darum kümmern, wie uns longest common extensions für Schritt 3 bzw. 4 im Algorithmus von Main und Lorentz helfen, werden wir uns erst überlegen, wie man solche longest common extensions effizient berechnen kann.

Dazu betrachten wir zunächst den Suffix-Baum $T = T(t\$)$ für t . In T betrachten wir die beiden Blätter, die zu den Suffixen $t_i \cdots t_n \$$ und $t_j \cdots t_n \$$ gehören. Zu diesen muss es nach Konstruktion des Suffix-Baumes T jeweils einen Pfad von der Wurzel zum entsprechenden Blatt geben. Solange diese beiden Pfade von der Wurzel gemeinsam verlaufen, verfolgen wir ein gemeinsames Wort, dass sowohl ab Position i als auch ab Position j in t beginnt (common extension). Sobald sich die Pfade trennen, unter-

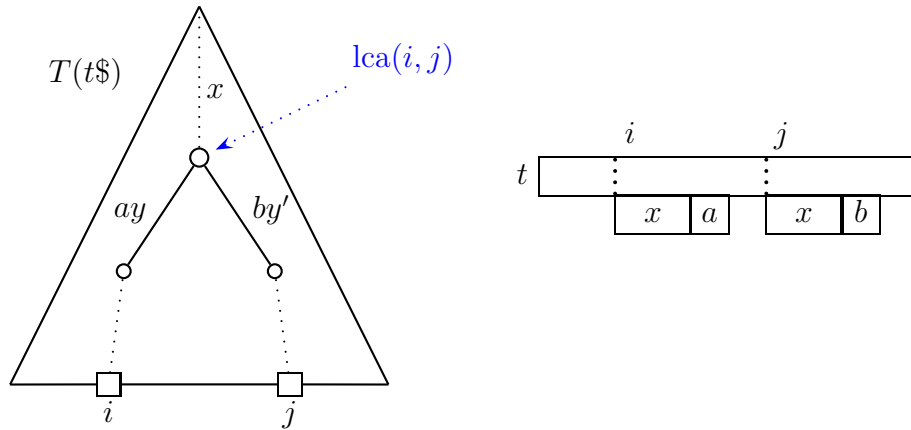


Abbildung 3.22: Skizze: Eine longest common (forward) extension in t und der niedrigste gemeinsame Vorfahre im Suffix-Baum von $t\$$

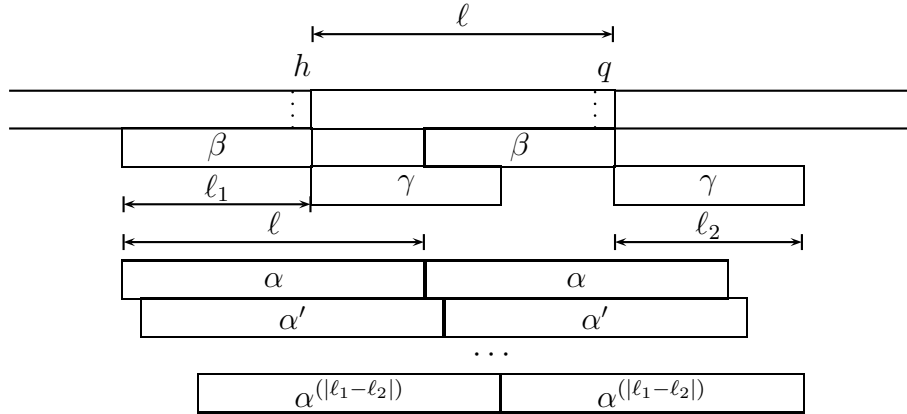
scheiden sich der folgende Buchstabe und wir haben an diesem Knoten $v \in V(T)$ die longest common extension $\text{lce}_f(i, j) = \ell = |\text{path}(v)| = |\text{path}(\text{lca}(i, j))|$ gefunden. Dies ist in Abbildung 3.22 noch einmal illustriert.

Wir können also das Problem der longest common extension im Wort t auf eine Anfrage zum niedrigsten gemeinsamen Vorfahren (eng. lowest common ancestor, kurz lca) zweier Blätter im Suffix-Baum für $t\$$ reduzieren.

Wir gehen in Kapitel 4 näher darauf ein, wie man solche lowest common ancestor Anfragen effizient beantworten kann, insbesondere wenn für einen Baum sehr viele Anfragen gestellt werden. Für die weitere Analyse des Algorithmus von Main und Lorentz geben wir hier noch kurz das hierfür interessierende Ergebnis an. Mit einer Vorverarbeitung des Baumes T mit n Knoten in Zeit von $O(n)$, kann jede nachfolgende LCA-Query in konstanter Zeit beantwortet werden.

Für die Lösung von Schritt 3 (bzw. analog für Schritt 4) gehen wir wie folgt vor. Wir setzen $h = \lfloor n/2 \rfloor$ und $q = h + \ell$ wobei $\ell \in [1 : \lfloor n/2 \rfloor]$ die Länge einer Hälfte eines Tandem-Repeats ist. Dann ermitteln wir die longest common backward extension in t von den Positionen h und q :

$$\ell_1 := \text{lce}_b(h, q).$$



$$\ell_1 := \text{lce}_b(h, q) \quad \ell_2 := \text{lce}_f(h + 1, q + 1)$$

Abbildung 3.23: Skizze: Longest common extensions in t ab Position $h = \lfloor n/2 \rfloor$ und $q = h + \ell$ und die induzierten Tandem-Repeats

Weiterhin ermitteln wir die longest common forward extension in t von den Positionen $h + 1$ und $q + 1$:

$$\ell_2 := \text{lce}_f(h + 1, q + 1).$$

Gilt für die longest common extensions $\ell_1 + \ell_2 \geq \ell$, dann überlappen sich die longest common extensions $\beta = t_{h-\ell_1+1} \cdots t_h$ und $\gamma = t_{h+1} \cdots t_{h+\ell_2}$ (bzw. sind unmittelbar aufeinander folgend) im Bereich $t_{h+1} \cdots t_q$. Aufgrund dieser Überlappung beginnen ab Position $h - \ell_1 + 1$ genau $|\ell_1 - \ell_2| + 1$ Tandem-Repeats. Gilt weiter, dass $\ell_1 > 0$, dann beginnt dieser Tandem-Repeat im Wort $t_1 \cdots t_h$, wie für Schritt 3 gefordert. Man beachte noch, dass all diese Tandem-Repeats bis auf das letzte nichtrechtsverzweigend sind. Falls $\ell_1 + \ell_2 < \ell$ gilt, dann gibt es offensichtlich kein Tandem-Repeat, der im Teilwort $t_1 \cdots t_h$ beginnt und in das Teilwort $t_{h+1} \cdots t_n$ hineinragt. Dies ist in Abbildung 3.23 noch einmal illustriert.

Wir kümmern uns nun noch um die Laufzeit von Schritt 3 (bzw. Schritt 4). Wie noch zu zeigen ist, können wir nach einer linearen Vorverarbeitung jede longest common extension Anfrage in konstanter Zeit bearbeiten. Für ein festes ℓ können somit die zwei LCE-Anfragen in konstanter Zeit beantwortet werden. Die Ausgabe eines rechtsverzweigenden Tandem-Repeat-Paares geht dann ebenfalls in konstanter Zeit, da es maximal eines geben kann. Für alle $\ell \in [1 : \lfloor n/2 \rfloor]$ ist somit Schritt 3 (bzw. Schritt 4) in Zeit $O(n)$ lösbar.

Theorem 3.32 *Der Algorithmus von Main und Lorentz kann alle rechtsverzweigenden Tandem-Repeat-Paare (bzw. alle rechtsverzweigenden Tandem-Repeats) in Zeit $O(n \log(n))$ und Platz $O(n)$ in einem Wort $t \in \Sigma^n$ finden.*

Will man alle Tandem-Repeat-Paare ausgeben, so erhöht sich die Laufzeit um die Anzahl der ausgegeben Tandem-Repeat-Paare, da diese ja nach Beobachtung 3.19 Linksrotationen der rechtsverzweigenden Tandem-Repeat-Paare sind.

Somit beträgt die Laufzeit des Algorithmus von Main und Lorentz $O(n \log(n)) + k$, wobei k die Anzahl der Tandem-Repeats in t ist. Da es nach Korollar maximal $O(n \log(n))$ Tandem-Repeat-Paare geben kann erhalten wir den folgenden Satz.

Theorem 3.33 *Der Algorithmus von Main und Lorentz kann alle Tandem-Repeat-Paare (bzw. alle Tandem-Repeats) in $t \in \Sigma^n$ in Zeit $O(n \log(n))$ und Platz $O(n)$ finden.*

Wir merken noch Folgendes ohne Beweis an: Die Anzahl primitiver Tandem-Repeats ist ebenfalls durch $O(n \log(n))$ beschränkt. Die Anzahl so genannter maximaler (verzweigender) Tandem-Repeats ist jedoch $O(n)$ beschränkt. Wir verweisen hierfür auf die Originalliteratur von R. Kolpakov und G. Kucherov.

28.11.06

3.2.7 Algorithmus von Landau und Schmidt

In diesem Abschnitt wollen wir einen Algorithmus zur Erkennung von Tandem-Repeats mit Fehlern angeben. Dazu zunächst einmal die Definition, welche Fehler wir tolerieren wollen.

Definition 3.34 *Sei $t = t_1 \cdots t_n \in \Sigma^*$. Ein k -mismatch Tandem-Repeat der Länge 2ℓ ist ein Paar (i, ℓ) , so dass $\delta_H(t_i \cdots t_{i+\ell-1}, t_{i+\ell} \cdots t_{i+2\ell-1}) \leq k$ gilt.*

Hierzu betrachten wir das folgende Beispiel:

$$t = b \text{ } \textcolor{red}{a} \text{ } \textcolor{blue}{b} \text{ } \textcolor{red}{a} \text{ } \textcolor{blue}{a} \text{ } \textcolor{red}{b} \text{ } \textcolor{blue}{a} \text{ } \textcolor{red}{b} \text{ } \textcolor{blue}{a} \text{ } \textcolor{red}{a} \text{ } \textcolor{blue}{b}.$$

Das Wort t enthält ein 2-mismatch Tandem-Repeat der Länge 12 mit einem Repeat an Position 2 mit Länge 6. Dabei gibt es jeweils einen **Mismatch** an den Positionspaaren (4, 10) und (7, 13).

Wir werden zur Lösung die Idee des Algorithmus von Main und Lorentz wiederverwenden. Wir werden also wieder einen Divide-and-Conquer-Ansatz benutzen. Auch hier müssen wir uns nur um Schritt 3 (bzw. 4) Gedanken machen. Sei also wiederum $t = t_1 \cdots t_n$ und sei $h = \lfloor n/2 \rfloor$ und $q = h + \ell$ für ein festes $\ell \in [1 : \lfloor n/2 \rfloor]$.

Anstatt jetzt an Position $h-1$ und $q-1$ nur eine longest common backward extension zu finden, iterieren wir dieses Verfahren jetzt $k+1$ mal. Dazu definieren wir

$$\begin{aligned}\ell_i &:= \text{lce}_b(h - L_{i-1}, q - L_{i-1}) \quad \text{für } i \in [1 : k+1], \\ L_i &:= \sum_{j=1}^i (\ell_j + 1) \quad \text{für } i \in [0 : k+1].\end{aligned}$$

Beachte, dass $L_i \geq i$ gilt. Anschaulich ist L_i die Anzahl der Buchstaben, die ab Position h bzw. q nach links gelesen werden muss, bis der i -te Mismatch auftritt. Somit erhalten wir mit

$$(t_{h-L_{k+1}+2} \cdots t_h, t_{q-L_{k+1}+2} \cdots t_q)$$

quasi eine *longest common backward extension with k mismatches*.

Analog führen wir ab den Positionen $h+1$ und $q+1$ ebenfalls $k+1$ longest common forward extension aus. Dazu definieren wir

$$\begin{aligned}\ell'_i &:= \text{lce}_f(h+1 + L'_{i-1}, q+1 + L'_{i-1}) \quad \text{für } i \in [1 : k+1], \\ L'_i &:= \sum_{j=1}^i (\ell'_j + 1) \quad \text{für } i \in [0 : k+1].\end{aligned}$$

Beachte, dass $L'_i \geq i$ gilt. Anschaulich ist L'_i die Anzahl Buchstaben, die ab Position h bzw. q nach rechts gelesen werden muss, bis der i -te Mismatch auftritt. Somit erhalten wir mit

$$(t_{h+1} \cdots t_{h+L'_{k+1}-1}, t_{q+1} \cdots t_{q+L'_{k+1}-1})$$

quasi eine *longest common forward extension with k mismatches*. Dieser Ansatz zur Lösung von Schritt 3 ist in Abbildung 3.24 noch einmal illustriert.

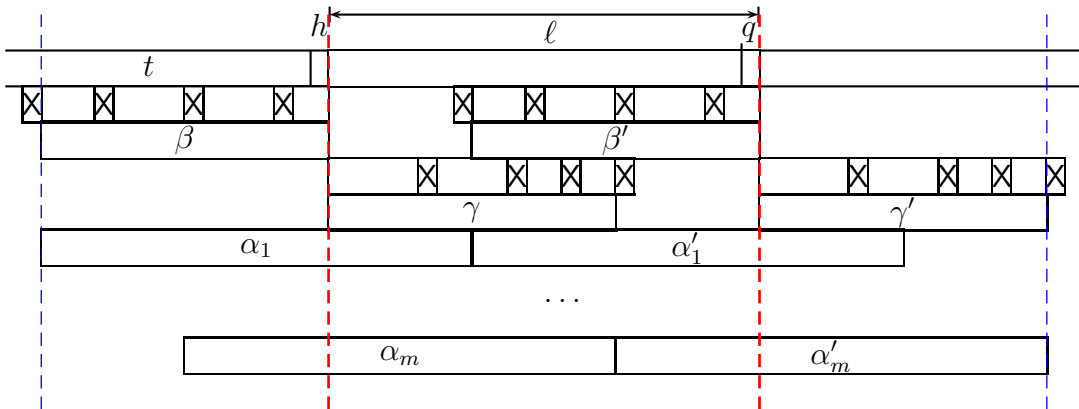


Abbildung 3.24: Skizze: Anwendung von je $k+1$ lce zur Tolerierung $\leq k$ Mismatches

Nur im Intervall $[h - L_{k+1} + 2 : q + L'_{k+1} - 1]$ kann jetzt ein k -mismatch Tandem-Repeat der Länge 2ℓ auftreten. Damit dieses k -mismatch Tandem-Repeat die Positionen h und $h + 1$ beinhaltet (wie im Conquer-Schritt gefordert), muss es sogar im Intervall $[A : B]$ enthalten sein, wobei $A := \max\{h - L_{k+1} + 2, h - 2\ell + 2\}$ und $B := \min\{q + L'_{k+1} - 1, h + 2\ell - 1\}$ ist.

Im Gegensatz zum Algorithmus von Main und Lorentz muss jedoch nicht jede in diesem Intervall $[A : B]$ enthaltene Zeichenreihe der Länge 2ℓ ein k -mismatch Tandem-Repeats sein, da sich im Inneren die Mismatches häufen können.

Für $i \in [A : B - 2\ell + 1]$ betrachten wir jetzt für jedes Wort $t_i \cdots t_{i+2\ell-1}$ der Länge 2ℓ die Anzahl $m(i)$ der Mismatches im zugehörigen Tandem-Repeat, d.h.:

$$m(i) := |\{j \in [1 : \ell] : t_{i+j-1} \neq t_{i+\ell+j-1}\}|.$$

Für $m(i)$ können wir auf jeden Fall festhalten, dass Folgendes gilt:

$$|m(i) - m(i + 1)| \leq 1 \quad \text{für } i \in [A : B - 2\ell + 1].$$

Die Funktion $m(i)$ kann sich auch nur an den Positionen ändern, an denen wir ein explizites Mismatch feststellen. Diese Funktion ist in Abbildung 3.25 illustriert.

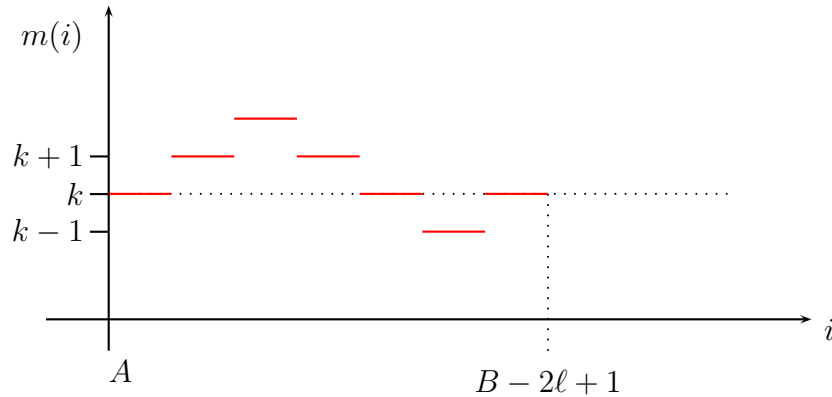


Abbildung 3.25: Skizze: Diagramm der Anzahl Mismatches $m(i)$ bei Teilwörtern der Länge 2ℓ im Kandidaten-Intervall von k -mismatch Tandem-Repeats

Algorithmisch benötigen wir nur die folgenden Listen von Positionen von Mismatches $L = \{h - L_k + 1, \dots, h - L_1 + 1\}$ und $L' = \{h + L'_1, \dots, h + L'_k\}$, die jeweils aufsteigend sortiert sind. Sei M die sortierte Vereinigung dieser beiden Listen. Nur an diesen Positionen kann sich $m(i)$ ändern. Als Tandem-Repeats bzw. Tandem-Repeat-Paar müssen wir also nur für $m \in M$ die Anzahl Mismatches feststellen.

Für das kleinste $m \in M$ können wir dies in linearer Zeit in ℓ tun. Für die folgenden $m \in M$ müssen wir immer nur testen, ob wir vorne einen Mismatch verlieren oder

hinten einen neuen erhalten. Somit ist die Laufzeit für Schritt 3 für ein festes ℓ gerade $O(k)$. Der Conquer-Schritt benötigt also insgesamt Laufzeit $O(kn + z)$, wobei z die Anzahl der ausgegeben k -mismatch Tandem-Repeats bzw. Tandem-Repeat-Paare ist.

Die Gesamtlaufzeit $T(n)$ des Algorithmus (ohne die Ausgabe der Tandem-Repeats) erfüllt dann die Rekursionsgleichung:

$$T(n) = \begin{cases} O(1) & \text{für } n \leq 1, \\ T(\lfloor \frac{n}{2} \rfloor) + T(\lceil \frac{n}{2} \rceil) + O(kn) & \text{für } n \geq 2. \end{cases}$$

Die Lösung dieser Rekursionsgleichung lautet $O(kn \log(n))$. Somit benötigt der Algorithmus inklusive Ausgabe $O(kn \log(n) + z)$, wobei z die Anzahl der k -mismatch Tandem-Repeats ist.

Theorem 3.35 *Sei $t \in \Sigma^*$. Alle k -mismatch Tandem-Repeats von t können mit dem Algorithmus von Landau und Schmidt in Zeit $O(kn \log(n) + z)$ gefunden werden, wobei z die Anzahl der k -mismatch Tandem-Repeats angibt.*

3.3 Vokabulare von Tandem-Repeats

In diesem Abschnitt wollen wir noch einen echten Linearzeit-Algorithmus zum Auffinden aller Tandem-Repeats vorstellen, wobei wir als Ausgabe einen mit Tandem-Repeats dekorierten Suffix-Baum generieren werden.

3.3.1 Vokabulare und Überdeckungen

Wir werden zunächst ein paar grundlegenden Definitionen angeben, um dann die Idee der Darstellung der Tandem-Repeats und das Prinzip des Algorithmus erläutern zu können.

Definition 3.36 *Sei $t = t_1 \cdots t_n \in \Sigma^*$. Zwei Tandem-Repeat-Paare (i, ℓ) und (i', ℓ') von t sind vom gleichen Typ, wenn $\ell = \ell'$ und $t_i \cdots t_{i+\ell-1} = t_{i'} \cdots t_{i'+\ell-1}$ gilt.*

Anschaulich sind zwei Tandem-Repeat-Paare vom gleichen Typ, wenn sie denselben Tandem-Repeat beschreiben.

Wir führen jetzt noch den Begriff des Vokabulars ein, den wir eigentlich schon kennen.

Definition 3.37 *Sei $t \in \Sigma^*$. Die Menge aller Tandem-Repeats $\overline{\mathcal{T}}(t)$ von t bezeichnen wir auch als das Vokabular $\mathcal{V}(t)$ von t .*

Wie wir ja bereits gesehen haben, entstehen normale Tandem-Repeats immer durch eine geeignete Anzahl von Linksrotationen aus einem rechtsverzweigenden Tandem-Repeat. Solchen konsekutiven Tandem-Repeats wollen wir zur besseren Orientierung noch einen Namen geben.

Definition 3.38 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Das Intervall $[i : j] \subseteq [1 : n]$ ist ein Run von Tandem-Repeats der Länge 2ℓ von t , wenn $(k, \ell) \in \mathcal{T}(t)$ für alle $k \in [i : j]$.

Mit Hilfe dieser Runs können wir nun auch beschreiben, wann ein Tandem-Repeat durch mehrere Linksrotationen aus einem anderen Tandem-Repeat hervorgeht.

Definition 3.39 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Ein Tandem-Repeat-Paar (i, ℓ) von t überdeckt ein Tandem-Repeat-Paar (j, ℓ) von t , wenn $[i : j]$ ein Run von Tandem-Repeats der Länge 2ℓ ist.

Damit können wir sinnvolle Teilmengen der Menge aller Tandem-Repeat-Paare definieren, die uns wieder alle Tandem-Repeat-Paare reproduzieren können.

Definition 3.40 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Eine Menge $P \subseteq \mathcal{T}(t)$ heißt Überdeckung von $\mathcal{T}(t)$, wenn für jedes $w \in \mathcal{V}(t)$ ein $(i, \ell) \in P$ existiert, so dass (i, ℓ) ein Tandem-Repeat-Paar $(j, \ell) \in \mathcal{T}(t)$ mit $w = t_j \cdots t_{j+2\ell-1}$ überdeckt. Gilt zusätzlich für jedes $(j', \ell) \in \mathcal{T}(t)$ mit $w = t_{j'} \cdots t_{j'+2\ell-1}$ die Beziehung $j' \geq j$, dann heißt P eine linkeste Überdeckung.

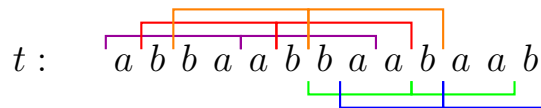


Abbildung 3.26: Beispiel: $t = abbaabbaabaab$

In Abbildung 3.26 sind für die Zeichenreihe $t = abbaabbaabaab$ alle Tandem-Repeats der Länge größer als zwei angegeben. Des Weiteren gilt dort beispielsweise:

- $(1, 4)$ überdeckt $(2, 4)$,
- $(1, 4)$ überdeckt $(3, 4)$,
- $(7, 3)$ überdeckt $(8, 3)$,
- $\mathcal{T}(t) = \{(1, 4), (2, 4), (3, 4), (2, 1), (4, 1), (6, 1), (8, 1), (7, 3), (8, 3), (11, 1)\}$,

- $\mathcal{V}(t) = \overline{\mathcal{T}}(t) = \{abbaabba, aa, baabaa, bbaabbaa, bb, aabaab, baabbaab\},$
- $P = \{(1, 4), (7, 3), (4, 1), (6, 1)\}$ ist eine Überdeckung von $\mathcal{T}(t)$.
- P ist keine linkeste Überdeckung
- $P' = \{(1, 4), (7, 3), (4, 1), (2, 1)\}$ oder $P'' = \{(1, 4), (2, 4), (7, 3), (4, 1), (2, 1)\}$ ist eine linkeste Überdeckung von $\mathcal{T}(t)$.

Man beachte, dass eine linkeste Überdeckung nicht notwendigerweise eine kleinste Überdeckung der Tandem-Repeat-Paare sein muss. Dazu sind in Abbildung 3.27 für das Wort $t = aabaabbaabaabaa$ die Tandem-Repeats der Länge größer als zwei dargestellt. Hier ist $P = \{(1, 1), (6, 1), (8, 3)\}$ eine minimale Überdeckung, die keine linkeste ist. $P' = \{(1, 1), (6, 1), (1, 3), (8, 3)\}$ oder $P'' = \{(1, 1), (6, 1), (1, 3), (9, 3)\}$ sind beispielsweise linkeste Überdeckung.

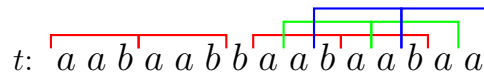


Abbildung 3.27: Beispiel: $t = aabaabbaabaabaa$

Theorem 3.41 *Das Vokabular einer Zeichenreihe der Länge n besitzt maximal $O(n)$ Wörter und kann in einem Suffixbaum repräsentiert werden.*

Der nicht ganz einfache Beweis ist dem Leser als Übungsaufgabe überlassen. Aus dem vorhergehenden Satz folgt sofort das folgenden Korollar.

02.12.06

Korollar 3.42 *Sei $t = t_1 \cdots t_n \in \Sigma^*$. Es gibt eine linkeste Überdeckung der Tandem-Repeat-Paare von t mit Größe $O(n)$.*

Wie schon erwähnt, lässt sich das Vokabular mit Hilfe eines Suffix-Baumes darstellen. Diese Darstellung hat dann offensichtlich eine Platzkomplexität von $O(n)$. In Abbildung 3.28 ist ein Beispiel für die Darstellung eines Vokabular mit Hilfe eines Suffix-Baumes für $t = abbaabbaabaab$ angegeben. Hierbei sind die Worte des Vokabular als Worte im Suffix-Baum dargestellt, die zu Pfaden von der Wurzel bis zum Trennzeichen | bzw. zu einem roten Knoten (hier im Beispiel nur zu einem Blatt) gehen. Offensichtlich wird durch den Suffix-Baum das folgende Vokabular dargestellt:

$$\mathcal{V}(t) = \{aa, aabaab, abbaabba, baabaa, baabbaab, bb, bbaabbaa\}.$$

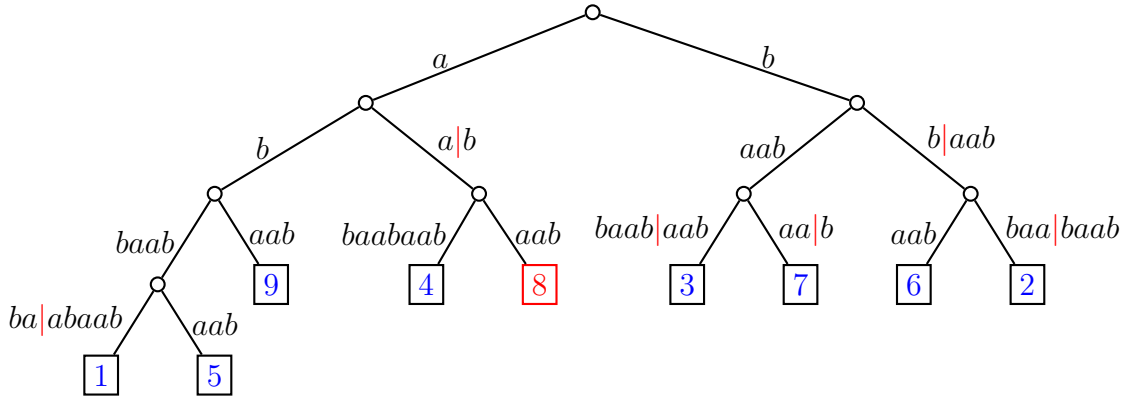


Abbildung 3.28: Beispiel: Der Suffix-Baum zu *abbaabbaabaab* und die Dekorierung seines Vokabulars

3.3.2 Skizze des Algorithmus von Gusfield und Stoye

Der Algorithmus von Gusfield und Stoye zur Dekorierung eines Suffix-Baumes von t mit dem Vokabular von t wird in drei Phasen vorgehen:

Phase I: Zuerst wird eine linkeste Überdeckung P der Tandem-Repeat-Paare $\mathcal{T}(t)$ konstruiert.

Phase II: Dann wird der Suffix-Baum $T(t)$ konstruiert und mit einigen Tandem-Repeats aus der linkesten Überdeckung P dekoriert.

Phase III: Letztendlich wird die vorgefundene Dekorierung zu einer vollständigen Dekorierung von $T(t)$ mit dem Vokabular von t erweitert.

3.3.3 Tandem-Repeats und Lempel-Ziv-Zerlegungen

Bevor wir zur Bestimmung einer linkesten Überdeckung kommen, benötigen wir noch ein Hilfsmittel, die Lempel-Ziv-Zerlegung, sowie einige fundamentale Beziehung von Tandem-Repeats zu diesen.

Notation 3.43 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Für jede Position $i \in [1 : n]$ ist ℓ_i und s_i wie folgt definiert:

$$\begin{aligned} \ell_i &:= \max \{k \in [0 : n - i + 1] : \exists j < i : t_j \cdots t_{j+k-1} = t_i \cdots t_{i+k-1}\}, \\ s_i &:= \min \{j \in [0 : i - 1] : t_j \cdots t_{j+\ell_i-1} = t_i \cdots t_{i+\ell_i-1}\}. \end{aligned}$$

Beachte, dass in den Definitionen von ℓ_i bzw. s_i jeweils $k = 0$ bzw. $j = 0$ in den Mengen, aus denen das Maximum bzw. Minimum gebildet wird, enthalten sind. Somit sind die Extrema wohldefiniert.

Anschaulich ist $t_i \cdots t_{i+\ell_i-1}$ die längste Zeichenreihe, die in t ab Position i in t bereits zum zweiten Mal auftritt; s_i gibt dabei an, ab welcher Position diese Zeichenreihe vorher schon einmal aufgetreten ist und ℓ_i gibt deren Länge an. Dabei ist $s_i = 0$, wenn $\ell_i = 0$ gilt. In Abbildung 3.29 ist diese Notation noch einmal illustriert.

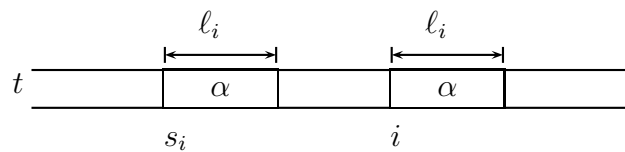


Abbildung 3.29: Skizze: Die Werte ℓ_i und s_i in t

In Abbildung 3.30 sind für das Wort $t = abbaabbaabaab$ noch einmal beispielsweise die Werte für ℓ_i und s_i angegeben.

	a	b	b	a	a	b	b	a	a	b	a	a	b
ℓ_i	0	0	1	1	6	5	4	3	2	4	3	2	1
s_i	0	0	2	1	1	2	3	4	1	3	4	1	2

Abbildung 3.30: Beispiel: Die Werte ℓ_i und s_i für $t = abbaabbaabaab$

Wir halten zunächst die fundamentale Beobachtung fest, dass die Folgenglieder der Folge $(\ell_i)_{i \in \mathbb{N}_0}$ nur um 1 sinken, aber beliebig steigen können.

Beobachtung 3.44 Für jedes $t \in \Sigma^n$ und jedes $i \in [2 : n]$ gilt, dass $\ell_i \geq \ell_{i-1} - 1$ sowie $\ell_i \leq n - i + 1$ und dass diese Schranken scharf sind.

Nun können wir die so genannte Lempel-Ziv-Zerlegung einer Zeichenreihe über Σ definieren.

Definition 3.45 Die Lempel-Ziv-Zerlegung einer Zeichenreihe $t = t_1 \cdots t_n \in \Sigma^*$ ist eine Folge von Zahlen $(i_1, \dots, i_{k+1})$ mit

$$\begin{aligned} i_1 &:= 1, \\ i_{j+1} &:= i_j + \max(1, \ell_{i_j}) \quad \text{für } i_j \leq n, \end{aligned}$$

wobei ℓ_i wie oben definiert ist. Das Wort $t_{i_j} \cdots t_{i_{j+1}-1}$ wird als j -ter Block der Lempel-Ziv-Zerlegung bezeichnet.

Beachte, dass für eine Lempel-Ziv-Zerlegung eines Worte $t \in \Sigma^n$ nach Definition immer $i_1 = 1$, $i_2 = 2$ und $i_{k+1} = n + 1$ gilt. Wir geben als Beispiel für das Wort $abbaabbaabaab$ die Folge $(i_1, \dots, i_7)$ der Lempel-Ziv-Zerlegung an:

$$\begin{aligned} i_1 &= 1 \\ i_2 &= 1 + \max(1, 0) = 2 \\ i_3 &= 2 + \max(1, 0) = 3 \\ i_4 &= 3 + \max(1, 1) = 4 \\ i_5 &= 4 + \max(1, 1) = 5 \\ i_6 &= 5 + \max(1, 6) = 11 \\ i_7 &= 11 + \max(1, 3) = 14 \end{aligned}$$

In Abbildung 3.31 sind für das Beispiel $abbaabbaabaab$ die Lempel-Ziv-Blöcke noch einmal illustriert.

$t =$	<table><tr><td>a</td><td>b</td><td>b</td><td>a</td><td>$a\ b\ b\ a\ a\ b$</td><td>$a\ a\ b$</td></tr></table>	a	b	b	a	$a\ b\ b\ a\ a\ b$	$a\ a\ b$
a	b	b	a	$a\ b\ b\ a\ a\ b$	$a\ a\ b$		

Abbildung 3.31: Beispiel: Lempel-Ziv-Zerlegung von $t = abbaabbaabaab$

Nun können wir einige fundamentale Eigenschaften von Tandem-Repeats in t bezüglich der Lempel-Ziv-Zerlegung von t festhalten.

Lemma 3.46 *Die rechte Hälfte eines Tandem-Repeats überlappt maximal zwei Lempel-Ziv-Blöcke.*

Beweis: Für einen Widerspruchsbeweis nehmen wir an, dass die rechte Hälfte eines Tandem-Repeats mindestens drei Lempel-Ziv-Blöcke überlappt. Sei also (i, ℓ) ein Tandem-Repeat-Paar zum Tandem-Repeat $\alpha\alpha$. Dann gibt es ein $j \in [1 : k - 1]$, so dass $i_j, i_{j+1} \in [i + \ell + 1, i + \ell - 2]$ gilt. Dies ist in Abbildung 3.32 illustriert.

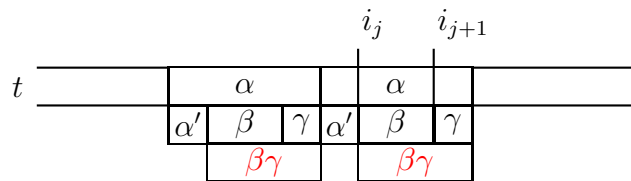


Abbildung 3.32: Skizze: Tandem-Repeats in t und die Lempel-Ziv-Blöcke

Der j -te Block ist dann also β und es gibt $\alpha', \gamma \in \Sigma^+$ mit $\alpha = \alpha'\beta\gamma$. Offensichtlich taucht aber $\beta\gamma$ schon vorher einmal auf, also wäre der j -te Block falsch definiert und wir erhalten den gewünschten Widerspruch. ■

Lemma 3.47 *Das linkeste Vorkommen eines Tandem-Repeats überlappt mindestens zwei Lempel-Ziv-Blöcke.*

Beweis: Für einen Widerspruchsbeweis nehmen wir an, dass es ein linkestes Vorkommen eines Tandem-Repeats gibt, das innerhalb eines Lempel-Ziv-Blocks liegt. Sei $\alpha\alpha$ dieses Tandem-Repeat und sei (i, ℓ) das linkeste Tandem-Repeat-Paar für $\alpha\alpha$ (d.h. mit minimalem i). Sei weiter j so gewählt, dass β der j -te Lempel-Ziv-Block ist, der das Tandem-Repeat-Paar (i, ℓ) enthält. Dies ist auch in Abbildung 3.33 illustriert.

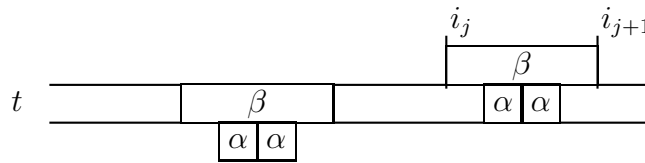


Abbildung 3.33: Skizze: Das linkeste Vorkommen eines Tandem-Repeats in t und die Lempel-Ziv-Zerlegung von t

Nach Definition des j -ten Lempel-Ziv-Blocks muss β ein Teilwort von t sein, dass bereits ab Position $j < i_j$ schon einmal vorkommt. Somit muss auch $\alpha\alpha$ noch einmal in t an einer Position $j' < i$ vorkommen, was offensichtlich ein Widerspruch zu der Tatsache ist, dass (i, ℓ) ein linkestes Tandem-Repeat-Paar für $\alpha\alpha$ ist. ■

Für unser weiteres Vorgehen definieren wir noch das Zentrum eines Tandem-Repeats.

Definition 3.48 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ mit der Lempel-Ziv-Zerlegung $(i_1, \dots, i_{k+1})$ und sei $\alpha\alpha$ ein Tandem-Repeat in t mit $t_i \cdots t_{i+2\ell-1} = \alpha\alpha$ zum Tandem-Repeat-Paar (i, ℓ) . Dann ist das Zentrum von (i, ℓ) im Block B der Lempel-Ziv-Zerlegung, wenn $i + \ell \in [i_B : i_{B+1} - 1]$.*

Aus den beiden vorhergehenden Lemmata ergibt sich nun der folgende Satz.

Theorem 3.49 *Wenn das linkeste Vorkommen eines Tandem-Repeats $\alpha\alpha$ sein Zentrum im Block B hat, dann gilt:*

- Entweder ist das linke Ende von $\alpha\alpha$ im Block B und sein rechtes Ende im Block $B + 1$
- oder das linke Ende von $\alpha\alpha$ liegt links vom Block B .

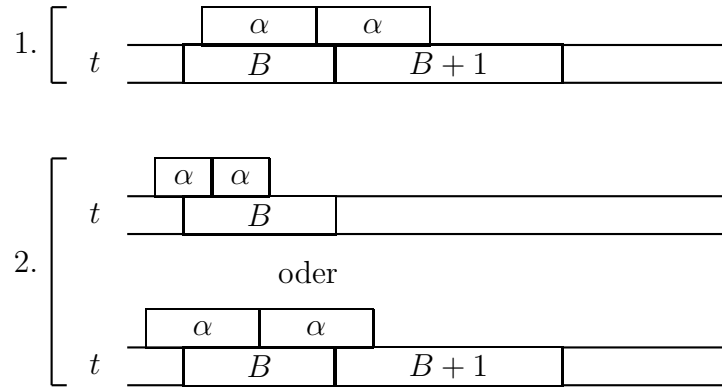


Abbildung 3.34: Skizze: Lage eines Tandem-Repeats $\alpha\alpha$ in t mit Zentrum im Block B zur Lempel-Ziv-Zerlegung von t

Diese verschiedenen Möglichkeiten eines Vorkommens eines Tandem-Repeats ist in der folgenden Abbildung 3.34 noch einmal illustriert.

Wir halten noch als Lemma fest, dass sich für eine Zeichenreihe seine zugehörige Lempel-Ziv-Zerlegung sehr effizient berechnen lässt.

Lemma 3.50 *Sei $t = t_1 \cdots t_n \in \Sigma^*$. Die Lempel-Ziv-Zerlegung von t kann in Zeit $O(n)$ berechnet werden.*

Der Beweis ist dem Leser als Übungsaufgabe überlassen.

3.3.4 Phase I: Bestimmung einer linkesten Überdeckung

In der ersten Phase versuchen wir eine linkeste Überdeckung von $\mathcal{T}(t)$ zu konstruieren. Dazu betrachten wir den Block B der Lempel-Ziv-Zerlegung von t und versuchen alle Tandem-Repeat-Paare, deren Zentrum im Block B liegt und eine der beiden Bedingungen in Theorem 3.49 erfüllt, auszugeben.

Hierzu bezeichne im Folgenden $h := i_B$ bzw. $h' := i_{B+1}$ den Anfang des Blocks B bzw. des Blocks $B+1$ in der Lempel-Ziv-Zerlegung. Dies ist in Abbildung 3.35 illustriert.

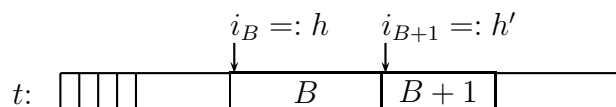


Abbildung 3.35: Skizze: Lempel-Ziv-Blöcke von t

Wir unterscheiden jetzt zwei Fälle, die zu den beiden Charakterisierungen von linken Vorkommen von Tandem-Repeats in Theorem 3.49 korrespondieren.

Fall 1: Wir suchen nur nach linken Vorkommen von Tandem-Repeats mit Zentrum im Block B , deren linkes Ende sich im Block B und dessen rechtes Ende sich im Block $B + 1$ befindet. In diesem Fall gilt für die Halblänge ℓ eines Tandem-Repeats, dass $\ell \in [2 : \ell(B) - 1]$ mit $\ell(B) := i_{B+1} - i_B$. Für $\ell = 1$ kann kein Tandem-Repeat der Länge 2 existieren, dessen Zentrum im Block B und dessen Ende im Block $B + 1$ liegt. Weiterhin kann die maximale Länge nur $\ell(B) - 1$ betragen, da sonst ein im Block B beginnendes Tandem-Repeat sein Zentrum bereits im Block $B + 1$ hätte.

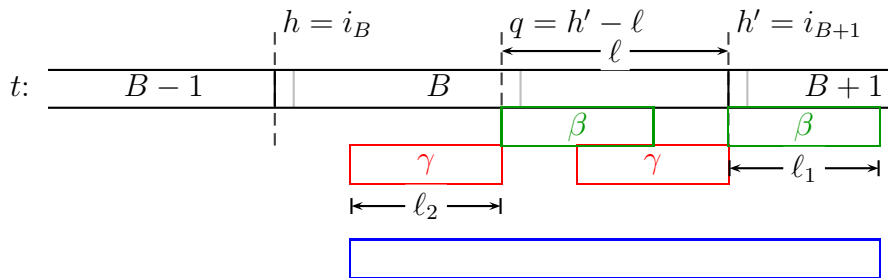


Abbildung 3.36: Skizze: Tandem-Repeats mit Zentrum im Block B , die die erste Bedingung von Theorem 3.49 erfüllen

Dieser Fall ist in Abbildung 3.36 illustriert. Wir setzen $q := h' - \ell$ und bestimmen mit dem Algorithmus für die longest common extensions in konstanter Zeit die folgenden Werte:

$$\begin{aligned}\ell_1 &:= \text{lce}_f(q, h'), \\ \ell_2 &:= \text{lce}_b(q - 1, h' - 1).\end{aligned}$$

Dabei sei β bzw. γ die longest common forward (bzw. backward) extension ab dem Positionspaar (q, h') bzw. $(q - 1, h' - 1)$. Aufgrund von Theorem 3.49 wissen wir, dass ein linkstes Tandem-Repeat das Ende von Block B überlappen muss. Es müssen also die folgenden Bedingungen gelten:

$\ell_1 + \ell_2 \geq \ell$: Auch hier gilt wieder, dass wir nur dann einen Tandem-Repeat der Länge 2ℓ gefunden haben, wenn $\ell_1 + \ell_2 \geq \ell$ gilt. Die Tandem-Repeats der Länge 2ℓ befinden sich dann im blauen Block in Abbildung 3.36.

$\ell_1 > 0$: Gilt weiter $\ell_1 > 0$, dann hat mindestens eines dieser Tandem-Repeats sein rechtes Ende im Block $B + 1$.

$\ell_2 > 0$: Gilt weiter $\ell_2 > 0$, dann gibt es mindestens ein Tandem-Repeat, dessen Zentrum im Block B liegt.

Wir geben dann das Tandem-Repeat-Paar $(\max\{q - \ell_2, q - \ell + 1\}, \ell)$ aus. Hierbei ist $(q - \ell_2, \ell)$ ein Tandem-Repeat-Paar, dass alle gefundenen überdeckt (also von diesem Run das linkeste ist). Das Tandem-Repeat-Paar $(q - \ell + 1, \ell)$ ist das erste, das in den Block $B + 1$ hineinragen kann, da

$$(q - \ell + 1) + 2\ell - 1 = q + \ell = (h' - \ell) + \ell = h'$$

gilt. Man beachte, dass aufgrund von $\ell \in [2 : i_{B+1} - i_B - 1]$ das ausgegebene Tandem-Repeat-Paar sein Zentrum im Block B hat, da

$$\begin{aligned} (q - \ell + 1) + \ell = q + 1 = h' - \ell + 1 &\in [i_{B+1} - (i_{B+1} - i_B - 1) + 1 : h' - 2 + 1] \\ &= [i_B + 2 : i_{B+1} - 1]. \end{aligned}$$

Der Algorithmus ist in Abbildung 3.37 noch einmal angegeben.

FINDLEFTMOSTTRCOND1 (block B , LZ $(i_1, \dots, i_{k+1})$)

```
{
  for ( $\ell \in [2 : i_{B+1} - i_B - 1]$ )
  {
     $h' = i_{B+1}$ ;
     $q = h' - \ell$ ;
     $\ell_1 = \text{lce}_f(q, h')$ ;
     $\ell_2 = \text{lce}_b(q - 1, h' - 1)$ ;
    if ( $((\ell_1 + \ell_2 \geq \ell) \ \&\& \ (\ell_1 > 0) \ \&\& \ (\ell_2 > 0))$ )
      output ( $\max\{q - \ell_2, q - \ell + 1\}, \ell$ );
  }
}
```

Abbildung 3.37: Algorithmus: Auffinden linkerster Tandem-Repeats mit Zentrum im Block B , die die erste Bedingung von Theorem 3.49 erfüllen

Fall 2: Jetzt suchen wir nach linkesten Vorkommen von Tandem-Repeats mit Zentrum im Block B , deren linkes Ende links vom Block B liegt. In diesem Fall gilt für die Halblänge ℓ des Tandem-Repeats, dass $\ell \in [1 : \ell(B) + \ell(B + 1)]$ mit $\ell(B) = i_{B+1} - i_B$.

Dieser Fall ist in Abbildung 3.38 illustriert. Wir setzen $q := h + \ell$ und bestimmen mit dem Algorithmus für die longest common extensions in konstanter Zeit die folgenden Werte:

$$\begin{aligned} \ell_1 &:= \text{lce}_f(h, q), \\ \ell_2 &:= \text{lce}_b(h - 1, q - 1). \end{aligned}$$

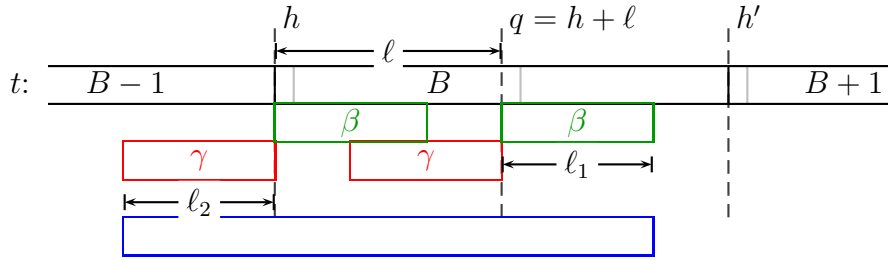


Abbildung 3.38: Skizze: Tandem-Repeats mit Zentrum im Block B , die die zweite Bedingung von Theorem 3.49 erfüllen

Dabei sei β bzw. γ die longest common forward (backward) extension ab dem Positionspaar (h, q') bzw. $(h - 1, q - 1)$. Aufgrund von Theorem 3.49 wissen wir, dass ein linkstes Tandem-Repeat den Anfang von Block B überlappen muss. Es müssen also folgende Bedingungen gelten:

$\ell_1 + \ell_2 \geq \ell$: Auch hier gilt wieder, dass es nur dann einen Tandem-Repeat mit Zentrum im Block B geben kann, wenn $\ell_1 + \ell_2 \geq \ell$ gilt. Die Tandem-Repeats der Länge 2ℓ befinden sich dann im blauen Block in Abbildung 3.38.

$\ell_2 > 0$: Weiter muss $\ell_2 > 0$ sein, damit mindestens ein Tandem-Repeat existiert, dessen linkes Ende links vom Block B liegt.

$\ell_1 \geq 0$: Auch muss $\ell_1 \geq 0$ sein, sonst lässt sich kein Tandem-Repeat mit Zentrum in Block B finden. Dies ist jedoch trivialerweise erfüllt.

$\max\{h - \ell_2, h - \ell\} + \ell < h'$: Damit das Zentrum für das zugehörige Tandem-Repeat wirklich im Block B liegt (hier gilt ja nur $\ell \leq i_{B+2} - i_B$).

Als Tandem-Repeat-Paar geben wir dann $(\max\{h - \ell_2, h - \ell\}, \ell)$ aus. Hierbei ist $(h - \ell_2, \ell)$ ein Tandem-Repeat-Paar, dass alle gefundenen überdeckt (also von diesem Run das linkeste ist). Das Tandem-Repeat-Paar $(h - \ell, \ell)$ ist das linkeste, dessen Zentrum gerade noch im Block $B + 1$ liegt, da $(h - \ell) + \ell = h = i_B$. Der Algorithmus ist in Abbildung 3.39 noch einmal angegeben.

Lemma 3.51 *Die beiden oben angegebenen Algorithmen liefern eine linkeste Überdeckung in Zeit $O(n)$.*

Beweis: Die Korrektheit haben wir bereits bewiesen, da wir nach Theorem 3.49 jedes linkeste Vorkommen eines Tandem-Repeats mit einem Tandem-Repeat-Paar überdeckt haben.

FINDLEFTMOSTTRCOND2 (block B , LZ $(i_1, \dots, i_{k+1})$)

```

{
  for ( $\ell \in [1 : i_{B+2} - i_B]$ )
  {
     $h = i_B$ ;
     $q = h + \ell$ ;
     $\ell_1 = \text{lce}_f(h, q)$ ;
     $\ell_2 = \text{lce}_b(h - 1, q - 1)$ ;
    if ( $(\ell_1 + \ell_2 \geq \ell) \ \&\& \ (\ell_2 > 0) \ \&\& \ (h - \min\{\ell_2, \ell\} < h' - \ell)$ )
      output ( $\max\{h - \ell_2, h - \ell\}, \ell$ );
  }
}

```

Abbildung 3.39: Algorithmus: Auffinden linkerster Tandem-Repeats mit Zentrum im Block B , die die zweite Bedingung von Theorem 3.49 erfüllen

Die Laufzeit zur Bearbeitung eines Blockes B lässt sich mit $c \cdot (\ell(B) + \ell(B + 1))$ für eine geeignete Konstante c abschätzen. Für die Laufzeit $T(n)$ gilt dann, wobei $(i_1, \dots, i_{k+1})$ eine Lempel-Ziv-Zerlegung von t und $\mathcal{B} = \{B_1, \dots, B_k\}$ die Menge aller Lempel-Ziv-Blöcke von t ist:

$$\begin{aligned}
T(n) &\leq \sum_{j=1}^{k-1} c \cdot (\ell(B_j) + \ell(B_{j+1})) + c \cdot (\ell(B_k)) \\
&= \sum_{j=1}^{k-1} c \cdot (i_{j+1} - i_j + i_{j+2} - i_{j+1}) + c \cdot (i_{k+1} - i_k) \\
&= \sum_{j=1}^{k-1} c \cdot (i_{j+2} - i_j) + c \cdot (i_{k+1} - i_k) \\
&= c(i_{k+1} + i_k - i_2 - i_1) + c \cdot (i_{k+1} - i_k) \\
&= c \cdot (2i_{k+1} - i_2 - i_1) \\
&\quad \text{da } i_1 = 1, i_2 = 2 \text{ und } i_{k+1} = n + 1 \\
&\leq c \cdot (2(n + 1) - 2 - 1) \\
&\leq 2c \cdot n.
\end{aligned}$$

Damit ist das Lemma bewiesen. ■

Aus diesem Lemma erhalten wir unmittelbar das folgende Korollar.

Korollar 3.52 *Sei $t \in \Sigma^n$. Die konstruierte linkeste Überdeckung von $\mathcal{T}(t)$ hat eine Größe von $O(n)$*

Im Folgenden führen wir noch eine Partition der Menge P ein, die die angegebenen Algorithmen ohne besonderen Mehraufwand generieren werden.

Notation 3.53 Sei $t \in \Sigma^n$ und sei $P \subseteq \mathcal{T}(t)$ eine Überdeckung von $\mathcal{T}(t)$. Dann ist

$$P(i) := \{(i, \ell) : (i, \ell) \in P\}$$

und es gilt offensichtlich $P = \bigcup_{i=1}^n P(i)$.

Mit dieser Notation können wir zeigen, dass die angegebenen Algorithmen auch die Mengen $P(i)$ als sortierte Listen generieren können.

Lemma 3.54 Der Algorithmus aus Phase I kann so modifiziert werden, dass er die Partition $P(1), \dots, P(n)$ ausgibt und jedes $P(i)$ als aufsteigend sortierte Liste (nach der Länge der zugehörigen Tandem-Repeats) erzeugt wird.

Beweis: Wir müssen bei der Ausgabe nur Folgendes beachten: Jedes neu generierte Tandem-Repeat-Paar (i, ℓ) wird an das Ende der Liste $P(i)$ angehängt.

Wir betrachten jetzt ein festes, neu erzeugtes Tandem-Repeat $\alpha\alpha$ und sein zugehöriges Tandem-Repeat-Paar (i, ℓ) mit Zentrum im Block B . Sei $\alpha'\alpha'$ ein weiteres Tandem-Repeat mit zugehörigem Tandem-Repeat-Paar (i, ℓ') in $P(i)$. Liegt dessen Zentrum in einem Block B' links von B , dann muss $\alpha\alpha$ offensichtlich länger als $\alpha'\alpha'$ sein, wie man der Abbildung 3.40 entnehmen kann

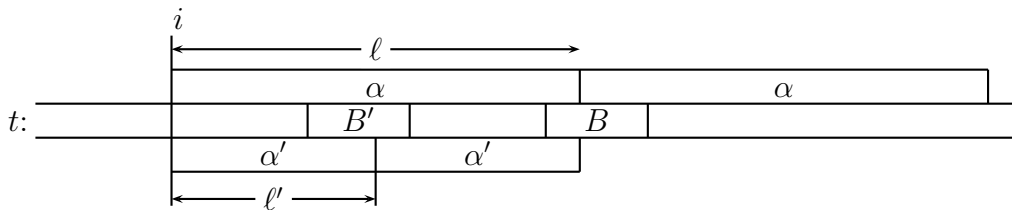


Abbildung 3.40: Skizze: Zwei Tandem-Repeats, die an Position i starten und Zentren in verschiedenen Blöcken besitzen

Sei nun $\alpha'\alpha'$ ein weiteres Tandem-Repeat mit zugehörigem Tandem-Repeat-Paar $(i, \ell') \in P(i)$ und Zentrum in B . Da in beiden Algorithmen die Längen von Tandem-Repeats in aufsteigender Größe betrachtet werden, folgt die Behauptung sofort. ■

Für das Folgende werden wir jedoch benötigen, dass die Listen $P(i)$ absteigend sortiert sind.

Korollar 3.55 *Der Algorithmus aus Phase I kann so modifiziert werden, dass er die Partition $P(1), \dots, P(n)$ ausgibt und jedes $P(i)$ als absteigend sortierte Liste (nach der Länge der zugehörigen Tandem-Repeats) erzeugt wird.*

Beweis: Die erzeugten Tandem-Repeat-Paare müssen nur am Beginn statt am Ende an die Listen angehängt werden. ■

3.3.5 Phase II: Dekorierung einer Teilmenge

Nun wollen wir einige Tandem-Repeat-Paare aus der Menge der in Phase I konstruierten linkensten Überdeckung P im Suffix-Baum $T(t\$)$ markieren.

Sei $t \in \Sigma^n$ und $T = T(t\$)$ der zugehörige Suffix-Baum. Mit $P(v)$ bezeichnen wir für $v \in V(T)$ eine Liste von Tandem-Repeat-Paaren. Für jedes Blatt $v \in V(T)$ soll dabei $P(v) = P(i)$ für $v = t_i \cdots t_n\$$ gelten. Für interne Knoten $v \in V(T)$ hätten wir gerne die folgende Zuordnung:

$$P(v) = \left\{ (j, \ell) : \exists i. \overline{t_i t_i \cdots t_n} \in V(T(v)) : (j, \ell) \in P(i) \wedge 2\ell < |\text{path}(v)| \right\},$$

wobei $T(v)$ den am Knoten v gewurzelten Teilbaum von T bezeichnet. Anschaulich beinhaltet dann $P(v)$ alle Tandem-Repeat-Paare aus P , deren zugehörige Tandem-Repeats echte Präfixe von $\text{path}(v)$ sind. Wir nehmen dabei an, dass die Menge $P(v)$ als Liste dargestellt wird, wobei die Listenelemente absteigend nach der Länge der Tandem-Repeats sortiert sind. Die Definition ist in Abbildung 3.41 noch einmal illustriert.

Der Algorithmus zur Dekorierung mit den Listen $P(v)$ in $T = T(t\$)$ ist in Abbildung 3.42 angegeben. Hierbei ist zunächst angenommen, dass die Listen $P(v)$ wie

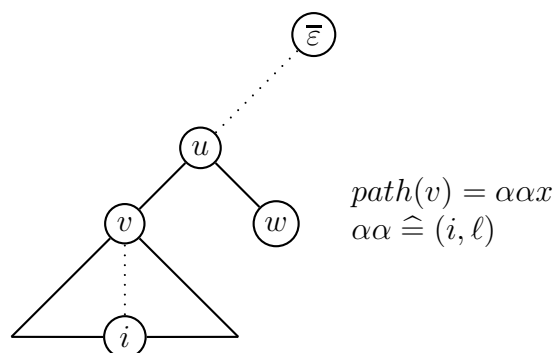


Abbildung 3.41: Skizze: Zur idealen Definition von $P(v)$

```

DECORATETREE (tree  $T = T(t\$)$ )
{
  for each ( $v \in V(T)$ )    /* using DFS */
  {
    /* after returning from all subtrees */
    if ( $v = \overline{t_i \cdots t_n\$}$ )
       $P(v) = P(i)$ ;
    else
      generate  $P(v)$ ;    /* for details see text */
     $u := \text{parent}(v)$ ;
    while ( $P(v) \neq \emptyset$ )
    {
       $(i, \ell) := \text{head}(P(v))$ ;
      if ( $2\ell > |\text{path}(u)|$ )
      {
        add mark  $(2\ell - |\text{path}(u)|)$  at edge  $(u, v)$ ;
         $P(v) := \text{tail}(P(v))$ ;
      }
      else
        exit;    /* while-loop */
    }
  }
}

```

Abbildung 3.42: Algorithmus: Dekoration von $T(t\$)$ mit einer Teilmenge der Tandem-Repeats aus der linkensten Überdeckung von Phase I

oben ideal gegeben sind. Man sieht leicht, dass die Listen $P(v)$ dabei korrekt konstruiert werden, wenn diese aus den Restlisten der Kinder des Knotens v richtig zusammengemischt werden, so dass die Listen $P(v)$ auch absteigend nach der Länge der korrespondierenden Tandem-Repeats sortiert sind.

Für die Dekorierung müssen wir nur die Tandem-Repeat-Paare (i, ℓ) von der Liste $P(v)$ entfernen, für die $2\ell > |\text{path}(u)|$ gilt. Dies sind genau diejenigen, die auf der Kanten (u, v) enden, wobei u der Elter von v ist. Hier wird nun klar, warum die erzeugten Listen aus Phase I absteigend sortiert sein sollten. Wir müssen ja immer die Tandem-Repeat-Paare mit der längsten Halblänge zuerst entfernen.

Idealerweise würde man, nachdem alle Kinder von u abgearbeitet wurden, die restlichen Listen der Kinder von u zur Liste $P(u)$ mischen. Leider kann dies zu einer Gesamtlaufzeit von $O(n^2)$ führen, da das Mischen selbst schon lineare Zeit kosten kann und man im schlimmsten Fall zu oft mischen muss.

Wir werden jetzt folgenden Trick verwenden, um bei einer linearen Laufzeit zu bleiben. Wir wählen als Liste $P'(v)$ für einen inneren Knoten $v \in V(T)$ gerade $P'(v) := P'(v')$, wobei v' ein Kind von v ist. Dabei wählen wir das Kind v' so aus, dass $T(v')$ das Blatt mit der kleinsten Nummer (Index-Position eines Suffixes von t) enthält. Die Listen $P'(v)$ für die Blätter werden genauso wie bei den Listen $P(v)$ gewählt.

Anders ausgedrückt ist die Liste $P'(v)$ eines inneren Knotens genau gleich der Liste des Blattes im Teilbaum von v , dessen zugehörige Indexposition minimal ist.

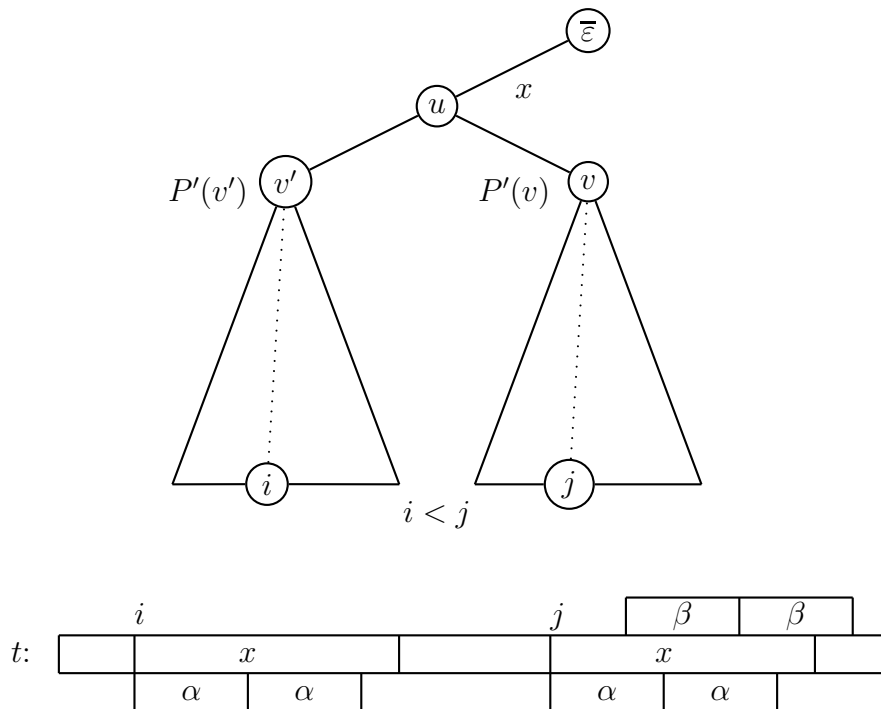


Abbildung 3.43: Skizze: intuitive Korrektheit der Definition von $P'(u)$ als Liste des linken Kindes von u

Warum verschenken wir hierbei keine wertvolle Information? Nehmen wir dazu an, wir wollen $P'(u)$ bestimmen. Seien v und v' zwei Kinder von u , wobei v' das Kind ist, das das Blatt mit dem längsten Suffix enthält. Sei weiter $x = \text{path}(u)$ und $(j, \ell) \in P'(v)$. Dann beschreibt (j, ℓ) einen Tandem-Repeat, der ein Präfix von x ist. Siehe dazu auch Abbildung 3.43.

Da x sowohl als Teilwort von t ab Position i als auch ab Position j vorkommt, wird (j, ℓ) auch von $(i, \ell) \in P'(v')$ dargestellt. Damit wird also das Tandem-Repeat zu (j, ℓ) weiterhin dargestellt und (j, ℓ) kann anscheinend ohne Informationsverlust weggelassen werden. Es ist jedoch möglich, dass andere Tandem-Repeat-Paare wie z.B. $\beta\beta$ in Abbildung 3.43, die von (j, ℓ) überdeckt werden, nun von (i, ℓ) nicht mehr überdeckt werden.

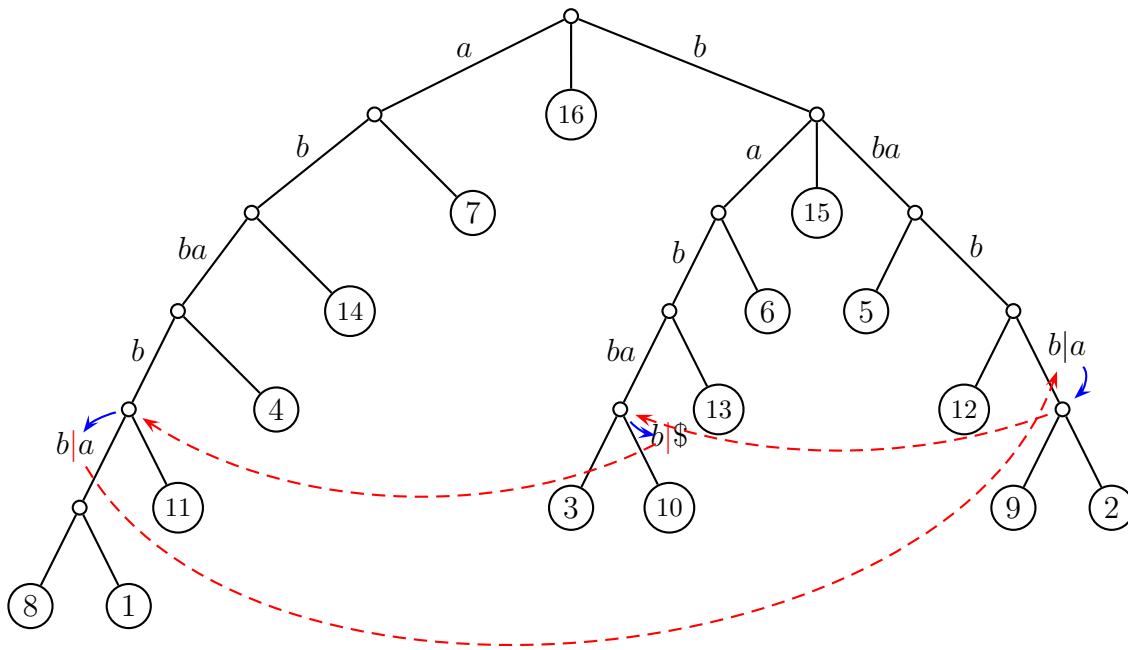


Abbildung 3.44: Beispiel: Auffinden der verlorenen Überdeckung mit Hilfe von Suffix-Link-Walks in $S(abbabbaabbabbab\$)$

Der letzte Fall kann tatsächlich eintreten, wie das folgende Beispiel zeigt. Sei dazu $t = abbabbaabbabbab$, dann ist z.B. $P = \{(1, 3), (9, 3), \dots\}$ eine linkeste Überdeckung. Hier würde am Knoten $v = \overline{abbabba}$ die Beziehung $(1, 3) \in P'(v)$ gelten und damit würden $abbabb$ an der Kante $(\overline{abba}, \overline{abbabba})$ dekoriert. Das Tandem-Repeat-Paar $(9, 3)$ würde jedoch am Knoten $v = \overline{bbabba}$ verloren gehen, d.h. $(9, 3) \notin P'(v)$. Damit wäre zwar das Tandem-Repeat $abbabb$ dargestellt, aber die Überdeckung $babbab$ vom Tandem-Repeat-Paar $(10, 3)$ ginge verloren. Dieses Beispiel ist noch einmal in Abbildung 3.44 illustriert.

Schauen wir uns das Beispiel in Abbildung 3.44 noch einmal an. Wir haben das Tandem-Repeat-Paar $(9, 3)$ in der Liste nicht berücksichtigt und damit die Dekoration von $babbab$ aufgegeben und stattdessen nur das Tandem-Repeat-Paar $(1, 3)$ beibehalten und damit den Tandem-Repeat $abbabb$ dekoriert.

Was passiert nun, wenn wir die Rechtsrotation von $abbabb$ betrachten, dies wäre das Tandem-Repeat $bbabba$. Wir können dieses im Suffix-Baum finden, indem wir von der Lokation für $abbabb$ zuerst dem (eventuell virtuellen) Suffix-Link folgen und müssen dann das Zeichen b noch ablaufen. Im Beispiel in Abbildung 3.44 sind solche (virtuellen) Suffix-Links durch gestrichelte rote Pfeile und das Ablaufen eines Zeichens durch blaue Pfeile dargestellt. Wie man sieht würden wir den Knoten $\overline{bbabba}$ erreichen.

Eine weitere Rechtsrotation liefert das Tandem-Repeat *babbab*. Da auch diese Lokation im Suffix-Baum existiert, haben wir damit auch das nicht überdeckte Tandem-Repeat *babbab* dekoriert. Eine weitere Rechtsrotation führt uns zur Ausgangslokation *abbabb* zurück.

Somit können wir die Hoffnung haben, dass sich ausgehend von den dekorierten Tandem-Repeat durch Rechtsrotationen (Folgen von Suffix-Links von Lokationen und Ablaufen des zugehörigen Zeichens im Suffix-Baum) alle Tandem-Repeat dekorieren lassen. Diese Idee, nicht dekorierte Tandem-Repeat, aufzufinden wollen wir nun noch formalisieren.

07.12.06

Definition 3.56 Sei $t \in \Sigma^*$ und $T = T(t\$)$ der zugehörige Suffix-Baum. Sei weiter $\text{loc}(aw)$ eine Lokation im Suffix-Baum T mit $a \in \Sigma$ und $w \in \Sigma^*$. Das Ablaufen von $\text{loc}(aw) \rightarrow \text{loc}(w) \rightarrow \text{loc}(wa)$ im Baum T heißt Suffix-Link-Walk in T . Er heißt erfolgreich, wenn $\text{loc}(wa)$ existiert, d.h. $wa \sqsubseteq t$, und erfolglos sonst. Eine Folge von Suffix-Link-Walks heißt Kette von Suffix-Link-Walks.

Nun formalisieren wir Mengen von Tandem-Repeat, die wir in Phase II gerne im Suffix-Baum T dekoriert hätten.

Definition 3.57 Sei $t \in \Sigma^*$ und $T = T(t\$)$ der zugehörige Suffix-Baum. Eine Teilmenge $Q \subseteq \mathcal{V}(t)$ heißt ausreichend, wenn für jedes Wort $w \in \mathcal{V}(t)$ die zugehörige Lokation in T durch eine Kette von Suffix-Links-Walks von einer Lokation in T , die zu einem Wort aus Q gehört, erreicht werden kann.

Haben wir in Phase II eine ausreichende Menge von Tandem-Repeat markiert, dann können wir das Vokabular im Suffix-Baum mit Hilfe von Suffix-Link-Walks vervollständigen. Wir zeigen jetzt, dass die in Phase II dekorierte Menge ausreichend ist.

Theorem 3.58 Die dekorierte Teilmenge $Q' \subseteq \mathcal{V}(t)$ aus Phase II ist ausreichend.

Beweis: Sei $Q \subseteq \mathcal{V}(t)$ die zur linkensten Überdeckung P von $\mathcal{T}(t)$ gehörige Menge von Tandem-Repeat. Nach Konstruktion ist Q ausreichend. Sei weiter $Q' \subseteq Q$ die Teilmenge von Tandem-Repeat, die in Phase II im Suffix-Baum T dekoriert wurde. Letztendlich sei $Q'' \subseteq \mathcal{V}$ die Teilmenge von Tandem-Repeat, die nicht mit Hilfe von Ketten von Suffix-Link-Walks beginnend an einem Tandem-Repeat aus Q' erreicht werden kann. Für den Beweis des Satzes genügt es also zu zeigen, dass $Q'' = \emptyset$.

Für einen Widerspruchsbeweis nehmen wir an, dass $Q'' \neq \emptyset$. Sei $P' \subseteq P$ die Teilmenge von Tandem-Repeat-Paaren, die genau den Wörtern aus Q' entsprechen, d.h. die zur Dekoration in Phase II verwendet wurden. Weiter sei $P'' \subseteq P \setminus P'$ eine Menge von Tandem-Repeat-Paaren, die alle Wörter aus Q'' überdecken. Sei weiter $(j, \ell) \in P''$ mit kleinster Startposition und sei $\alpha\alpha$ der Tandem-Repeat von (j, ℓ) , d.h. $t_j \cdots t_{j+2\ell-1} = \alpha\alpha$. Wir unterscheiden jetzt 2 Fälle, je nachdem, ob (j, ℓ) das linkeste Auftreten des Tandem-Repeats $\alpha\alpha$ beschreibt oder nicht.

Fall 1: Sei also (j, ℓ) das linkeste Auftreten von $\alpha\alpha$ in t . Wir betrachten den Knoten $v \in V(T)$ mit minimaler Worttiefe, so dass $\alpha\alpha$ ein Präfix von $\text{path}(v)$ ist. Dann beschreibt jedes Blatt in $T(v)$ einen Suffix von t der mit $\alpha\alpha$ beginnt. Also muss j das Blatt mit der kleinsten Index-Position im Teilbaum $T(v)$ sein. Nach Konstruktion ist (j, ℓ) in einer Liste, die in Phase II nicht gelöscht wird, und wird daher in T dekoriert. Das liefert den gewünschten Widerspruch.

Fall 2: Sei also (j, ℓ) nicht das linkeste Auftreten von $\alpha\alpha$ in t . Sei also (i, ℓ) mit $i < j$ das linkeste Auftreten von $\alpha\alpha$ in t . Dann muss nach Konstruktion ein $(h, \ell) \in P$ existieren, so dass (h, ℓ) das Paar (i, ℓ) überdeckt. Also gilt $h \leq i < j$. Dies ist in Abbildung 3.45 illustriert.

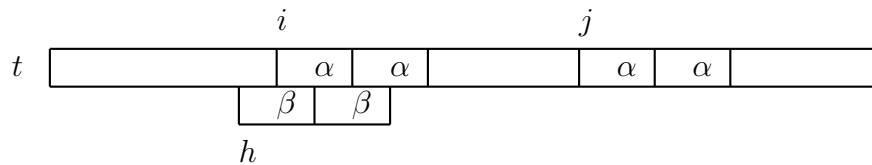


Abbildung 3.45: Skizze: (j, ℓ) beschreibt nicht das linkeste Vorkommen von $\alpha\alpha$ in t

Nach Wahl von j gilt, dass $(h, \ell) \in P'$ und dass (h, ℓ) dekoriert wurde. Da (h, ℓ) aber (i, ℓ) überdeckt und nach Definition dann $[h : i]$ ein Run ist, muss das Tandem-Repeat zu (i, ℓ) durch eine Kette von Suffix-Link-Walks ausgehend vom Tandem-Repeat zu (h, ℓ) , das in Phase II dekoriert wurde, erreicht werden. Damit erhalten wir wieder den gewünschten Widerspruch. ■

3.3.6 Phase III: Vervollständigung der Dekorierung von $\mathcal{V}(t)$

Im letzten Abschnitt haben wir gezeigt, dass die in Phase II dekorierten Tandem-Repeats in T eine ausreichende Menge darstellen. Wir müssen also in Phase III diese nur noch mit Hilfe von Ketten von Suffix-Link-Walks vervollständigen. Der Algorithmus zur Vervollständigung der Dekorierung von $\mathcal{V}(t)$ geht dabei wie folgt vor.

1. Durchlaufe $T(t\$)$ mittels einer Tiefensuche.

2. An dekorierten Lokation aus Q' in T startet eine Kette von Suffix-Link-Walks bis entweder ein erfolgloser Suffix-Link-Walk auftritt oder der Suffix-Link-Walk auf eine bereits erfolgte Dekorierung trifft.
3. Dabei werden alle besuchten Lokationen dekoriert.

Lemma 3.59 *Der Algorithmus aus Phase III vervollständigt eine Dekorierung aus Phase II zum Vokabular der gegebenen Zeichenreihe im zugehörigen Suffix-Baum.*

Beweis: Wir haben bereits gezeigt, dass die Menge Q' ausreichend ist. Endet ein Suffix-Link-Walk erfolglos, dann haben wir alle von diesem Tandem-Repeat durch Suffix-Link-Walks erreichbaren Tandem-Repeats in $\mathcal{V}(t)$ gefunden.

Endet ein Suffix-Link-Walk erfolgreich an einer Lokation, die bereits in Q' ist, können wir auch aufhören, da die folgenden Rechtsrotationen des zuletzt aufgefunden Tandem-Repeats von dieser Lokation in Q' gefunden werden oder wurden.

Dummerweise könnte ein Suffix-Link-Walk auch an einer dekorierten Lokation in T enden, die nicht zu Q' gehört. Wir werden jetzt zeigen, dass dies jedoch nicht passieren kann. Für einen Widerspruchsbeweis nehmen wir an, dass ein Suffix-Link-Walk, gestartet an einem Tandem-Repeat in Q' , an einer dekorierten Lokation endet, die kein Tandem-Repeat in Q' beschreibt.

Eine solche Situation kann nur eintreten, wenn wir beim Folgen eines Suffix-Link-Walks eine nachträgliche Dekorierung erwischen, die wir bei einem früheren Suffix-Link-Walk dekoriert hatten. Sei $wawa$ mit $a \in \Sigma$ und $w \in \Sigma^*$ ein Tandem-Repeat, an dessen Lokation in T sich zwei Suffix-Link-Walks treffen. Nach Definition der Suffix-Links, müssen aber beide von der Lokation gekommen sein, die zum Tandem-Repeats $awaw$ gehört hat. Somit können sich zwei Suffix-Link-Walks nur treffen, wenn einer davon einer leerer Suffix-Link-Walk ist und wir erhalten den gewünschten Widerspruch. ■

In den Übungen wurde bereits das folgende Lemma bewiesen.

Lemma 3.60 *Sei $t \in \Sigma^n$. An jeder Position $i \in [1 : n]$ können maximal zwei rechteste Vorkommen eines Tandem-Repeats beginnen.*

Damit erhalten wir sofort eine Aussage über die Größe des Vokabulars einer Zeichenreihe.

Korollar 3.61 *Sei $t \in \Sigma^n$. Dann besitzt t maximal $2n$ viele Tandem-Repeats.*

Weiter erhalten wir eine Aussage über die Anzahl von Dekorierungen pro Kante im zugehörigen Suffix-Baum.

Korollar 3.62 *Sei $t \in \Sigma^*$ und $T = T(t\$)$ der zugehörige Suffix-Baum. Jede Kante $(u, v) \in E(T)$ besitzt maximal zwei Dekorierungen.*

Beweis: Bezeichne $R(v)$ für jeden Knoten $v \in V(T)$ die maximale Startposition eines Suffixes von t , dessen zugehöriges Blatt im Teilbaum $T(v)$ liegt. Betrachten wir jetzt eine Kante $(u, v) \in E(T)$. Das rechteste Vorkommen der dort endenden Tandem-Repeats muss auch ab Position $R(v)$ beginnen. Da dort nach dem vorhergehenden Lemma nur zwei beginnen können, können auf dieser Kante nur zwei Tandem-Repeats enden. ■

Damit können wir nun die Laufzeit des Algorithmus von Gusfield und Stoye analysieren.

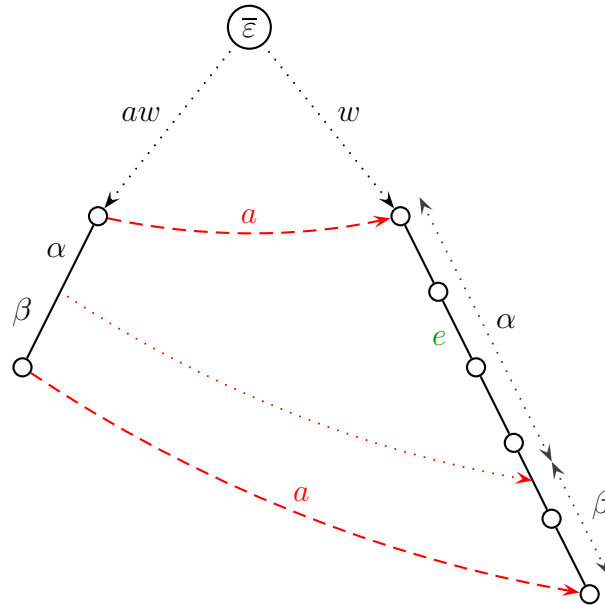
Lemma 3.63 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und $T = T(t\$)$ der zugehörige Suffix-Baum. Die Anzahl der traversierten Kanten in T in Phase III durch die Suffix-Link-Walks ist beschränkt durch $O(|\Sigma| \cdot n)$.*

Beweis: Da es nur maximal $2n$ Tandem-Repeats in t geben kann, kann es nur maximal $O(n)$ Suffix-Link-Walks geben, da jeder Suffix-Link-Walk an einer zugehörigen Lokation eines Tandem-Repeat startet und an solchen Lokationen nur jeweils ein Suffix-Link-Walk gestartet wird. Für die Laufzeit müssen wir jedoch analysieren, wie viele Kanten im Suffix-Baum überlaufen werden. Diese können dummerweise pro Suffix-Link-Walk nicht durch eine Konstante beschränkt werden (ähnlich wie bei Ukkonens Algorithmus).

Es kann also passieren, dass bei einem Suffix-Link-Walk von $\text{loc}(ax)$ über $\text{loc}(x)$ zu $\text{loc}(xa)$ mehrere Kanten im Suffix-Baum überlaufen werden, und zwar beim Ablaufen von $\text{loc}(ax)$ zu $\text{loc}(x)$ (analog wie bei Ukkonens Algorithmus).

Betrachten wir also eine Kante $e \in E(T)$ und ein Zeichen $a \in \Sigma$. Wir fragen uns, wie oft die Kante e bei einem Suffix-Link-Walk übersprungen werden kann, wenn der Suffix-Link-Walk eine Rechtsrotation des Zeichens a beschreibt. Diese Situation ist in Abbildung 3.46 illustriert.

Im Folgenden sagen wir, ein Suffix-Link hat die Markierung $a \in \Sigma$, wenn der Suffix-Link von $\overline{a}z$ nach $\overline{z}$ für ein $z \in \Sigma^*$ ist. Für eine übersprungene Kante $e = (u, v)$ bei einer Rechtsrotation um a ist der gefolgte Suffix-Link derjenige, der als erster auf dem Pfad zur Wurzel in einen Knoten mit Markierung a hineingeht. Wir können also

Abbildung 3.46: Skizze: Ein Suffix-Link-Walk von $\text{loc}(aw\alpha)$ zu $\text{loc}(w\alpha\alpha)$

den zum Suffix-Link-Walk gehörenden Suffix-Link eindeutig identifizieren. Dieser Suffix-Link führt also von $\overline{aw}$ nach $\overline{w}$ für ein $w \in \Sigma^*$.

Weiterhin wissen wir, dass ein Suffix-Link-Walk immer am Ende eines Tandem-Repeats startet. Sei also $aw\alpha$ das auslösende Tandem-Repeat. Dann liegt die Kante e auf dem Pfad von $\overline{w}$ zu $\text{loc}(w\alpha)$.

Dies können wir auch umdrehen. Ist α' die Markierung von $\overline{w}$ zum Endpunkt v der Kanten e , dann wissen wir, dass ein Suffix-Link-Walk, der e überspringt, auf der Kante von $\overline{w}$ begonnen haben muss, dessen Präfix der Kantenmarkierung mit α' beginnt. Es kann also nur eine Kante als Ausgangspunkt des Suffix-Link-Walks zum Überspringen der Kanten e bei einer Rechtsrotation um a in Frage kommen.

Da auf jeder Kante nur maximal zwei Tandem-Repeats markiert sind, kann diese Kante für jedes Zeichen $a \in \Sigma$ nur zweimal übersprungen werden. Also kann diese Kante insgesamt maximal $2 \cdot |\Sigma|$ mal übersprungen werden. Da im Suffix-Baum T nur $O(n)$ Kanten enthalten sind, folgt die Behauptung. ■

Theorem 3.64 *Das Vokabular $\mathcal{V}(t)$ einer Zeichenreihe $t \in \Sigma^n$ kann im Suffix-Baum $T(t\$)$ in Zeit $O(|\Sigma|n)$ mit Platz $O(n)$ dekoriert werden.*

Lowest Common Ancestors

4

4.1 Algorithmus von Bender und Farach-Colton

In diesem Kapitel wollen wir zwei Algorithmen zur Bestimmung eines niedrigsten gemeinsamen Vorfahren vorstellen. Im ersten Abschnitt werden wir dabei das Problem auf so genannte Range Minimum Queries reduzieren.

Der hier vorgestellte Algorithmus wurde zuerst 1989 von Berkman und Vishkin in einer anderen Darstellung beschrieben und in der hier dargestellten Fassung von Bender und Farach-Colton im Jahre 2000 wiederentdeckt.

4.1.1 Lowest Common Ancestor und Range Minimum Queries

Bevor wir die Problemstellung definieren können, benötigen wir noch den Begriff des niedrigsten gemeinsamen Vorfahren von zwei Knoten in einem Baum.

Definition 4.1 Sei $T = (V, E)$ ein gewurzelter Baum und seien $v, w \in V$ zwei Knoten von T . Der niedrigste gemeinsame Vorfahre von v und w , bezeichnet mit $\text{lca}(v, w)$ (engl. least common ancestor, lowest common ancestor oder auch nearest common ancestor), ist der Knoten $u \in V$, so dass u sowohl ein Vorfahre von v als auch von w ist und es keinen echten Nachfahren von u gibt, der ebenfalls ein Vorfahre von v und w ist.

Damit können wir die Lowest Common Ancestor Queries formalisieren.

LOWEST COMMON ANCESTOR QUERY

Eingabe: Ein gewurzelter Baum T und zwei Knoten $v, w \in V(T)$.

Gesucht: Ein Knoten u , der der niedrigste gemeinsame Vorfahre von v und w ist.

Wir werden uns nun überlegen, wie wir in einem Baum mit n Knoten ℓ Lowest Common Ancestor Queries in Zeit $O(n + \ell)$ beantworten können. Dazu werden wir das Lowest Common Ancestor Problem auf das Range Minimum Query Problem reduzieren, das wie folgt definiert ist.

RANGE MINIMUM QUERY

Eingabe: Eine Feld F der Länge n von reellen Zahlen und $i \leq j \in [1 : n]$.

Gesucht: Ein Index k mit $F[k] = \min \{F[\ell] : \ell \in [i : j]\}$.

Man kann Range Minimum Queries natürlich auch allgemeiner für eine total geordnete Menge definieren. Wir werden später zeigen, wie wir mit einer Vorverarbeitung in Zeit $O(n)$ jede solche Anfrage in konstanter Zeit beantworten können. Für die Reduktion betrachten wir die so genannte *Euler-Tour* eines Baumes.

4.1.2 Euler-Tour eines gewurzelten Baumes

Wir definieren nun, was wir unter einer Euler-Tour eines gewurzelten Baumes verstehen wollen. Diese Euler-Tour ist nicht mit Euler-Kreisen oder -Pfaden in Graphen zu verwechseln.

Definition 4.2 Sei $T = (V, E)$ ein gewurzelter Baum mit Wurzel r und seien $T_1, \dots, T_\ell$ die Teilbäume, die an der Wurzel r hängen. Die Euler-Tour durch T ist eine Liste von $2|V| - 1$ Knoten, die wie folgt rekursiv definiert ist. Für $\ell \geq 0$ seien $E_1, \dots, E_\ell$ mit $E_i = (v_1^{(i)}, \dots, v_{n_i}^{(i)})$ für $i \in [1 : \ell]$ die Euler-Touren von $T_1, \dots, T_\ell$. Die Euler-Tour von T ist dann definiert durch

$$(r, v_1^{(1)}, \dots, v_{n_1}^{(1)}, r, v_1^{(2)}, \dots, v_{n_2}^{(2)}, r, \dots, r, v_1^{(\ell)}, \dots, v_{n_\ell}^{(\ell)}, r).$$

Man beachte, dass in der obigen Definition auch $\ell = 0$ sein kann, d.h. der Baum besteht nur aus dem Blatt r . Dann ist diese Liste durch (r) gegeben. Der Leser sei dazu aufgefordert zu verifizieren, dass die oben definierte Euler-Tour eines Baumes mit n Knoten tatsächlich eine Liste mit $2n - 1$ Elementen ist.

EULER-TOUR (tree $T = (V, E)$)

```
{
  node  $v = \text{root}(T)$ ;
  output  $v$ ;
  for all  $((v, w) \in E(T))$ 
  {
    EULER-TOUR( $T(w)$ ) /*  $T(w)$  denotes the subtree rooted at  $w$  */
    output  $v$ ;
  }
}
```

Abbildung 4.1: Algorithmus: Konstruktion einer Euler-Tour

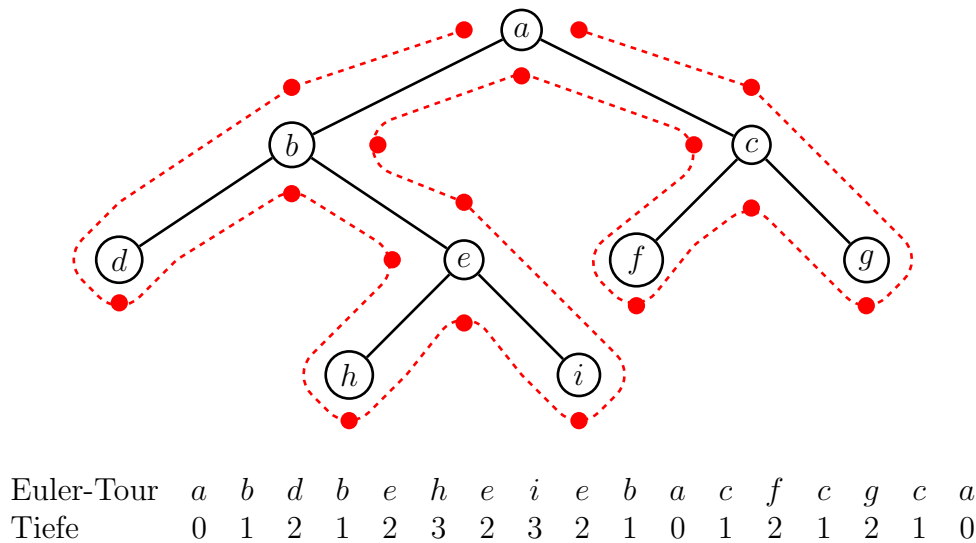


Abbildung 4.2: Beispiel: Euler-Tour

Die Euler-Tour kann sehr leicht mit Hilfe einer Tiefensuche in Zeit $O(n)$ berechnet werden. Der Algorithmus hierfür ist in Abbildung 4.1 angegeben. Man kann sich die Euler-Tour auch bildlich sehr schön als das Abmalen der Bäume anhand ihrer äußeren Kontur vorstellen, wobei bei jedem Antreffen eines Knotens des Baumes dieser in die Liste aufgenommen wird. Dies ist in Abbildung 4.2 anhand eines Beispiels illustriert.

Definition 4.3 Sei $T = (V, E)$ ein gewurzelter Baum. Die Tiefe eines Knotens ist die Anzahl der Kanten auf dem einfachen Weg von der Wurzel zu diesem Knoten. Die maximale Tiefe eines Knotens im Baum T bezeichnet man als die Tiefe des Baumes.

Zusammen mit der Euler-Tour, d.h. der Liste der abgelaufenen Knoten, betrachten wir zusätzlich noch die Tiefe des entsprechenden Knotens, die bei der Tiefensuche in der Regel mitberechnet werden (siehe auch das Beispiel in der Abbildung 4.2).

4.1.3 Reduktion LCA auf RMQ

Betrachten wir jetzt die Anfrage an zwei Knoten des Baumes i und j . Zuerst bemerken wir, dass diese Knoten, sofern sie keine Blätter sind, in der Euler-Tour mehrfach vorkommen. Wir wählen jetzt für i und j willkürlich einen der Knoten, der in der Euler-Tour auftritt, als einen Repräsentanten aus. Weiter stellen wir fest, dass in der Euler-Tour der niedrigste gemeinsame Vorfahre von i und j in der Teilliste, die

durch die beiden Repräsentanten definiert ist, vorkommen muss, was man wie folgt sieht.

Nehmen wir an, dass i in der Euler-Tour vor j auftritt. Betrachten wir die Teilliste der Euler-Tour von einem Repräsentanten i' von i zu einem Repräsentanten j' von j . Sei k der niedrigste gemeinsame Vorfahre von i und j . Da die Tiefensuche von k bereits aktiv sein muss, wenn i' in die Euler-Tour aufgenommen wird, und k noch aktiv sein muss, wenn j' in die Euler-Tour aufgenommen wird, kann also in die Teilliste zwischen den Repräsentanten i' und j' kein Repräsentant eines Knoten mit einer Tiefe kleiner oder gleich der Tiefe des Knotens k aufgenommen werden.

Andererseits muss bei der Rückkehr von i auf dem Weg zu j nach der Definition eines niedrigsten gemeinsamen Vorfahrs der Knoten k besucht werden und in die Euler-Tour aufgenommen werden. Somit befindet sich in der Teilliste der Euler-Tour zwischen beliebigen Repräsentanten von i und von j der Knoten k mindestens einmal und es ist der einzige Knoten mit minimaler Tiefe in dieser Teilliste. Damit können wir das so erhaltene Zwischenergebnis im folgenden Lemma festhalten.

Lemma 4.4 *Gibt es eine Lösung für das Range Minimum Query Problem, dass für die Vorverarbeitung Zeit $O(p(n))$ und für eine Anfrage $O(q(n))$ benötigt, so kann das Problem des niedrigsten gemeinsamen Vorfahren mit einem Zeitbedarf für die Vorverarbeitung in Zeit $O(n + p(2n - 1))$ und für eine Anfrage in Zeit $O(q(2n - 1))$ gelöst werden.*

Damit können wir uns jetzt ganz auf das Range Minimum Query Problem konzentrieren. Offensichtlich kann ohne eine Vorverarbeitung eine einzelne Anfrage mit $O(j - i + 1) = O(n)$ Vergleichen beantwortet werden. Das Problem der Range Minimum Queries ist jedoch insbesondere dann interessant, wenn für ein gegebenes Feld eine Vielzahl von Range Minimum Queries durchgeführt werden. In diesem Fall können mit Hilfe einer Vorverarbeitung die Kosten der einzelnen Queries gesenkt werden.

4.1.4 Ein quadratischer Algorithmus für RMQ

Eine triviale Lösung würde alle möglichen Anfragen vorab berechnen. Dazu könnte eine zweidimensionale Tabelle $Q[i, j]$ für $i \leq j \in [1 : n]$ wie folgt angelegt werden:

$$Q[i, j] = \operatorname{argmin} \{F[\ell] : \ell \in [i : j]\}.$$

Dazu würde für jedes Paar (i, j) das Minimum der Bereichs $F[i : j]$ mit $j - i$ Vergleichen bestimmt werden. Dies würde zu einer Laufzeit für die Vorverarbeitung

```

RMQ (int  $F[]$ , int  $n$ )
{
    for ( $i = 1$ ;  $i \leq n$ ;  $i++$ )
         $Q[i, i] = i$ ;

    for ( $i = 1$ ;  $i \leq n$ ;  $i++$ )
        for ( $j = i + 1$ ;  $i + j \leq n$ ;  $j++$ )
            if ( $F[Q[i, j - 1]] \leq F[i + j]$ )
                 $Q[i, j] = Q[i, j - 1]$ ;
            else
                 $Q[i, j] = i + j$ ;
}

```

Abbildung 4.3: Algorithmus: Vorverarbeitung für Range Minimum Queries

von

$$\sum_{i=1}^n \sum_{j=i}^n (j - i) = \Theta(n^3)$$

führen. In der Abbildung 4.3 ist ein einfacher, auf dynamischer Programmierung basierender Algorithmus angegeben, der diese Tabelle in Zeit $O(n^2)$ berechnen kann. Damit erhalten wir das folgende erste Resultat für die Komplexität des Range Minimum Query Problems.

Theorem 4.5 *Für das Range Minimum Query Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n^2)$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

4.1.5 Eine verbesserte Variante

Eine schnellere Variante erhalten wir, wenn wir nicht alle Werte $Q[i, j]$ berechnen, sondern nur für solche Paare (i, j) deren Differenz $j - i + 1$ eine Zweierpotenz ist. Dazu definieren wir eine Tabelle Q' für $i \in [1 : n]$ und $k \in [0 : \lfloor \log(n) \rfloor]$ wie folgt:

$$Q'[i, k] = \operatorname{argmin} \{ F[\ell] : \ell \in [i : i + 2^k - 1] \}.$$

Da die Tabelle nur $\Theta(n \log(n))$ Einträge besitzt, kann diese Tabelle mittels dynamischer Programmierung in Zeit $O(n \log(n))$ berechnet werden, wie in Abbildung 4.4 dargestellt.

Wie beantworten wir jetzt eine gegebene Anfrage $\text{RMQ}(i, j)$? Wir berechnen zuerst $k := \lfloor \log(j - i + 1) \rfloor$. Also gilt $2^k \leq j - i + 1 < 2^{k+1}$. Dann ermitteln wir die Indizes

```

RMQ2 (int  $F[]$ , int  $n$ )
{
    for ( $i = 1$ ;  $i \leq n$ ;  $i++$ )
         $Q'[i, 0] = i$ ;

    for ( $k = 0$ ;  $k \leq \lfloor \log(n) \rfloor$ ;  $k++$ )
        for ( $i = 1$ ;  $i + 2^k \leq n$ ;  $i++$ )
            if ( $F[Q'[i, k]] \leq F[Q'[i + 2^k, k]]$ )
                 $Q'[i, k + 1] = Q'[i, k]$ ;
            else
                 $Q'[i, k + 1] = Q'[i + 2^k, k]$ ;
}

```

Abbildung 4.4: Algorithmus: Besseres Vorverarbeitung für Range Minimum Queries

der Minima in F im Bereich $[i : i + 2^k - 1]$ bzw. $[j - 2^k + 1 : j]$ mittels $r := Q'[i, k]$ bzw. $s := Q'[j - 2^k + 1, k]$. Nach Wahl von k gilt $[i : i + 2^k - 1] \cup [j - 2^k + 1 : j] = [i : j]$. Gilt dann $F[r] \leq F[s]$, dann ist die Antwort r , ansonsten s .

Theorem 4.6 *Für das Range Minimum Query Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n \log(n))$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

4.1.6 Incremental Range Minimum Query (*)

Betrachten wir das Feld F näher, das aus der Euler-Tour eines Baumes entsteht. Man beobachtet sofort, dass sich die Werte von konsekutiven Feldelementen um genau 1 unterscheiden. Daher wollen wir eigentlich das folgende eingeschränkte Problem lösen.

INCREMENTAL RANGE MINIMUM QUERY

Eingabe: Ein Feld F der Länge n von reellen Zahlen, wobei $|F[i + 1] - F[i]| = 1$ für alle $i \in [1 : n - 1]$ gilt, und $i \leq j \in [1 : n]$.

Gesucht: Ein (minimaler) Index k mit $F[k] = \min \{F[\ell] : \ell \in [i : j]\}$.

Wir halten zunächst die folgende fundamentale Eigenschaften fest.

Beobachtung 4.7 *Seien F und G zwei Felder der Länge n mit reellen Einträgen und sei $c \in \mathbb{R}$ eine Konstante. Gilt $F[k] - G[k] = c$ für alle $k \in [1 : n]$, dann gilt $RMQ_F(i, j) = RMQ_G(i, j)$.*

Somit können wir für Incremental Range Minimum Queries annehmen, dass für ein zu verarbeitendes Feld F gilt, dass $F[1] = 0$, indem man von allen Feldelementen $F[1]$ abzieht.

Definition 4.8 *Sei F ein Feld der Länge n mit reellen Einträgen. Das Feld F heißt inkrementell, wenn $|F[i+1] - F[i]| = 1$ für alle $i \in [1 : n-1]$. Das Feld F heißt normalisiert, wenn $F[1] = 0$.*

Ausgehend von dieser Definition sieht man sofort, dass es nur eine beschränkte Anzahl normalisierter, inkrementeller Felder gibt.

Lemma 4.9 *Es gibt 2^{n-1} verschiedene normalisierte, inkrementelle Felder der Länge n .*

Beweis: Da das Feld F normalisiert ist, gilt $F[1] = 0$. Da das Feld F inkrementell ist, muss man sich für jedes Folglied nur merken, ob der Wert um 1 erhöht oder erniedrigt wird. Dies sind insgesamt $n - 1$ Bit-Informationen. ■

4.1.7 Ein optimaler Algorithmus (*)

Für einen optimalen Algorithmus mit linearer Vorverarbeitungszeit unterteilen wir das gegebene Feld F in n/k Blöcke der Länge jeweils k (außer eventuell dem letzten) mit $k := \lceil \frac{1}{2} \log(n) \rceil$. Sei dann F' ein Feld der Länge $n' := \lceil n/k \rceil$, wobei gilt:

$$F'[i] = \min \{ F[j] : j \in [(i-1)k + 1 : \min\{i \cdot k, n\}] \}.$$

Anschaulich steht im i -ten Eintrag von Feld F' das Minimum des i -ten Blocks von Feld F . Weiter sei P' ein Feld der gleichen Länge wie F' , wobei $P'[i]$ die Position angibt, an welcher Position im i -ten Block von F das Minimum angenommen wird, d.h.

$$P'[i] = \operatorname{argmin} \{ F[j] : j \in [(i-1)k + 1 : \max\{i \cdot k, n\}] \} - (i-1)k.$$

Die Felder F' und P' können offensichtlich in linearer Zeit ermittelt werden. In der Regel wird man das Feld F' nicht abspeichern, da $F'[i] = F[P'[i] + (i-1)k]$ gilt.

Wir können jetzt das Feld F' der Länge $n' = \lceil n/k \rceil$ mit unserem Algorithmus mit Laufzeit $O(n' \log(n'))$ vorverarbeiten, dann lässt sich jede Range Minimum Query

auf F' in konstanter Zeit beantworten. Die Laufzeit für diese Vorverarbeitung von F' beträgt

$$O(n' \log(n')) = O\left(\frac{n}{\log(n)} \log\left(\frac{n}{\log(n)}\right)\right) = O(n).$$

Wie beantworten wir jetzt eine Anfrage $\text{RMQ}(i, j)$? Zuerst ermitteln wir die Blöcke i' bzw. j' in denen sich i bzw. j befinden. Gilt $i' = j'$, dann ist die Position $p_i = \text{RMQ}_F(i, j)$ (innerhalb des Blockes $i' = j'$ von F) die Lösung. Gilt andererseits $i' < j'$, dann ermitteln wir die folgenden drei Positionen:

- Die Position $p_i = \text{RMQ}_F(i, k \cdot i')$;
- Falls $i' + 1 < j'$ ist, die Position $p' = \text{RMQ}_F(k \cdot i' + 1, k(j' - 1))$, dabei gilt $p' = (q - 1) \cdot k + P'[q]$ mit $q := \text{RMQ}_{F'}(i' + 1, j' - 1)$;
- Die Position $p_j = \text{RMQ}_F(k(j' - 1) + 1, j)$.

Aus diesen drei Werten kann dann die Position mit dem minimalen Element mit zwei weiteren Vergleichen ermittelt werden.

Dabei kann p' nach der Vorverarbeitung von F' in konstanter Zeit ermittelt werden. Wir müssen also nur für p_i bzw. p_j ein effizientes Verfahren angeben. Dabei finden diese Anfragen immer innerhalb eines Blockes der Länge k statt.

Nach Lemma 4.9 gibt es nur

$$2^{k-1} \leq 2^{\lceil \log(n)/2 \rceil - 1} \leq 2^{\log(n)/2} = O(\sqrt{n})$$

viele normalisierte, inkrementelle Felder der Länge k . Anstatt für alle n/k Teilfelder jeweils eine Lookup-Tabelle für die Incremental Range Minimum Queries zu konstruieren, konstruieren wir für alle mögliche normalisierten, inkrementellen Blöcke jeweils nur eine Tabelle mittels der einfachen dynamischen Programmierung, von denen es ja nur $O(\sqrt{n})$ viele gibt. Die Vorverarbeitungszeit hierfür beträgt offensichtlich:

$$O(\sqrt{n} \cdot \log^2(n)) = O(n).$$

Für jeden der $n' = \lceil n/k \rceil$ Blöcke müssen wir nur noch einen Verweis auf die entsprechende Tabelle angeben, was sich leicht in linearer Zeit bestimmen lässt. Dann können wir eine Incremental Range Minimum Query innerhalb eines Blocks in mit Hilfe der konstruierten Tabellen in konstanter Zeit beantworten.

Theorem 4.10 *Für das Incremental Range Minimum Query Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n)$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

Damit erhalten wir das für uns wichtige Theorem.

Theorem 4.11 *Für das Lowest Common Ancestor Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n)$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

4.1.8 Optimale Lösung für RMQ (*)

In diesem letzten Teil wollen wir noch zeigen, dass wir auch beim allgemeinen Range Minimum Query Problem mit einer linearen Vorverarbeitungszeit jede Anfrage in konstanter Zeit beantworten können. Hierfür reduzieren wir das Problem lustigerweise auf ein LCA-Problem. Um die Reduktion beschreiben zu können, benötigen wir noch das Konzept eines Kartesischen Baumes.

Definition 4.12 *Sei F ein Feld der Länge n mit reellen Einträgen. Der Kartesische Baum des Feldes F ist ein geordneter gewurzelter Baum $T = (V, E)$ mit n Knoten. Die Wurzel ist markiert mit dem Index $k = \operatorname{argmin} \{F[i] : i \in [1 : n]\}$. An der Wurzel hängen zwei Teilbäume (die auch leer sein können), wobei der linke bzw. rechte Teilbaum ein Kartesischer Baum für die Teilfelder $F[1 : k - 1]$ bzw. $F[k + 1 : n]$ ist.*

Man beachte, dass die Definition nicht eindeutig ist, da es durchaus mehrere Minima in einem Feld geben kann. Wir wählen hier im Folgenden immer das Minimum mit dem kleinsten Index. Die folgende Beobachtung folgt unmittelbar aus der Definition eines Kartesischen Baumes.

Beobachtung 4.13 *Sei F ein Feld der Länge n mit reellen Einträgen und T der zugehörige Kartesische Baum. Eine Inorder-Nummerierung von T ergibt die Folge $1, \dots, n$.*

Wir zeigen jetzt, wie man einen Kartesischen Baum für ein Feld F in linearer Zeit konstruieren kann. Sei T ein Kartesischer Baum für das Feld F mit den Indizes aus $[1 : n - 1]$. Wir werden jetzt T so modifizieren, dass auch das Feldelement $F[n]$ berücksichtigt wird.

Wir betrachten dazu den Pfad vom Knoten v mit Markierung $n - 1$ zur Wurzel r des Baumes T . Für die Markierungen $(i_1, \dots, i_\ell)$ der Knoten auf diesem Pfad mit $n - 1 = i_1$ und $i_\ell = r$ gilt nach Konstruktion $F[i_1] \geq \dots \geq F[i_\ell]$. Wir laufen nun diesen Pfad vom Blatt v zur Wurzel ab und vergleichen, ob $F[n] < F[i_k]$ für jeden Knoten mit Markierung i_k gilt. Sobald diese Bedingung nicht mehr erfüllt ist (also $F[i_k] \leq F[n]$ gilt), erzeugen wir einen neuen Knoten mit Markierung n und machen

ihn zum rechten Kind des Knotens mit Markierung i_k . Der Teilbaum mit Wurzel i_{k-1} wird zum linken Teilbaum des neuen Knotens mit Markierung n . Falls wir über die Wurzel hinauslaufen (also eigentlich $k = \ell + 1$ wird), erzeugen wir eine neue Wurzel mit Markierung n die nur einen linken Teilbaum besitzt, nämlich den alten Kartesischen Baum.

Man überlegt sich leicht, dass dieser Algorithmus wirklich einen Kartesischen Baum konstruiert. Wir müssen uns nur noch überlegen, ob dieser auch in linearer Zeit konstruiert werden kann. Hierfür bezeichne $\ell(i)$ die Länge des Pfades (Anzahl der Knoten des Pfades) von der Wurzel zum Knoten mit Markierung i im Kartesischen Baum für $F[1 : i]$. Nach Konstruktion gilt $1 \leq \ell(i) \leq \ell(i-1) + 1$. Das heißt, dass beim Einfügen des i -ten Knotens der rechteste Pfad auf eine Länge von eins schrumpfen kann (nämlich dann, wenn der neue Knoten zur Wurzel wird), aber maximal um ein länger werden kann (nämlich dann, wenn der neue Knoten als ein neues Blatt an den rechtesten Pfad angehängt wird).

Die Kosten, um den i -ten Knoten anzuhängen, sind also proportional zur Anzahl der Knoten auf dem rechtesten Pfad des abgehängten Teilbaum plus eins, also $O(\ell(i) - \ell(i-1) + 1)$. Somit ergibt sich für die Gesamtlaufzeit:

$$\sum_{i=1}^n O(\ell(i) - \ell(i-1) + 1) = O(\ell(n) - \ell(0) + n) \leq O(n - 0 + n) = O(n).$$

Damit haben wir das folgende Lemma bewiesen.

Lemma 4.14 *Sei F ein Feld der Länge n mit reellen Einträgen. Der Kartesische Baum für F kann in Zeit und Platz $O(n)$ konstruiert werden.*

Für die Lösung unseres Problems behaupten wir jetzt, dass $\text{RMQ}_F(i, j) = \text{lca}_T(i, j)$ gilt, wobei ohne Beschränkung der Allgemeinheit $i \leq j$ gilt. Für den Beweis sei $k = \text{lca}_T(i, j)$.

Wir nehmen zuerst an, dass $k \notin \{i, j\}$. Im Kartesischen Baum T ist auf dem Weg von der Wurzel zu i bzw. j der Knoten k der letzte auf dem gemeinsamen Weg. Sei v bzw. w das linke bzw. rechte Kind von k . Dann muss sich i im Teilbaum von v und j im Teilbaum von w befinden. Sei $\underline{k}$ bzw. $\bar{k}$ die kleinste bzw. die größte Markierung im Teilbaum von v bzw. w . Nach Konstruktion des Kartesischen Baumes gilt für alle $x \in [\underline{k} : \bar{k}]$, dass $F[x] \geq F[k]$. Somit ist $k = \text{RMQ}_F(\underline{k}, \bar{k})$. Da $[i : j] \subseteq [\underline{k} : \bar{k}]$ und $k \in [i : j]$, muss auch $k = \text{RMQ}_F(i, j)$ gelten.

Der Fall $k \in \{i, j\}$ ist sehr ähnlich und der Beweis bleibt dem Leser zur Übung überlassen. Somit haben wir den folgenden Satz bewiesen.

Theorem 4.15 *Für das Range Minimum Query Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n)$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

Als Korollar erhalten wir außerdem noch das Ergebnis, dass man auch das RMQ-Problem auf ein LCA-Problem reduzieren kann.

Korollar 4.16 *Gibt es eine Lösung für das Lowest Common Ancestor Problem, dass für das Vorverarbeitung Zeit $O(p(n))$ und für eine Anfrage $O(q(n))$ benötigt, so kann das Range Minimum Query Problem mit einem Zeitbedarf für das Vorverarbeitung in Zeit $O(n + p(2n - 1))$ und für eine Anfrage in Zeit $O(q(2n - 1))$ gelöst werden.*

4.1.9 Eine einfachere optimale Variante nach Alstrup et al.

In diesem Abschnitt stellen wir noch eine einfachere Variante vor, die auf die Arbeit von Alstrup et al. aus dem Jahre 2002 zurückgeht. Hierbei nehmen wir allerdings an, dass einige einfache Bit-Operationen, wie bitweises Und, bitweises Oder sowie Bitscans in konstanter Zeit auf Bit-Vektoren der Länge $\log(n)$ durchgeführt werden können. Intel-Prozessoren stellen diese Operationen beispielsweise zur Verfügung. Man beachte, dass auf 64-Bit-Prozessoren so lange Felder (mit 2^{64} Elementen, also mehreren Trilliarden) schon gar nicht mehr im Hauptspeicher gehalten werden können.

Für einen optimalen Algorithmus mit linearer Vorverarbeitungszeit unterteilen wir das gegebene Feld F in n/k Blöcke der Länge jeweils k (außer eventuell dem letzten) mit $k := \lceil \log(n) \rceil$. Sei dann F' ein Feld der Länge $n' := \lceil n/k \rceil$, wobei gilt:

$$F'[i] = \min \{ F[j] : j \in [(i-1)k + 1 : \min\{i \cdot k, n\}] \}.$$

Anschaulich steht im i -ten Eintrag von Feld F' das Minimum des i -ten Blocks von Feld F . Weiter sei P' ein Feld der gleichen Länge wie F' , wobei $P'[i]$ die Position angibt, an welcher Position im i -ten Block von F das Minimum angenommen wird, d.h.

$$P'[i] = \operatorname{argmin} \{ F[j] : j \in [(i-1)k + 1 : \max\{i \cdot k, n\}] \} - (i-1)k.$$

Die Felder F' und P' können offensichtlich in linearer Zeit ermittelt werden. In der Regel wird man das Feld F' nicht abspeichern, da $F'[i] = F[P'[i] + (i-1)k]$ gilt.

Wir können jetzt das Feld F' der Länge $n' = \lceil n/k \rceil$ mit unserem Algorithmus mit Laufzeit $O(n' \log(n'))$ vorverarbeiten, dann lässt sich jede Range Minimum Query

auf F' in konstanter Zeit beantworten. Die Laufzeit für diese Vorverarbeitung von F' beträgt

$$O(n' \log(n')) = O\left(\frac{n}{\log(n)} \log\left(\frac{n}{\log(n)}\right)\right) = O(n).$$

Wie beantworten wir jetzt eine Anfrage $\text{RMQ}(i, j)$? Zuerst ermitteln wir die Blöcke i' bzw. j' in denen sich i bzw. j befinden. Gilt $i' = j'$, dann ist die Position $p_i = \text{RMQ}_F(i, j)$ (innerhalb des Blockes $i' = j'$ von F) die Lösung. Gilt andererseits $i' < j'$, dann ermitteln wir die folgenden drei Positionen:

- Die Position $p_i = \text{RMQ}_F(i, k \cdot i')$;
- Falls $i' + 1 < j'$ ist, die Position $p' = \text{RMQ}_F(k \cdot i' + 1, k(j' - 1))$, dabei gilt $p' = (q - 1) \cdot k + P'[q]$ mit $q := \text{RMQ}_{F'}(i' + 1, j' - 1)$;
- Die Position $p_j = \text{RMQ}_F(k(j' - 1) + 1, j)$.

Aus diesen drei Werten kann dann die Position mit dem minimalen Element mit zwei weiteren Vergleichen ermittelt werden.

Dabei kann p' nach der Vorverarbeitung von F' in konstanter Zeit ermittelt werden. Wir müssen also nur für p_i bzw. p_j ein effizientes Verfahren angeben. Dabei finden diese Anfragen immer innerhalb eines Blockes der Länge k statt.

Betrachten wir nun einen Block B der Länge k und definieren hierfür k Bit-Vektoren $V^{(B,j)}$ für $j \in [1 : k]$ der Länge jeweils k wie folgt:

$$V_i^{(B,j)} = \begin{cases} 1 & \text{falls } i < j \wedge \forall \ell \in [i : j] : B[i] \leq B[\ell], \\ 0 & \text{sonst.} \end{cases}$$

Intuitiv bedeutet dies, dass genau dann $\text{RMQ}_B(i, j) = i$, wenn $V_i^{(B,j)} = 1$.

Angenommen, wir hätten für jeden Block jeweils diese k Bit-Vektoren, dann lässt sich eine RMQ-Anfrage innerhalb dieses Blockes wie folgt bestimmen

$$\text{RMQ}_B(i, j) = \begin{cases} j & \text{falls } V^{(B,j)} \wedge 0^{i-1}1^{k-i+1} = 0^k, \\ \mathbf{bsf}(V^{(B,j)} \wedge 0^{i-1}1^{k-i+1}) & \text{falls } V^{(B,j)} \wedge 0^{i-1}1^{k-i+1} \neq 0^k. \end{cases}$$

wobei $\wedge$ das bitweise Und ist sowie $\mathbf{bsf}(v)$ die Position (aus $[1 : k]$) der linkensten 1 eines Bitvektors v liefert (Bit Scan Forward):

$$\mathbf{bsf}(v) = i \quad :\Leftrightarrow \quad v_i = 1 \wedge \forall \ell \in [1 : i - 1] : v_\ell = 0.$$

Beachte, dass nach Definition $\mathbf{bsf}(0^k) = k + 1$.

	i	$RMQ_B(i, j)$	j
$V^{(B,j)}$	$* \dots * 0$	1	$* \dots * 0$
$0^{i-1} 1^{k-i+1}$	$0 \dots 0$	1	$1 \dots 1$

Abbildung 4.5: Skizze: Bit Scan Forward für RMQ-Anfrage

Mit Hilfe von Abbildung 4.5 wird nun die Berechnung der RMQ-Anfrage klar. Wir suchen das linkeste gesetzte Bit rechts von Position i , da für diese Position i' gilt $RMQ_B(i', j) = i'$ und diese die linkeste Position im Intervall $[i : j]$ ist.

Wie zu Beginn dieses Abschnittes erwähnt, haben wir angenommen, dass diese Operationen in konstanter Zeit durchführbar sind. Somit lässt sie eine RMQ-Anfrage innerhalb eines Blockes mit Hilfe der Bit-Vektoren $V^{(B,j)}$ in konstanter Zeit beantworten.

Wir müssen uns nur noch überlegen, wie man diese Bit-Vektoren $V^{(B,j)}$ bestimmt. Hierzu sei B zunächst ein beliebiger, aber fester Block. Wir bestimmen folgenden Vektor $L^B \in [0 : k]^k$, der wie folgt definiert ist:

$$L^B[j] = \max \{i \in [0 : j - 1] : B[i] \leq B[j]\}.$$

Hierbei nehmen wir an, dass $B[0] = -\infty$ ist, damit die Maximumsbildung über eine nichtleere Menge gebildet wird. Anschaulich ist $L^B[j]$ der größte Index $i < j$, so dass $B[i]$ kleiner als $B[j]$ ist. Somit ist $L^B[j]$ die größte Bit-Position an der der Vektor $V^{(B,j)}$ den Wert 1 annimmt. Daraus ergibt sich die folgende rekursive Definition von $V^{(B,j)}$:

$$V^{(B,j)} = \begin{cases} V^{(B,L^B[j])} \vee 0^{L^B[j]-1} 1 0^{k-L^B[j]+1} & \text{falls } L^B[j] > 0, \\ 0^k & \text{falls } L^B[j] = 0. \end{cases}$$

Wir müssen also nur noch den Vektor L^B ermitteln. Als Wert für $L^B[j]$ ist eine Indexposition i nur dann interessant, wenn $L^B[i] \leq L^B[\ell]$ für alle $\ell \in [i + 1 : j - 1]$. Diese Indexpositionen $(i_0, \dots, i_\nu)$ mit

$$0 = i_0 < i_1 < \dots < i_\nu \quad \text{und} \quad -\infty = L^B[i_0] \leq L^B[i_1] \leq \dots \leq L^B[i_\nu]$$

merken wir uns im Algorithmus in Abbildung 4.6 mit Hilfe eines Kellers (der mit dem Wert $i_0 = 0$ initialisiert wird).

Der Algorithmus in Abbildung 4.6 hat eine Laufzeit von $O(k)$, da jeder Index aus $[1 : k]$ nur einmal auf dem Keller abgelegt wird und somit auch nur einmal vom Keller entfernt werden kann. Somit ist die Laufzeit $O(n' \cdot k) = O(\frac{n}{k} \cdot k) = O(n)$ und wir erhalten wiederum den folgenden Satz.

Theorem 4.17 *Für das Range Minimum Query Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n)$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

```

CONSTRUCTBITVEC (int[] B, int k)
{
    stack S;
    S.push(0);
    for (j = 1; j ≤ k; j++)
    {
        while (S.top() > B[j]) S.pop();
        LB[j] = S.top();
        S.push(B[j]);
    }

    for (j = 1; j ≤ k; j++)
        if (LB[j] > 0)
            V(B,j) := V(B,LB[j]) ∨ 0LB[j]-110k-LB[j]+1;
        else
            V(B,j) = 0k;
}

```

Abbildung 4.6: Algorithmus: Erstellung der Bit-Vektoren eines Blockes B

Damit erhalten wir das für uns wichtige Theorem bereits zum zweiten Mal.

Theorem 4.18 *Für das Lowest Common Ancestor Problem kann mit Hilfe einer Vorverarbeitung, die mit einem Zeitbedarf von $O(n)$ ausgeführt werden kann, jede Anfrage in konstanter Zeit beantwortet werden.*

Wir merken an dieser Stelle noch, dass wir mit dieser Methode das allgemeine RMQ-Problem und nicht das IRMQ-Problem lösen können. Daher können wir statt der Euler-Tour auch einfach die Inorder-Liste der Knoten des Baumes verwenden (die der Euler-Tour allerdings ähnlich ist).

Definition 4.19 *Sei $T = (V, E)$ ein gewurzelter Baum mit Wurzel r und seien $T_1, \dots, T_\ell$ die Teilbäume, die an den Kindern der Wurzel gewurzelt sind, mit den Inorder-Listen $I_1, \dots, I_\ell$, wobei $I_j = (v_1^{(j)}, \dots, v_{n_j}^{(j)})$. Besitzt r keine Kinder, dann ist die Inorder-Liste von T gleich (r) ; Ansonsten ist die Inorder-Liste von T gegeben durch: $(v_1^{(1)}, \dots, v_{n_1}^{(1)}, r, v_1^{(2)}, \dots, v_{n_2}^{(2)}, r, \dots, r, v_1^{(\ell)}, \dots, v_{n_\ell}^{(\ell)})$.*

Den Beweis, dass man die LCA-Anfragen analog auf RMQ-Anfragen reduzieren kann, sind dem Leser überlassen. Der Leser möge sich auch überlegen, wie lange die Inorder-Listen werden können.

Wir wollen hier auch noch anmerken, dass der Platzbedarf der hier vorgestellten Varianten für RMQ-Anfragen (und damit auch für LCA-Anfragen) $O(n)$ Maschinenwörter, also $O(n \log(n))$ Bits beträgt. Man kann mit etwas aufwendigeren Techniken zeigen, dass sich der zusätzliche Platzbedarf (zusätzlich zum betrachteten Feld, auf dem die RMQ-Anfragen gestellt werden) auf $2n + o(n)$ Bits reduzieren lässt, ohne dass man dabei auf die lineare Vorverarbeitungszeit und die konstante Anfragezeit verzichten muss. Ferner kann man zeigen, dass dieser additive Platzbedarf optimal ist. Für die Details verweisen wir auf die Originalliteratur.

14.12.06

4.2 Algorithmus von Schieber und Vishkin (*)

In diesem Abschnitt behandeln wir eine andere, ältere Variante zur Bestimmung Lowest Common Ancestors. Dieser Abschnitt ist nur als Ergänzung im Skript, da er Teil der Vorlesung im Wintersemester 2003/04 war.

In diesem Abschnitt nehmen wir explizit an, dass sich alle Operationen auf Bit-Strings der Länge $O(\log(n))$ in konstanter Zeit erledigen lassen, wobei n die Anzahl Knoten im untersuchten Baum ist. Dies haben wir bisher in allen Algorithmen auch immer implizit angenommen. Da hier aber explizit scheinbar aufwendigere Operationen auf Bit-Strings vorkommen, wollen wir dies hier explizit anmerken.

4.2.1 LCA-Queries auf vollständigen binären Bäumen

Zuerst wollen wir den einfachen Fall betrachten, dass die Anfragen in einem vollständigen binären Baum stattfinden.

Definition 4.20 *Ein gewurzelter Baum heißt binär, wenn jeder Knoten maximal zwei Kinder besitzt. Ein binärer Baum heißt vollständig, wenn alle Blätter dieselbe Tiefe besitzen und jeder innere Knoten genau zwei Kinder besitzt.*

Kommen wir nun zur Definition der Inorder-Nummerierung.

Definition 4.21 *Sei $B = (V, E)$ ein binärer Baum mit n Knoten. Sei weiter $I : V \rightarrow [1 : n]$ eine bijektive Abbildung. Wenn für jeden Knoten $v \in V$ gilt, dass $I(w) < I(v) < I(w')$ für jeden Nachfahren w des linken Kindes und jeden Nachfahren w' des rechten Kindes, dann ist I die Inorder-Nummerierung von B .*

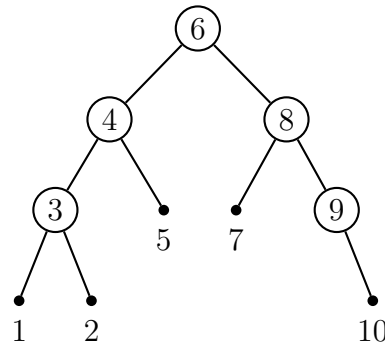


Abbildung 4.7: Beispiel: Inorder-Nummerierung eines binären Baumes

In Abbildung 4.7 ist das Beispiel einer Inorder-Nummerierung für einen beliebigen binären Baum angegeben.

Wir versuchen nun die Anfrage $\text{lca}(i, j)$ für einen vollständigen Baum zu beantworten. Wir unterscheiden zwei Fälle, je nachdem, ob die Knoten Nachfahren voneinander sind oder nicht.

Fall 1: Es gilt $\text{lca}(i, j) \in \{i, j\}$. Wir führen eine Tiefensuche aus und vergeben jedem Knoten seine DFS-Nummer. Für jeden Knoten v bestimmen wir zusätzlich das Intervall von DFS-Nummern aller Knoten, die Nachfahren von v sind. Dann gilt:

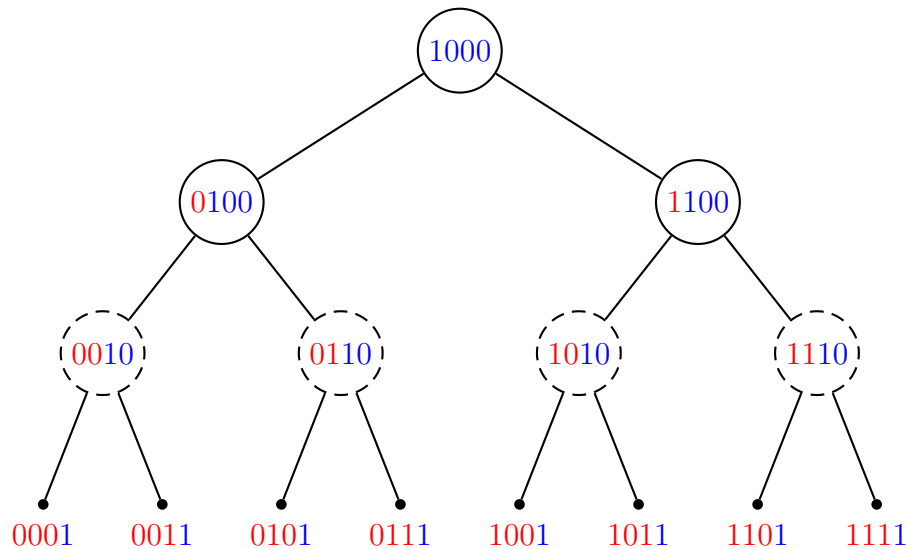
$$\text{lca}(i, j) = \begin{cases} i & \text{falls } \text{dfs}(j) \in \text{DFS-Intervall}(i), \\ j & \text{falls } \text{dfs}(i) \in \text{DFS-Intervall}(j), \\ 0 & \text{sonst.} \end{cases}$$

Wir merken noch an, dass sich die DFS-Nummern und DFS-Intervalle in linearer Zeit mit einer Tiefensuche berechnen lassen und sich dann jede solche Lowest Common Ancestor Anfrage in konstanter Zeit beantworten lässt. Bei Ergebnis 0 gehen wir natürlich in den folgenden Fall 2.

Fall 2: Es gilt $\text{lca}(i, j) \notin \{i, j\}$. Wir vergeben bei einer Tiefensuche die Inorder-Nummerierung des vollständigen binären Baumes. Eine Beispiel für eine solche Inorder-Nummerierung ist in Abbildung 4.8 angegeben. Hierbei sind die Inorder-Nummern als Bit-Strings dargestellt.

Für die Inorder-Nummerierung gibt es auch eine zweite Interpretation. Sei dazu v ein Knoten des vollständigen binären Baumes B . Ohne Beschränkung der Allgemeinheit, interpretieren wir für jeden inneren Knoten ein Kind als linkes und ein Kind als rechtes Kind. Dies können wir tun, da in einem vollständigen binären Baum jeder innere Knoten genau zwei Kinder besitzt.

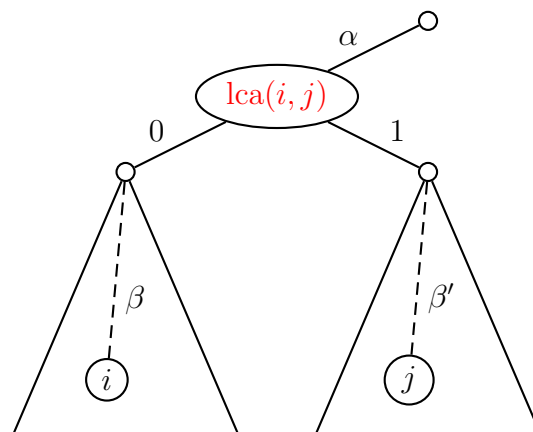
Wir betrachten die folgende Abbildung $\alpha : V(B) \rightarrow \{0, 1\}^*$ vermöge der folgenden Zuordnung. Folgen wir dem Pfad von der Wurzel zu einem Knoten v , so setzen wir



Abbildungung 4.8: Beispiel: Inorder-Nummerierung eines vollständigen binären Baumes der Tiefe 3

für jede Verzweigung zu einem linken Kind eine 0, zu einem rechten Kind eine 1. Die so konstruierte Zeichenkette $\alpha(v)$ ordnen wir dem Knoten v zu. Dabei erhält die Wurzel als Zuordnung das leere Wort ε . Dies ist in Abbildung 4.8 ebenfalls illustriert. Die roten Zeichenreihen sind dabei die dem Knoten zugeordnete Zeichenreihe.

Die Inorder-Nummerierung von v erhält man dann als $\alpha(v) \cdot 10^{d-d(v)}$, wobei d die Tiefe des vollständigen binären Baumes d ist und $d(v)$ die Tiefe des Knotens v im vollständigen binären Baum angibt. Somit haben die zugehörigen Bit-Strings (mit führenden Nullen) eine Länge von $d + 1$. In Abbildung 4.8 sind die angehängten Bit-Strings aus 10^* blau dargestellt. Es sei dem Leser zum Beweis überlassen, dass die obige Beschreibung der Inorder-Nummerierung korrekt ist.



Abbildungung 4.9: Skizze: Lowest Common Ancestor von i und j

Wie findet man jetzt mit Hilfe dieser Inorder-Nummerierung $\text{lca}(i, j) \notin \{i, j\}$? Wir betrachten dazu die Skizze in Abbildung 4.9. Wir bemerken, dass dann der Pfad von der Wurzel des vollständigen binären Baumes zu i bzw. j bis zum Knoten $\text{lca}(i, j)$ identisch ist. Sei die zugeordnete Zeichenreihe gerade α . Also beginnt die Inorder-Nummerierung von i und j jeweils mit α . Da $\text{lca}(i, j) \notin \{i, j\}$ gilt, müssen sich die Inorder-Nummern von i und j in der darauf folgenden Bit-Position unterscheiden.

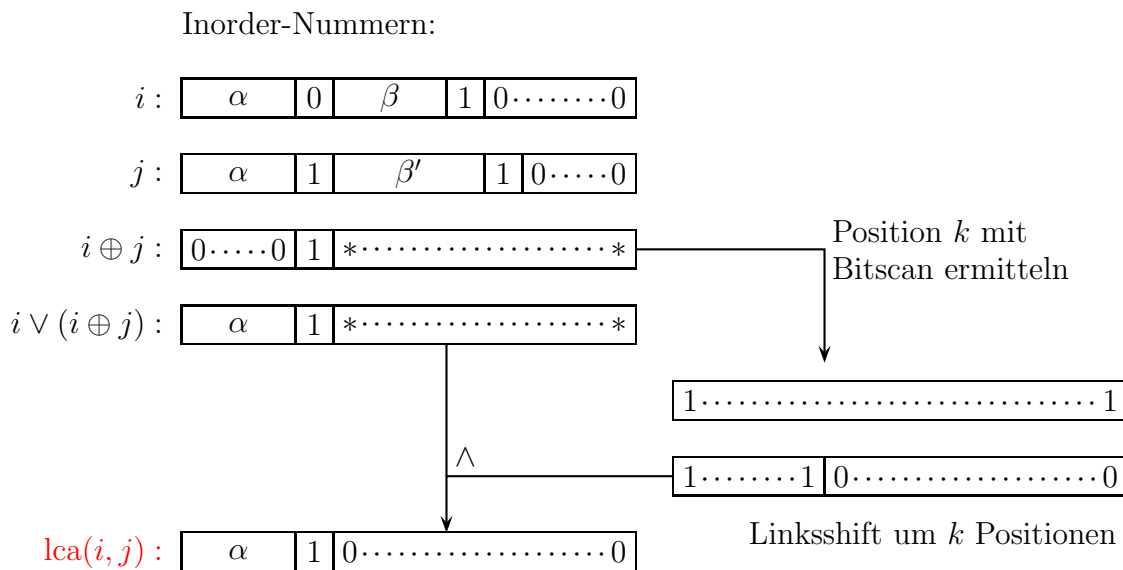


Abbildung 4.10: Skizze: Inorder-Nummern von i , j und $\text{lca}(i, j)$

Betrachten wir jetzt die Inorder-Nummern als Bit-Strings von i , j und $\text{lca}(i, j)$ wie in Abbildung 4.10. Dann können wir den Bereich von α aus $i \oplus j$ als den führenden 0er Block herauslesen ($\oplus$ bezeichnete hier das exklusive Oder). Mit Hilfe eines Bitscans lässt sich die Position der ersten 1 von links (und somit im Wesentlichen die Länge von α) aus $i \oplus j$ bestimmen. Der Bit-String $i \vee (i \oplus j)$ liefert schon im Wesentlichen den Bit-String für den $\text{lca}(i, j)$. Wir müssen nur noch die letzten Bit-Positionen mittels eines logischen Unds auf 0 setzen.

Auch hier gilt wieder, dass sich die Inorder-Nummerierung mit einer Tiefensuche in linearer Zeit konstruieren kann und dass sich jede anschließende Lowest Common Ancestor Anfrage in konstanter Zeit beantworten lässt. Man beachte, dass diese Methode nur für $\text{lca}(i, j) \notin \{i, j\}$ funktioniert.

Kombiniert man die beiden Methoden in den vorher diskutierten Fällen, erhalten wir als Ergebnis das folgenden Lemma.

Lemma 4.22 *Sei B ein vollständiger binärer Baum. Nach einer Vorverarbeitung, die in linearer Zeit ausgeführt werden kann, kann jede Lowest Common Ancestor Query $\text{lca}_B(i, j)$ in konstanter Zeit beantwortet werden.*

Wir wollen hier noch anmerken, dass Bitscans in modernen Prozessoren ein Standardfunktion (z.B. `bsf` und `bsr`), die in konstanter Zeit ausgeführt werden kann. Man überlegen sich, dass eine Implementierung mit Bitkomplexität $O(\log\log(n))$ möglich ist, also nicht aufwendiger ist, als die Addition von zwei Binärzahlen.

4.2.2 LCA-Queries auf beliebigen Bäumen

Wenden wir uns jetzt dem Fall zu, dass wir Lowest Common Ancestor Anfragen in einem beliebigen Baum beantworten wollen. Ohne Beschränkung der Allgemeinheit identifizieren wir im Folgenden die Knoten eines beliebigen Baumes durch seine DFS-Nummer. Ein Beispiel ist in Abbildung 4.11 angegeben.

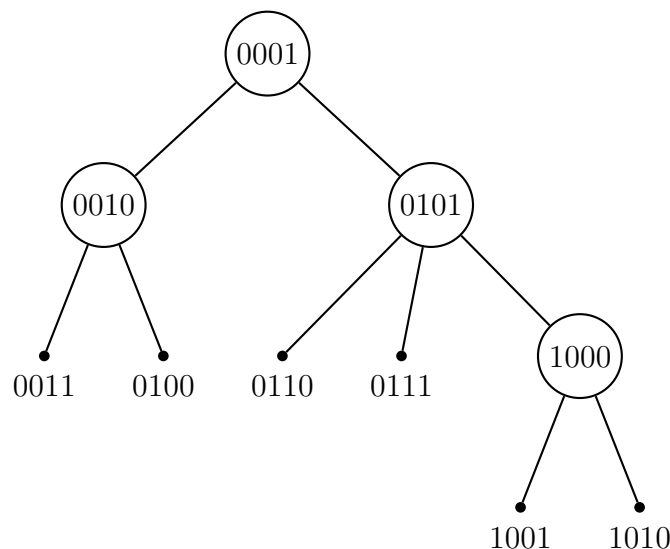


Abbildung 4.11: Beispiel: Ein beliebiger Baum mit einer DFS-Nummerierung seiner Knoten

Notation 4.23 Sei $k \in \mathbb{N}$. Dann ist $h(k)$ die kleinste Bitposition (von rechts), so dass in der Binärdarstellung von k an dieser Position eine 1 steht.

Man beachte, dass hierbei die rechteste Bitposition den Index 1 besitzt. Es gilt beispielsweise: $h(8) = 4$, $h(6) = 2$

Definition 4.24 In einem Baum ist die Höhe eines Knotens v die maximale Anzahl der Knoten eines längsten einfachen Pfades von v zu einem Blatt im Teilbaum von v .

Beobachtung 4.25 Sei v ein Knoten im vollständigen binären Baum, dann ist seine Höhe gleich $h(\text{inorder}(v))$, wobei $\text{inorder}(v)$ die Inorder-Nummer von v ist.

Notation 4.26 Sei v ein Knoten in einem beliebigen Baum, dann ist $I(v)$ der Knoten w im Teilbaum von v , dessen Wert $h(w)$ maximal ist.

Beobachtung 4.27 Ist w ein Nachfahre von v , dann gilt $h(I(w)) \leq h(I(v))$.

In Abbildung 4.12 sind für einen beliebigen Baum diejenigen Knoten mit gleicher Farbe gekennzeichnet, die denselben Knoten $I(v)$ besitzen. Die schwarzen Knoten sind dabei jeweils als einelementige Mengen zu verstehen.

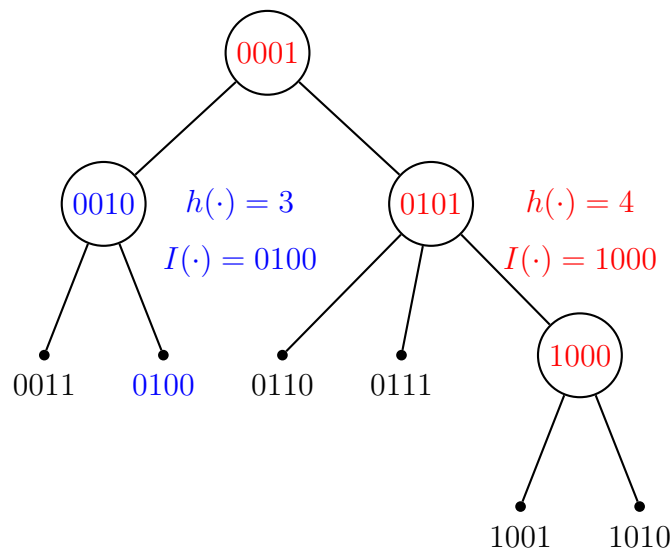


Abbildung 4.12: Beispiel: Knoten im Baum mit gleichem $I(v)$ -Wert

Lemma 4.28 Für jeden Knoten $v \in V(T)$ existiert ein eindeutiger Knoten w in $T(v)$, so dass $h(w)$ maximal ist.

Beweis: Für einen Widerspruchsbeweis nehmen wir an, dass $w \neq w' \in V(T(v))$ mit $h(w) = h(w')$ existieren, so dass für alle $u \in V(T(v))$ gilt $h(u) \leq h(w)$. Dann lassen sich w und w' wie folgt mit $\alpha, \alpha', \beta \in \{0, 1\}^*$ und $k = h(w) - 1 = h(w') - 1$ schreiben:

$$\begin{aligned} w &= \alpha \cdot 0 \cdot \beta \cdot 10^k, \\ w' &= \alpha' \cdot 1 \cdot \beta \cdot 10^k. \end{aligned}$$

Da w und w' DFS-Nummern sind und sich im Teilbaum von $T(v)$ befinden, muss sich der Knoten $u := \alpha 10^{|\beta|+k+1}$ ebenfalls im Teilbaum $T(v)$ befinden, da $u \in [w : w']$. Da offensichtlich $h(u) > h(w) = h(w')$ gilt, erhalten wir den gewünschten Widerspruch. ■

Damit haben wir auch gezeigt, dass die Abbildung $v \mapsto I(v)$ wohldefiniert ist.

Korollar 4.29 *Sei $v \in V(T)$ und seien w und w' zwei Kinder von v . Dann gilt $|\{h(I(v)), h(I(w)), h(I(w'))\}| \geq 2$.*

Beweis: Für einen Widerspruchsbeweis sei $h(I(v)) = h(I(w)) = h(I(w'))$. Sei weiter $u = I(w)$ und $u' = I(w')$. Nach Definition von I gilt $u \neq u'$. Siehe hierfür auch die Skizze in Abbildung 4.13.

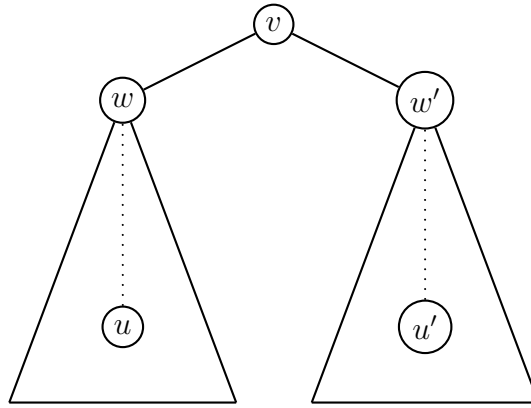


Abbildung 4.13: Skizze: Zum Beweis von Korollar 4.29

Gilt $h(u) = h(u')$, dann muss $h(I(v)) > h(I(w)) = h(I(w'))$ gelten, sonst hätten wir einen Widerspruch zu Lemma 4.28. Also erhalten wir den gewünschten Widerspruch zu $h(I(v)) = h(I(w)) = h(I(w'))$. ■

Korollar 4.30 *Die Partition der Knoten von T in Mengen mit gleichen $I(v)$ -Wert ist eine Partition in Pfade von T .*

Beweis: Aus dem vorhergehenden Korollare folgt, dass nur höchstens ein Kind denselben I -Wert besitzen kann wie sein Elter. ■

Definition 4.31 *Die Pfade in der durch $I(\cdot)$ induzierten Partition heißen Runs.*

Aus Beobachtung 4.27 folgt unmittelbar die folgende Beobachtung:

Beobachtung 4.32 Für jeden Knoten $v \in V(T)$ ist $I(v)$, der tiefste Knoten des Runs, der v enthält.

Definition 4.33 Der höchste Knoten eines Runs heißt Kopf des Runs.

Wir werden im Folgenden mit Hilfe der Funktion $I(\cdot)$ die Knoten aus dem beliebigen Baum T auf einen vollständigen binären Baum abbilden. Mit Hilfe der dort schon bekannten Methode wollen wir dann die Lowest Common Ancestor Anfragen in einem beliebigen Baum T auf Lowest Common Ancestor Anfragen in einem vollständigen binären Baum B reduzieren. Dazu betrachten wir die folgende Abbildung:

$$I : V(T) \rightarrow V(B) : v \mapsto I(v).$$

Hierbei ist v eine DFS-Nummer in T und $I(v)$ wird als der Knoten in B mit Inorder-Nummer $I(v)$ interpretiert. Der vollständige binäre Baum hat dann eine Tiefe von $\lfloor \log(n) \rfloor$, wobei n die Anzahl der Knoten in T bezeichnet. Man beachte, dass die Abbildung I weder injektiv noch surjektiv sein muss. Diese Idee der Abbildung ist in der Skizze in Abbildung 4.14 illustriert.

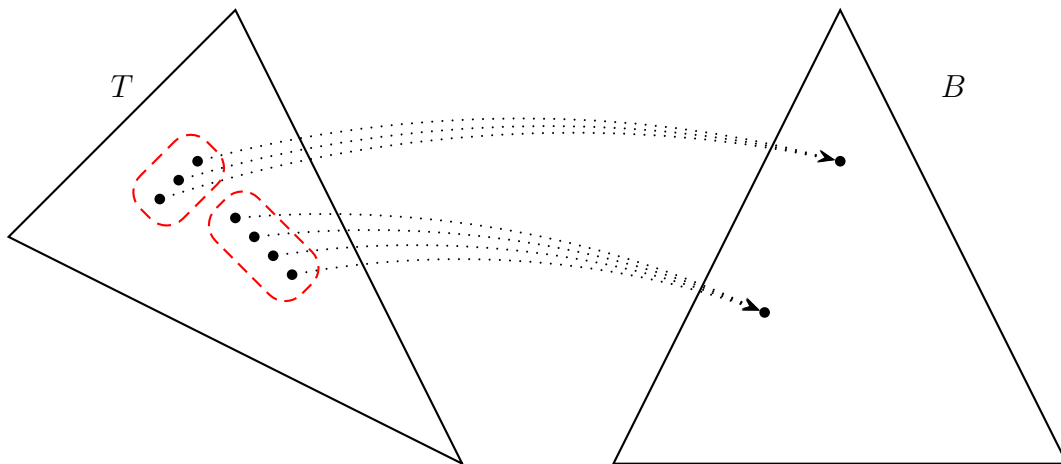
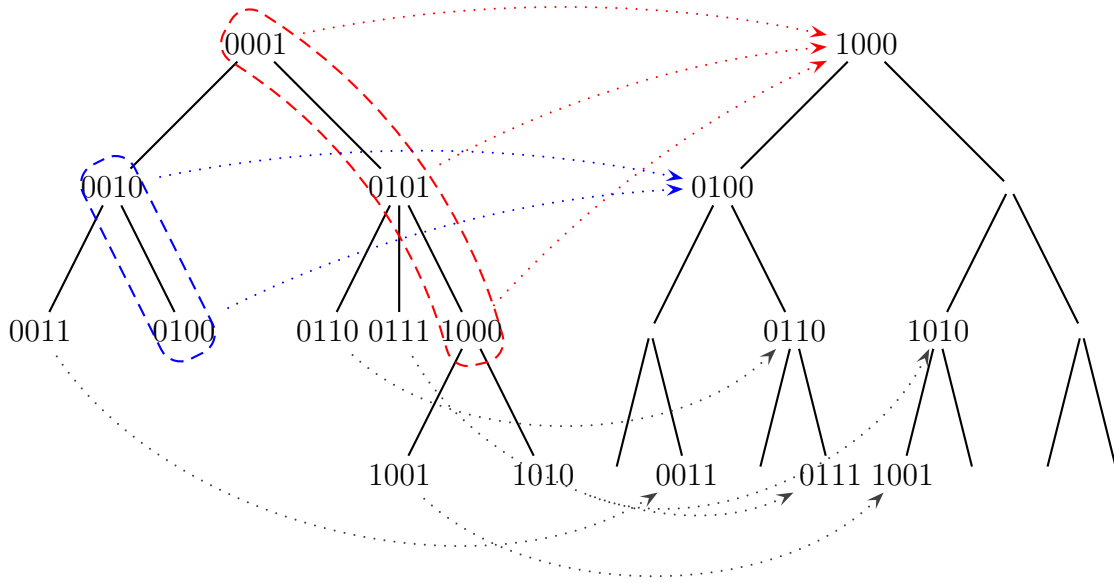


Abbildung 4.14: Skizze: Die Abbildung I eines beliebigen Baumes in einen vollständigen binären Baum

Diese Abbildung I ist in Abbildung 4.15 noch einmal anhand eines expliziten Beispiels dargestellt.

Abbildung 4.15: Beispiel: Die Abbildung I

4.2.3 Vorverarbeitung

Wir geben in diesem Abschnitt den Algorithmus für die Vorverarbeitung für die Lowest Common Ancestor Queries an. Der Algorithmus selbst ist in Abbildung 4.16 angegeben. Hierfür benötigen wir aber erst noch die folgende Notation.

Notation 4.34 Sei T ein beliebiger Baum mit n Knoten. Für jeden Knoten $v \in V(T)$ ist A_v ein Bit-Feld der Länge $\lfloor \log(n) \rfloor + 1$. Dabei ist $A_v(i)$ genau dann gleich 1, wenn der Knoten v einen Vorfahren in T besitzt, der auf einen Knoten in Höhe i im Baum B abgebildet wird. Ansonsten ist $A_v(i) = 0$. Kurz es gilt:

$$A_v(i) = \begin{cases} 1 & \text{falls ein Vorfahre } u \text{ von } v \text{ in } T \text{ existiert mit } h(I(u)) = i, \\ 0 & \text{sonst.} \end{cases}$$

Für die Definition von A_v betrachten wir noch folgenden Beispiel für $v = (0011)_2 = 3$ aus Abbildung 4.15. Der Knoten $v = 3$ hat dann zwei Vorfahren, nämlich die Runs mit I -Wert $4 = (0100)_2$ und mit I -Wert $8 = (1000)_2$. Somit ist $A_3(4) = A_3(3) = 1$. Weiter gilt $A_3(1) = 1$, da der Knoten v ja ein Vorfahre von sich selbst ist. Da v jedoch keinen Vorfahren u mit $h(I(u)) = 2$ besitzt, ist $A_3(2) = 0$.

Schauen wir uns jetzt an, ob die in Abbildung 4.16 angegebene Vorverarbeitung effizient berechnet werden kann. Die DFS-Nummerierung lässt sich offensichtlich während einer Tiefensuche nebenbei mitberechnen. Der Wert $h(v)$ für jeden Knoten

PREPROCESSING_LCA (tree T)

```

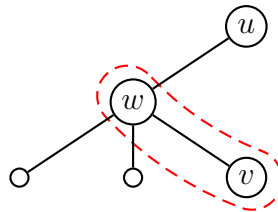
{
  for each ( $v \in V(T)$ ) do /* using DFS */
  {
    /* during descend */
    determine  $\text{dfs}(v)$ ; /* dfs-number of  $v$  */
    /* during ascend */
    determine  $h(v)$  and  $I(v)$ ;
    determine  $L(v) = w$ , where  $w$  is the head of a run containing  $v$ ;
  }
  for each ( $v \in V(T)$ ) do /* using DFS */
    determine  $A_v$ ;
}

```

Abbildung 4.16: Algorithmus: Vorverarbeitung zur LCA-Query

lässt sich ebenfalls pro Knoten in konstanter Zeit mit geeigneten Operationen auf Bit-Strings berechnen. Der Wert $I(v)$ lässt sich ebenfalls leicht berechnen, da nur für den Knoten und seinen Kinder derjenige Knoten bestimmt werden muss, der den maximalen h -Wert besitzt.

Für die Ermittlung des Feldes L hilft die folgende Beobachtung: Für jeden Knoten gibt es maximal ein Kind, das zu einem Run gehört, dessen Kopf noch unbekannt ist. Wir müssen also bei der Rückkehr aus der Tiefensuche nur maximal eine Liste von Knoten halten, deren Kopf noch unbekannt ist. Da der Knoten w genau dann ein Kopf eines Runs ist, wenn für seinen Elter u gilt $I(u) \neq I(w)$, lässt sich auch leicht feststellen, wann ein Kind zum Kopf eines Runs wird. Insgesamt kann das Feld L in linearer Zeit bestimmt werden. Diese Situation ist in Abbildung 4.17 skizziert.

Abbildung 4.17: Skizze: Knoten w mit maximal einem Kind, für den der Kopf seines Runs unbekannt ist

Es stellt sich die Frage, wie sich Bit-Strings A_v effizient berechnen lassen. Wir werden diese Bit-Strings während einer Tiefensuche berechnen. Beim erstmaligen Aufsuchen eines Knotens v übernehmen wir den Bit-String des Elters (bei der Wurzel den Bit-String aus lauter Nullen) und setzen eigenes Bit an Position $h(I(v))$. Beim endgültigen

Verlassen des Knotens v löschen wir das beim Betreten gesetzte Bit an Position $h(I(v))$. Damit wird A_v für jeden Knoten v nach der Definition korrekt berechnet.

Lemma 4.35 *Sei T ein beliebiger Baum. Die Vorverarbeitung für Lowest Common Ancestor Anfragen in T kann in linearer Zeit implementiert werden.*

Wir werden in den folgenden Abschnitten noch sehen, wozu wir die in der Vorverarbeitung berechneten Werte benötigen.

4.2.4 Beziehung zwischen Vorfahren in T und B

In diesem Abschnitt beweisen wir eine im Folgenden wichtige elementare Beziehung zwischen der Vorfahr-Relation in T und B .

Lemma 4.36 *Wenn z ein Vorfahre von x in T ist, dann ist $I(z)$ auch ein Vorfahre von $I(x)$ in B :*

Beweis: Wir unterscheiden zwei Fälle, je nachdem, ob $I(x) = I(z)$ gilt oder nicht.

Fall 1: Sei $I(x) = I(z)$. Da x und z auf denselben Knoten abgebildet werden, ist die Behauptung trivial.

Fall 2: Sei jetzt $I(x) \neq I(z)$. Da der Knoten z ein Vorfahre von x ist, gilt also aufgrund von Beobachtung 4.27, dass $h(I(z)) \geq h(I(x))$. Wäre $h(I(z)) = h(I(x))$, dann müsste $I(x) = I(z)$ sein, da der Knoten mit maximalen h -Wert im Teilbaum von z eindeutig ist. Also erhalten wir einen Widerspruch und es muss $h(I(z)) > h(I(x))$ gelten.

Wir betrachten jetzt die Bit-Strings von $I(z)$ und $I(x)$ wie in Abbildung 4.18. Dabei nehmen wir an, dass sich $I(z)$ und $I(x)$ an der Position k (gezählt von rechts) zum ersten Mal unterscheiden, wenn wir die Bit-Strings von links vergleichen.

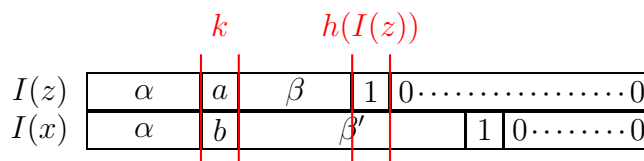


Abbildung 4.18: Skizze: Die Bit-Strings von $I(z)$ und $I(x)$

Wir zeigen jetzt, dass dann $k \leq h(I(z))$ gilt. Für einen Widerspruchsbeweis nehmen wir an, dass $k > h(I(z))$ gilt. Nach Voraussetzung gilt $I(z) \neq I(x)$. Wir unterscheiden jetzt noch zwei Unterfälle, je nachdem, ob $I(z) > I(x)$ oder $I(z) < I(x)$ gilt.

Fall 2a: Es gilt $I(z) > I(x)$. Wir betrachten jetzt den Wert $N := \alpha 10 \dots 0$. Dann gilt $N \in (I(x) : I(z))$ und N wird in der Tiefensuche von z aus gefunden. Dann wäre aber wieder $h(N) > h(I(z))$, was nicht sein kann.

Fall 2b: Es gilt jetzt $I(z) < I(x)$. Eine analoge Beweisführung liefert denselben Widerspruch.

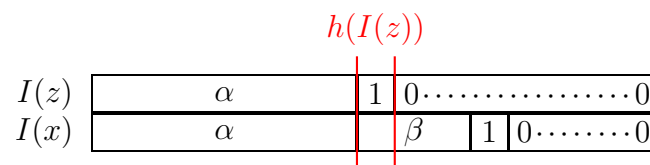


Abbildung 4.19: Skizze: Die Bit-Strings $I(z)$ und $I(x)$ im Falle $k \leq h(I(z))$

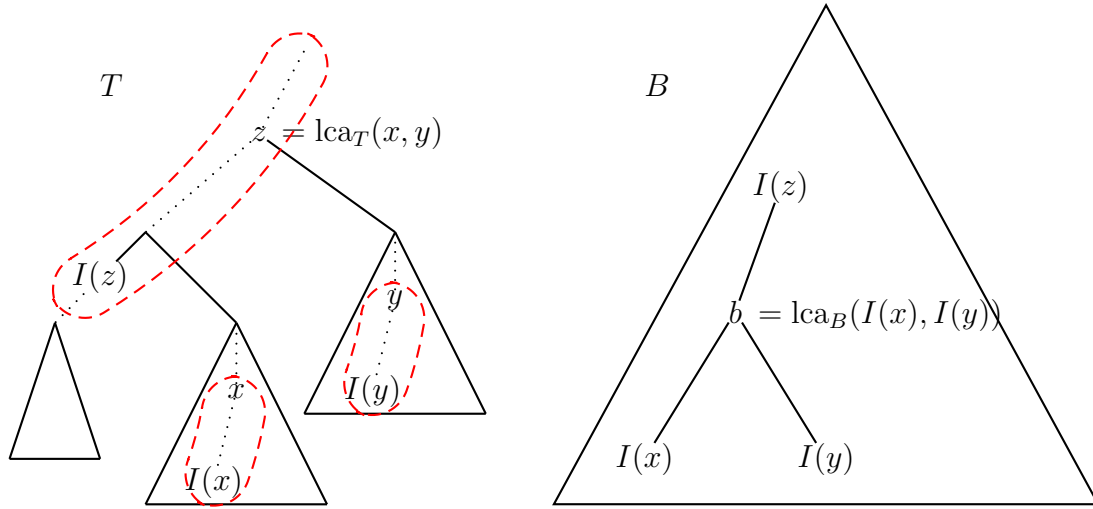
Wir erhalten also in jedem Fall einen Widerspruch und es muss daher $k \leq h(I(z))$ gelten. Also muss für die Bit-Strings von $I(z)$ und $I(x)$ die in Abbildung 4.19 angegebene Darstellung gelten. Das bedeutet aber nach unserer zweiten Interpretation der Inorder-Nummerierung des vollständigen binären Baumes, dass $I(z)$ ein Vorfahre von $I(x)$ ist. ■

4.2.5 Berechnung einer LCA-Query

Nun überlegen wir uns, wie wir eine Lowest Common Ancestor Query $\text{lca}_T(x, y)$ für den Baum T verarbeiten. Dazu werden wir für x bzw. y die Werte $I(x)$ und $I(y)$ bestimmen. Wie wir gesehen haben, gilt für den niedrigsten gemeinsamen Vorfahren $z = \text{lca}_T(x, y)$, dass $I(z)$ auch ein Vorfahre von $I(x)$ und $I(y)$ im vollständigen binären Baum B ist. Wir bestimmen also $b = \text{lca}_B(I(x), I(y))$ und hoffen, dass wir aus b den Knoten z (bzw. seine DFS-Nummer) rekonstruieren können. Diese Idee ist in Abbildung 4.20 noch einmal bildlich dargestellt.

Der Algorithmus zur Beantwortung von $z = \text{lca}_B(x, y)$ geht dann wie folgt vor:

1. Bestimme $I(x)$ und $I(y)$.
2. Bestimme $b = \text{lca}_B(I(x), I(y))$.
3. Bestimme $h(I(z))$ aus b mit Hilfe des Wertes $h(b)$.
4. Bestimme z aus $I(z)$.

Abbildung 4.20: Skizze: Reduktion einer LCA-Anfrage in T auf eine in B

Wir wissen, dass $I(z)$ ein Vorfahre von b ist und wir müssen uns nur den richtigen Vorfahren von b und dessen korrektes Urbild auffinden. Dabei hilft uns das folgende Lemma.

Lemma 4.37 *Sei T ein beliebiger Baum mit n Knoten und seien $x, y, z \in V(T)$ mit $z = \text{lca}_T(x, y)$. Sei B ein vollständiger binärer Baum mit Tiefe $\lfloor \log(n) \rfloor$ und sei $b = \text{lca}_B(I(x), I(y))$. Weiter sei $j = \min \{k \geq h(b) : A_x(k) = A_y(k) = 1\}$. Dann gilt $h(I(z)) = j$.*

Beweis: Sei $k := h(I(z))$. Da z ein Vorfahre von x und y in Baum T ist, gilt $A_x(k) = A_y(k) = 1$ nach Definition von A . Da $I(z)$ ein Vorfahre von $I(x)$ und $I(y)$ ist und somit auch von $b = \text{lca}_B(I(x), I(y))$, gilt also $j \leq k = h(I(z))$.

Da $A_x(j) = 1$, muss ein Vorfahre x' von x in T existieren, so dass $I(x')$ ein Vorfahre von $I(x)$ in B ist und dass $I(x')$ die Höhe $j \geq h(b)$ in B hat. Also ist $I(x')$ ein Vorfahre von b in B auf Höhe j . Analog gibt es einen Vorfahren y' von y in T , so dass $I(y')$ ein Vorfahre von b in B auf Höhe j ist. Also muss $I(x') = I(y')$ gelten. Somit gehören x' und y' zum selben Run in T .

Ohne Beschränkung der Allgemeinheit sei x' ein Vorfahre von y' . Dann ist x' ein Vorfahre von x und y . Da $z = \text{lca}_T(x, y)$, muss auch x' ein Vorfahre von z sein. Dann gilt auch $j = h(I(x')) \geq h(I(z)) = k$.

Somit gilt also insgesamt $j = k = h(I(z))$ und die Behauptung ist gezeigt. ■

Lemma 4.38 *Sei T ein beliebiger Baum und $x, y, z \in V(T)$. Sei $z = \text{lca}_T(x, y)$ und sei $h(I(z))$ bekannt, dann kann z in konstanter Zeit ermittelt werden.*

Beweis: Sei $\bar{x}$ bzw. $\bar{y}$ der erste Knoten auf dem Pfad von x bzw. y zu z , der im Run von z liegt. Es gilt dann $z = \bar{x}$ oder $z = \bar{y}$ und somit gilt $z = \min(\bar{x}, \bar{y})$. Um z zu ermitteln, benötigen wir also nur die Knoten $\bar{x}$ und $\bar{y}$. Siehe dazu auch die Abbildung 4.21.

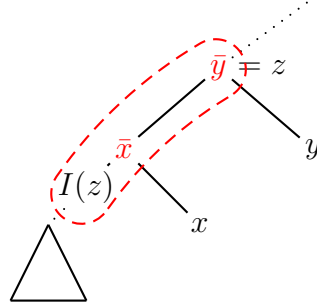


Abbildung 4.21: Skizze: Ermittlung von $z = \text{lca}_T(x, y)$

Wie findet man jedoch $\bar{x}$ bzw. $\bar{y}$? Sei dazu $j = h(I(z))$. Ist $h(I(x)) = j$, dann befinden sich x und z im selben Run. Also gilt $\bar{x} = x$. Analoges gilt natürlich auch für y .

Betrachten wir also den Fall $\bar{x} \neq x$. Sei w der erste Knoten von $\bar{x}$ zu x der nicht im Run von $\bar{x}$ liegt. Dann gilt $h(I(x)) \leq h(I(w)) < h(I(\bar{x})) = h(I(z))$. Diese Situation ist in Abbildung 4.22 skizziert.

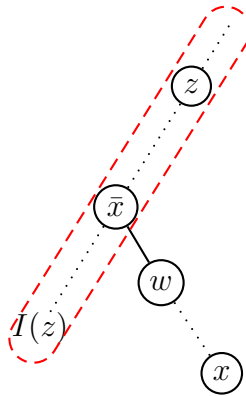


Abbildung 4.22: Skizze: Definition des Knotens w

Weiter gilt, dass $h(I(w))$ auf dem Pfad von x nach w sogar maximal ist. Dann ist $h(I(w)) = \max \{k < j : A_x(k) = 1\}$. Damit lässt sich $h(I(w))$ mit konstant vielen Operationen auf Bit-Strings berechnen. Zuerst müssen alle Bits ab Position j in A_x mittels einer geeigneten Verundung auf Null gesetzt werden. Mit einem anschließenden Bitscan kann dann die linkeste 1 gefunden werden.

Dann x ein Nachfahre von w ist, muss auch $I(x)$ ein Nachfahre von $I(w)$ in B sein. Nach der Inorder-Nummerierung in B ergibt sich die folgende Darstellung der Bit-Strings von $I(w)$ und $I(x)$ wie in Abbildung 4.23 angegeben.

$$\begin{array}{c}
 \textcolor{red}{h(I(w))} \\
 \begin{array}{|c|c|c|c|}
 \hline
 I(w) & \alpha & 1 & 0 \dots \dots \dots 0 \\
 \hline
 I(x) & \alpha & \beta & 1 \mid 0 \dots \dots \dots 0 \\
 \hline
 \end{array}
 \end{array}$$
Abbildung 4.23: Skizze: Bit-Strings von $I(w)$ und $I(x)$

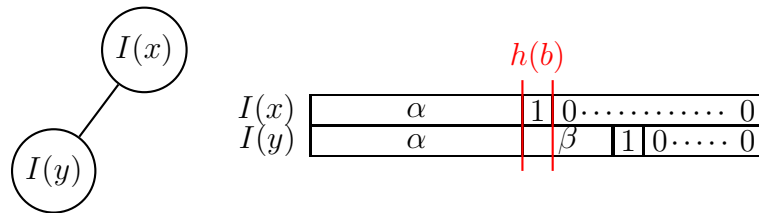
Aus $I(x)$ und $h(I(w))$ kann $I(w)$ leicht mit konstant vielen Operationen auf Bit-Strings berechnet werden. Weiterhin kann der Knoten w in T aus $I(w)$ mit Hilfe des Feldes L bestimmt werden, da ja nach Konstruktion w der Kopf eines Runs ist. Es gilt also $w = L(I(w))$ und $\bar{x} = \text{parent}(w)$ in T und die Behauptung ist bewiesen. ■

4.2.6 LCA-Query-Algorithmus

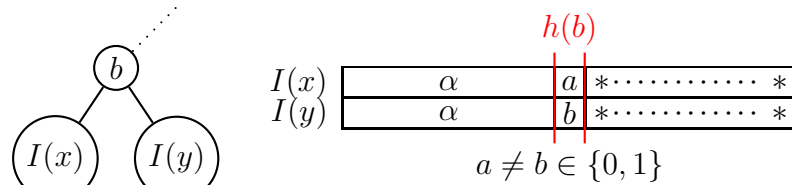
Wir wollen nun noch den vollständigen Algorithmus zu Berechnung der Anfrage $z = \text{lca}_T(x, y)$ in Abbildung 4.27 angeben.

Weiterhin wollen wir festhalten, dass der Baum B nur zu Illustration des Verfahrens dient, aber dessen Konstruktion nicht die Auswertung der Anfrage notwendig ist. Wie findet man also $h(b)$ ohne den Baum B zu konstruieren? Wir unterscheiden wieder zwei Fälle, je nachdem, ob $\text{lca}_B(I(x), I(y)) \in \{I(x), I(y)\}$ gilt oder nicht.

Fall 1: Es gilt $\text{lca}_B(I(x), I(y)) \in \{I(x), I(y)\}$. Diese Situation ist in Abbildung 4.24 dargestellt. Man sieht leicht, dass sich $h(b)$ mit einer konstanten Anzahl von Operationen auf Bit-Strings berechnen lässt.

Abbildung 4.24: Skizze: Bestimmung von $h(b)$ im ersten Fall

Fall 2: Es gilt $\text{lca}_B(I(x), I(y)) \notin \{I(x), I(y)\}$. Diese Situation ist in Abbildung 4.25 dargestellt. Man sieht leicht, dass sich $h(b)$ mit einer konstanten Anzahl von Operationen auf Bit-Strings berechnen lässt.

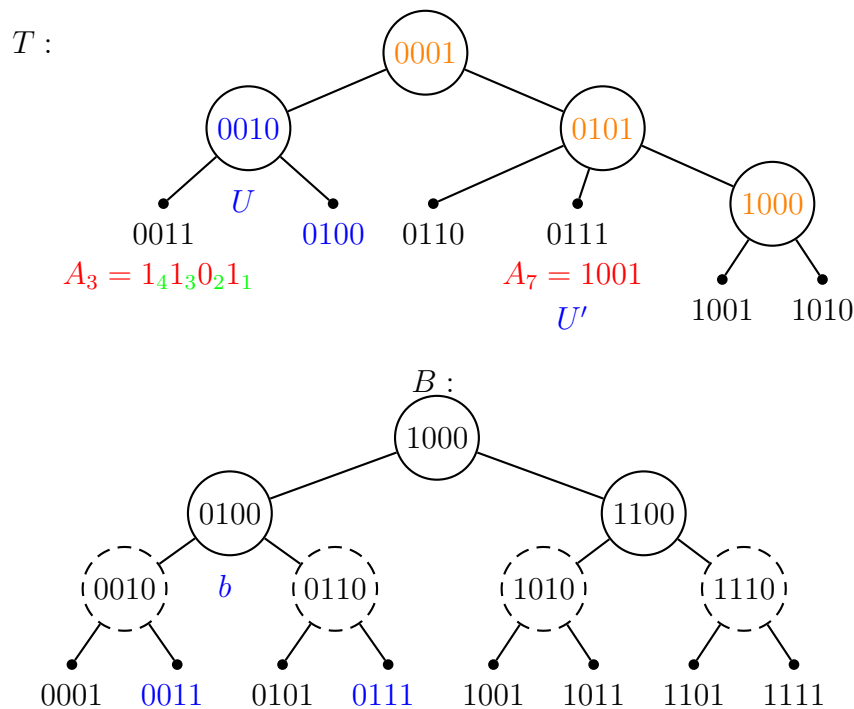
Abbildung 4.25: Skizze: Bestimmung von $h(b)$ im zweiten Fall

Halten wir das fundamentale Ergebnis dieses Abschnittes im folgenden Satz fest.

Theorem 4.39 *Mit geeigneten Preprocessing, das Zeit $O(n)$ benötigt, kann jede LCA-Query in einem Baum mit n Knoten in konstanter Zeit ausgeführt werden.*

Zum Abschluss geben wir noch ein Beispiel an. Wir wollen den niedrigsten gemeinsamen Vorfahren der Knoten 3 und 7 in Abbildung 4.26 bestimmen.

1. Es gilt $b = 0100$ und damit $h(b) = 3$. Man kann $h(b)$ auch direkt aus 0011 und 0111 ablesen, da an Position 3 zum ersten Mal von links die beiden Bit-Strings differieren.
2. Es gilt $h(I(z)) = j = 4$ nach Lemma 4.37

Abbildung 4.26: Beispiel: Die Anfrage $\text{lca}(3, 7) = \text{lca}(0011, 0111)$

3. Der Algorithmus in Abbildung 4.27 berechnet folgende Wert:

$x = 0011; \ell = 1; k = 2; u = 00|10; L(u) = 0010; \bar{x} = 0001;$

$y = 0111; \ell = 1; k = 1; u' = 011|1; L(u) = 0111; \bar{y} = 0101;$

Dann gilt $z = \min(\bar{x}, \bar{y}) = 0001$.

LCA (tree T , node x, y)

```
{
   $b = \text{lca}_B(I(x), I(y));$ 
   $j = \min \{k \geq h(b) : A_x(k) = A_y(k) = 1\};$ 
   $\ell = \min \{i : A_x(i) = 1 \text{ ist linkeste } 1\};$ 
  if ( $\ell = j$ )
     $\bar{x} = x;$ 
  else
  {
     $k := \max \{i < j : A_x(i) = 1\}$ 
     $u := \text{left}(I(x), m, k + 1) \cdot 1 \cdot 0 \cdots 0; \quad /* = I(w) */$ 
     $w = L(u);$ 
     $\bar{x} = \text{parent}(w);$ 
  }
   $\ell = \min \{i : A_y(i) = 1 \text{ ist linkeste } 1\};$ 
  if ( $\ell = j$ )
     $\bar{y} = y;$ 
  else
  {
     $k := \max \{i < j : A_y(i) = 1\}$ 
     $u := \text{left}(I(y), m, k + 1) \cdot 1 \cdot 0 \cdots 0; \quad /* = I(w) */$ 
     $w = L(u);$ 
     $\bar{y} = \text{parent}(w);$ 
  }
   $z = \min\{\bar{x}, \bar{y}\}$ 
  return  $z;$ 
}
```

Abbildung 4.27: Algorithmus: LCA-Query $\text{lca}(x, y)$

5.1 Grundlegende Eigenschaften von Suffix-Arrays

In diesem Abschnitt wollen wir auf eine andere, zum Suffix-Baum ähnliche Datenstruktur, dem Suffix-Array, eingehen. Hauptmerkmal des Suffix-Arrays gegenüber dem Suffix-Baum ist seine speicherplatzsparendere Darstellung.

5.1.1 Was sind Suffix-Arrays?

Zunächst einmal müssen wir uns überlegen, was Suffix-Arrays überhaupt sind.

Definition 5.1 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Ein Suffix-Array für t ist ein Feld der Länge $n+1$, welche die Startpositionen aller Suffixe von $t\$$ in geordneter (lexikographischer) Reihenfolge enthält. Dabei gilt $\$ < a \in \Sigma$ mit $\$ \notin \Sigma$.

Wir geben ein Beispiel für $t = \text{MISSISSIPPI}$ an.

$A[0]$	=	12	$\hat{=}$	$\$$
$A[1]$	=	11	$\hat{=}$	I $\$$
$A[2]$	=	08	$\hat{=}$	IPPI $\$$
$A[3]$	=	05	$\hat{=}$	ISSIPPI $\$$
$A[4]$	=	02	$\hat{=}$	ISSISSIPPI $\$$
$A[5]$	=	01	$\hat{=}$	MISSISSIPPI $\$$
$A[6]$	=	10	$\hat{=}$	PI $\$$
$A[7]$	=	09	$\hat{=}$	PPI $\$$
$A[8]$	=	07	$\hat{=}$	SIPPI $\$$
$A[9]$	=	04	$\hat{=}$	SISSIPPI $\$$
$A[10]$	=	06	$\hat{=}$	SSIPPI $\$$
$A[11]$	=	02	$\hat{=}$	SSISSIPPI $\$$

Man beachte, dass der erste Eintrag immer $n+1 \hat{=}$ $\$$ ist. Man könnte die Erweiterung um das Endzeichen $\$$ auch weglassen, wenn man explizit vereinbart, dass ein Präfix

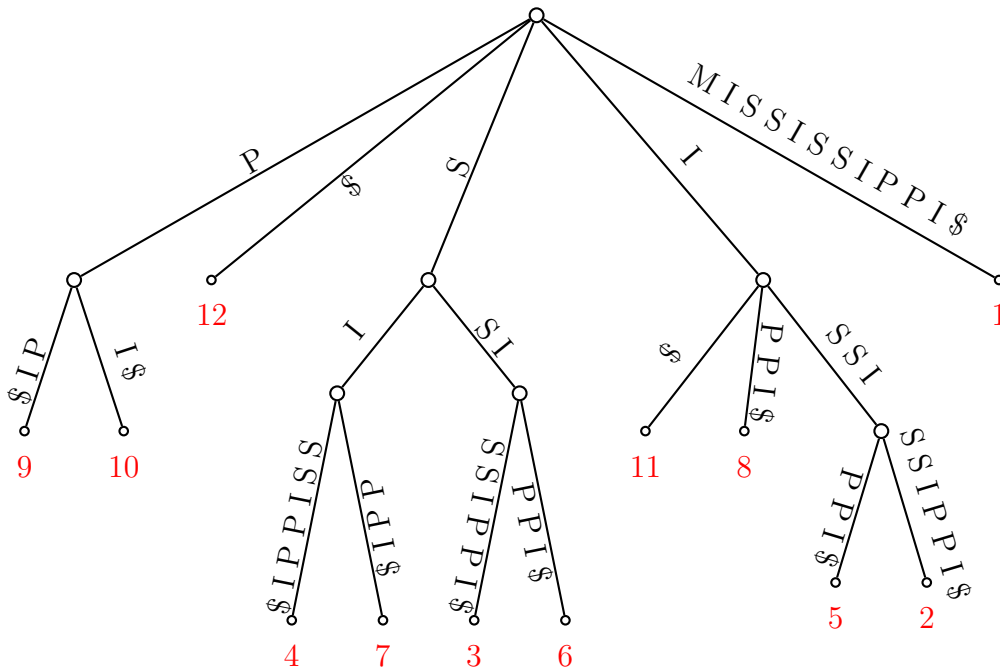


Abbildung 5.1: Beispiel: Suffix-Baum für MISSISSIPPI

eines anderen Wortes immer kleiner als das andere Wort ist. Oft ist dies in der Implementierung nur durch Beachtung von Sonderfällen zu handhaben, so dass die Ergänzung um das Endzeichen \$ oft hilfreich ist. Zum Vergleich ist in Abbildung 5.1 noch einmal der Suffix-Baum zu $t = \text{MISSISSIPPI}$ angegeben. Im Folgenden werden wir uns hauptsächlich damit beschäftigen, wie man solchen Suffix-Arrays effizient konstruiert und wie man ihn als Ersatz für Suffix-Bäume verwenden kann.

5.1.2 Konstruktion aus Suffix-Bäumen

Ein erster einfacher Ansatz zur Konstruktion eines Suffix-Arrays ist die Ableitung aus einem Suffix-Baum. Man konstruiert zunächst den Suffix-Baum und durchläuft diesen dann mit einer Tiefensuche. Berücksichtigt man dabei die Ordnung auf den Kindern eines Knotens, die durch die lexikographische Ordnung der Markierungen der ausgehenden Kanten gegeben ist, so werden alle Blätter (bzw. die zugehörigen Suffixe) in lexikographischer Ordnung aufgefunden.

Theorem 5.2 *Ein Suffix-Array für ein gegebenes Wort kann in linearer Zeit konstruiert werden.*

Da der Suffix-Baum jedoch selbst schon relativ viel Platz benötigt, ist man an speicherplatzsparenden Varianten zur Erstellung von Suffix-Arrays interessiert.

5.1.3 Algorithmus von Manber und Myers

In diesem Abschnitt stellen wir einen ersten direkten Algorithmus zur Konstruktion von Suffix-Arrays vor. Wir werden dabei zunächst auf die Idee des Bucket-Sorts zurückgreifen. Wir initialisieren das Feld A fortlaufend mit alle Suffixen, absteigend nach ihrer Länge, also $A[i] = n + 1 - i$ für $i \in [0 : n]$. Dann sortieren wir die Zeichenreihen nach dem ersten Zeichen. Eine platzsparende Implementierung mit maximal zwei Feldern der gleichen Länge sei dem Leser zur Übung überlassen. Das Ergebnis dieses Bucket-Sorts für $t = \text{MISSISSIPPI}$ ist in Abbildung 5.2 illustriert.

	\$	I					M	P		S			
A	12	2	5	8	11	1	9	10	3	4	6	7	

Abbildung 5.2: Beispiel: Bucket-Sort nach dem ersten Zeichen für MISSISSIPPI

Als Ergebnis des Bucket-Sorts nach dem ersten Zeichen erhalten wir eine Menge von Buckets, in denen sich nur Suffixe befinden, die mit demselben Zeichen beginnen. Jetzt müssen wir nach dem zweiten Zeichen sortieren, allerdings nur innerhalb der Buckets. Wir könnten hierfür wieder einen Bucket-Sort verwenden, es gibt aber eine trickreichere, effizientere Variante.

Seien $t^i = t_i \cdots t_n \$$ und $t^j = t_j \cdots t_n \$$ zwei Wörter eines Buckets, dann gilt $t_i = t_j$. Um einen Bucket mit den beiden Suffixe t^i und t^j nach dem zweiten Zeichen zu sortieren, müssen wir t_{i+1} mit t_{j+1} vergleichen. Diesen Vergleich haben wir jedoch schon implizit ausgeführt. Dieser Vergleich hat dasselbe Ergebnis wie der Vergleich von $t^{i+1} = t_{i+1} \cdots t_n \$$ mit $t^{j+1} = t_{j+1} \cdots t_n \$$. Wir müssen also nur in jedem Bucket die Zeichenreihen t^i nach dem Ergebnis der Sortierung von t^{i+1} nach dem ersten Zeichen anordnen.

Dabei werden wir uns etwas geschickter anstellen, indem wir alle Indexpositionen des sortierten Feldes A sortiert nach dem ersten Zeichen durchgehen. Ist $A[i] = j$, dann schieben wir den Index $j - 1$ in seinem Bucket an den aktuellen Anfang des Buckets und aktualisieren den Anfang des Buckets. Damit erhalten wir eine neue Einteilung von Buckets, in denen die enthaltenen Suffixe jeweils in den ersten beiden Zeichen übereinstimmen.

Nun sortieren wir innerhalb der Buckets nicht nach dem dritten, sondern gleich nach dem dritten und vierten Zeichen. Um diesen Bucket nach dem dritten und vierten Zeichen zu sortieren, müssen wir für t^i und t^j die Zeichenreihen $t_{i+2}t_{i+3}$ mit $t_{j+2}t_{j+3}$ vergleichen. Diesen Vergleich haben wir wiederum bereits implizit ausgeführt. Dieser Vergleich hat dasselbe Ergebnis wie der Vergleich von $t^{i+2} = t_{i+2} \cdots t_n \$$ mit $t^{j+2} = t_{j+2} \cdots t_n \$$. Wir müssen also nur in jedem Bucket die Zeichenreihen t^i nach dem Ergebnis der Sortierung von t^{i+2} nach den ersten beiden Zeichen anordnen.

A	1	2	3	4	5	6	7	8	9	10	11	12
	12	2	5	8	11	1	9	10	3	4	6	7
12		11										
2						1						
5									4			
8									4	7		
11							10					
1												
9		11	8									
10							10	9				
3		11	8	2								
4									4	7	3	
6		11	8	2	5							
7									4	7	3	6
	12	11	8	2	5	1	10	9	4	7	3	6
	\$	I	I	I	I	M	P	P	S	S	S	S
		\$	P	S	S	I	I	P	I	I	S	S
			P	S	S	S	\$	I	S	P	I	I

Abbildung 5.3: Beispiel: Sortierung während der ersten Phase nach dem zweiten Zeichen

Dies werden wir wieder etwas geschickter machen, indem wir alle Indexpositionen des nach den ersten beiden Zeichen sortierten Feldes A durchgehen. Ist $A[i] = j$, dann schieben wir den Index $j - 2$ in seinem Bucket an den aktuellen Anfang des Buckets und aktualisieren den Anfang des Buckets. Damit erhalten wir eine neue Einteilung von Buckets, die jeweils nach den ersten vier Zeichen sortiert sind.

Somit verdoppelt sich in jeder Phase die Anzahl der Zeichen, nach denen die Suffixe sortiert sind. Nach maximal $\log(n)$ Phasen ist das Suffix-Array also fertiggestellt, d.h. wir erhalten $n + 1$ Buckets, die jeweils genau einen Suffix beinhalten.

In Abbildung 5.3 ist die erste Phase noch einmal exemplarisch dargestellt. In der ersten Zeile sind die Buckets aus dem Bucket-Sort gegeben. In jeder Zeile wird das jeweils nächste Element j aus dem bisher sortierten Feld A verwendet. Dann wird der Index $j - 1$ innerhalb seines Buckets nach vorne bewegt. In der letzten Zeile sind dann die resultierende Sortierung und die zugehörigen Anfänge der Suffixe dargestellt.

In Abbildung 5.4 ist die zweite Phase noch einmal anhand unseres Beispiels dargestellt. In der ersten Zeile sind die Buckets aus der ersten Phase gegeben. In jeder Zeile wird das jeweils nächste Element j aus dem bisher sortierten Feld A verwendet. Dann wird der Index $j - 2$ innerhalb seines Buckets nach vorne bewegt. In der letzten Zeile sind dann die resultierende Sortierung und die zugehörigen Anfänge der Suffixe dargestellt.

A	1	2	3	4	5	6	7	8	9	10	11	12
	12	11	8	2	5	1	10	9	4	7	3	6
12							10					
11								9				
8											6	
2												
5											6	3
1												
10			8									
9									7			
4				2								
7				2	5							
3						1						
6									7	4		
	12	11	8	2	5	1	10	9	7	4	6	3
	\$	I	I	I	I	M	P	P	S	S	S	S
		\$	P	S	S	I	I	P	I	I	S	S
			P	S	S	S	\$	I	P	S	I	I
			I	I	I	S		\$	P	S	P	S
			\$	S	P	I			I	I	P	S

Abbildung 5.4: Beispiel: Sortierung während der zweiten Phase nach dem dritten und vierten Zeichen

In Abbildung 5.5 ist eine einfache Version des Algorithmus von Manber und Myers in Pseudo-Code angegeben. Hierbei bezeichnet A das Feld der Länge n , in dem am Ende das Suffix-Array stehen soll. Das Feld A' der Länge n ist ein Hilfsfeld, um die Buckets von A zu sortieren. Das Ergebnis wird in A' zwischengespeichert, bevor dann am Ende einer Phase A' wieder auf A kopiert wird. B ist ein Bit-Feld der Länge n , das an Position i genau dann TRUE ist, wenn an Position i ein neuer Bucket startet. Das Feld R der Länge n gibt an, an welcher Stelle in A der Index i steht. Es gilt also genau dann $A[i] = j$, wenn $R[j] = i$ ist. Damit gilt insbesondere auch $A[R[i]] = i$ und $R[A[i]] = i$. Das Feld L verweist immer auf die linkeste Position innerhalb eines Buckets. Es gilt also genau dann $L[i] = i$, wenn $B[i] = \text{TRUE}$ gilt. Ansonsten ist $L[i] < i$. Genauer gesagt gilt $L[i] = \max \{j \leq i : B[j] = \text{TRUE}\}$. Weiter verwenden wir hier die C-Notation $(b)?x : y$ die denn Wert x ergibt, wenn der Boolesche Ausdruck b wahr ist und y sonst.

Während des Algorithmus wird jedoch $L[i]$ für $B[i] = \text{TRUE}$ auch dazu missbraucht, um den ersten freien Platz innerhalb seines Bucket anzugeben. Auf diese Position wird dann der nächste Wert innerhalb des Buckets nach vorne geschoben. Weiterhin kann es leere Stelle in A' geben, nämlich genau dann, wenn in A an dieser Position ein Suffix mit Start-Position größer als $n - 2^p$ steht. Da diese Suffixe jedoch eine Länge

```

MANBER_MYERS (string  $t$ )
{
     $A = \text{bucketsort}(t)$ ; /* sort all suffixes of  $t$  according to their first characters */

     $B[0] = \text{TRUE}$ ;
    for ( $i = 1$ ;  $i \leq n$ ;  $i++$ )
         $B[i] = (t[A[i]] \neq t[A[i+1]])$ ; /* initialize  $B$  */

    for ( $p = 0$ ;  $p \leq \log(n)$ ;  $p++$ )
    {
        for ( $i = 0$ ;  $i \leq n$ ;  $i++$ )
        {
             $R[A[i]] = i$ ; /* compute  $R$  */
             $A'[i] = 0$ ; /* initialize  $A'$  */
             $L[i] = (B[i]) ? i : L[i-1]$ ; /* initialize  $L$  */
        }

        for ( $i = 0$ ;  $i \leq n$ ;  $i++$ )
            if ( $A[i] - 2^p > 0$ )
            {
                 $j = R[A[i] - 2^p]$ ; /* move suffix from position  $j$  */
                 $p = (L[j] = j) ? L[j] : L[L[j]]$ ; /* to position  $p$  */
                 $b = (L[i] = i) ? i : L[i]$ ; /* bucket name of current suffix */
                 $b' = A'[p]$ ; /* bucket name of previous suffix */
                 $A'[p] = A[j]$ ; /* move suffix to front of its bucket */
                if (not  $B[p+1]$ )
                     $A'[p+1] = b$ ; /* save bucket name of current suffix */
                 $B[p] = ((B[p] \text{ or } (b \neq b')))$ ; /* update  $B$  */
                 $j = (L[j] = j) ? j : L[j]$ ;
                 $L[j]++$ ; /* position for next suffix within bucket */
            }

        for ( $i = 0$ ;  $i \leq n$ ;  $i++$ )
            if ( $A'[i] > 0$ )
                 $A[i] = A'[i]$ ; /* copy  $A'$  to  $A$  */
    }
}

```

Abbildung 5.5: Algorithmus von Manber und Myers

kleiner als 2^p besitzen, sind diese in A bereits korrekt sortiert und müssen beim Umkopieren von A' nach A nicht mehr berücksichtigt werden. Um solche Positionen in A' leicht feststellen zu können, wurde A' mit 0 (einer nichtexistenten Startposition) initialisiert.

Um die neu entstehenden Bucket-Grenzen bei der Sortierung feststellen zu können, speichern wir hinter der Start-Position des Suffixes den Namen des Buckets, aus dem der Suffix stammt, um den vorhergehenden Suffix an den Anfang des Buckets zu bewegen. Damit können wir leicht feststellen, wann eine neue Bucketgrenze entsteht, nämlich genau dann, wenn zwei aufeinander folgende Suffixe aufgrund von Suffixen aus verschiedenen Buckets entstehen. Dies ist auch in Abbildung 5.6 illustriert.

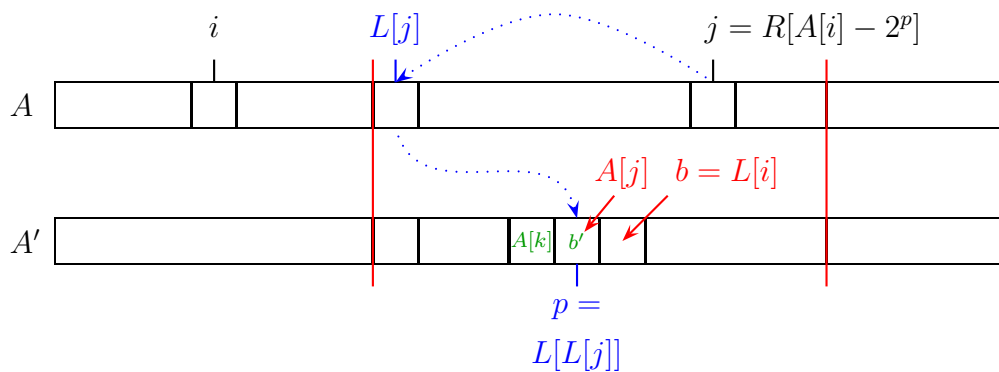


Abbildung 5.6: Skizze: Bucket-Sortierung im Algorithmus von Manber und Myers

Theorem 5.3 Mit dem Algorithmus von Manber und Myers kann ein Suffix-Array für eine gegebene Zeichenreihe der Länge n in Zeit $O(n \log(n))$ erstellt werden.

Man kann zeigen, dass man die Implementierung soweit verbessern kann, dass neben dem Eingabewort nur noch zwei Felder der gleichen Länge von Integers benötigt werden. (Sowie ein Bit-Feld der gleichen Länge, dass jedoch als Vorzeichen in den anderen Integer-Feldern versteckt werden kann.)

Unter geeigneten Annahmen über die Verteilung von Eingabewörtern kann man zeigen, dass man erwarten kann, dass sich alle Suffixe spätestens nach $O(\log(n))$ Zeichen unterscheiden. Im Erwartungswert sind also nur $O(\log \log(n))$ statt $O(\log(n))$ Phasen notwendig. Somit erhalten wir den folgenden Satz.

Theorem 5.4 Mit dem Algorithmus von Manber und Myers kann ein Suffix-Array für eine gegebene Zeichenreihe der Länge n mit erwarteter Laufzeit $O(n \log \log(n))$ erstellt werden.

Die erwartete Laufzeit kann mit einem Trick noch weiter auf eine lineare Laufzeit auf Kosten einer Speicherplatzvergrößerung um einen konstanten Faktor gesenkt werden. Hierfür verweisen wir den Leser jedoch auf die Originalliteratur von Manber und Myers, da wir im Folgenden noch einen einfachen Linearzeit-Algorithmus zur Konstruktion von Suffix-Arrays vorstellen werden.

5.2 Algorithmus von Ko und Aluru (*)

In diesem Abschnitt stellen wir einen direkten Algorithmus zur Konstruktion von Suffix-Arrays in Linearzeit vor. Dieser Abschnitt ist nur der Vollständigkeit halber Teil des Skripts, da der Stoff Bestandteil der Vorlesung im Wintersemester 2003/04 war.

5.2.1 Typ S und Typ L Suffixe

Zunächst einmal werden wir die Suffixe von $t\$$ in zwei verschiedene Typen klassifizieren.

Definition 5.5 Sei $t = t_1 \cdots t_n \in \Sigma^*$, dann heißt $t^i := t_i \cdots t_n\$$ vom Typ S, wenn $t^i < t^{i+1}$ gilt, und vom Typ L, wenn $t^i > t^{i+1}$ gilt. Der Suffix $t^{n+1} = \$$ ist sowohl vom Typ S als auch vom Typ L.

Man beachte, dass nur der Suffix $\$$ beiden Typen angehört. Wir zeigen jetzt noch, dass man alle Suffixe von $t\$$ in Linearzeit klassifizieren kann.

Lemma 5.6 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Alle Suffixe von $t\$$ können in Zeit $O(n)$ in die Typen S und L klassifiziert werden.

Beweis: Sei t^i ein Suffix von $t\$$. Wir unterscheiden jetzt zwei Fälle, je nachdem, ob $t_i = t_{i+1}$ gilt oder nicht.

Fall 1: Es gilt $t_i \neq t_{i+1}$. Dann ist t^i vom Typ S, wenn $t_i < t_{i+1}$ und andernfalls vom Typ L.

Fall 2: Es gilt $t_i = t_{i+1}$. Sei $j > i$ die kleinste Position mit $t_j \neq t_i$, d.h. es gilt $j = \min \{k \in [i : n + 1] : t_k \neq t_i\}$. Dieser Fall ist in Abbildung 5.7 illustriert.

Gilt nun $t_i < t_j$ dann sind $t^i, \dots, t^{j-1}$ alle vom Typ S. Gilt andererseits $t_i > t_j$, dann sind $t^i, \dots, t^{j-1}$ alle vom Typ L.

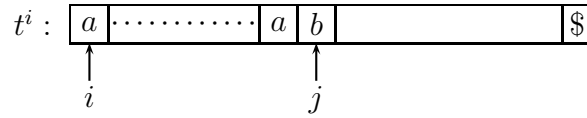


Abbildung 5.7: Skizze: 2. Fall im Beweis der Klassifizierung nach Typ S bzw. Typ L

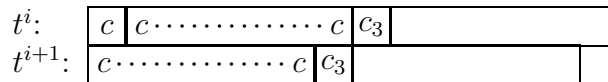
Offensichtlich lassen sich somit alle Suffixe mittels eines Durchlaufs über t in linearer Zeit in die Typen S und L klassifizieren. ■

Das folgende Lemma liefert uns für die Suffixe, die mit demselben Zeichen beginnen, eine einfache Aussage, ob $t^i < t^j$ gilt oder nicht, die nur auf der Klassifizierung nach den Typen S und L basiert.

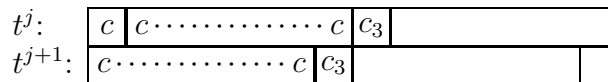
Lemma 5.7 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und seien t^i und t^j zwei verschiedene Suffixe von $t\$$. Sei weiter t^i vom Typ L und t^j vom Typ S. Gilt $t_i = t_j$, dann ist $t^i < t^j$.

Beweis: Wir führen den Beweis durch Widerspruch. Wir nehmen also an, dass $t^i > t^j$ gilt. Sei $c = t_i = t_j \in \Sigma$ und seien $c_1 \neq c_2 \in \Sigma$ die beiden Zeichen, in denen sich die Wörter t^i und t^j das erste Mal (von links nach rechts) unterscheiden. Sei weiter $t^i = c \cdot \alpha \cdot c_1 \cdot \beta$ und $t^j = c \cdot \alpha \cdot c_2 \cdot \beta'$ mit $\alpha, \beta, \beta' \in \Sigma^*$

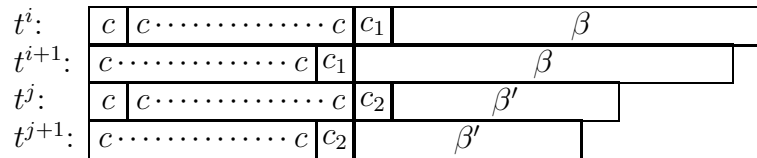
Fall 1: Es gilt $\alpha \notin \{c\}^*$. Sei c_3 das erste Zeichen in α , das ungleich c ist. Da t^i vom Typ L ist, muss $c > c_3$ gelten, wie man in der folgenden Abbildung 5.8 leicht sieht.

Abbildung 5.8: Skizze: t^i im Fall 1 $\alpha \notin \{c\}^*$

Da t^j vom Typ S ist, muss $c < c_3$ gelten, wie man in der folgenden Abbildung 5.9 leicht sieht.

Abbildung 5.9: Skizze: t^j im Fall 1 $\alpha \notin \{c\}^*$

Somit muss sowohl $c < c_3$ als auch $c > c_3$ gelten, was den gewünschten Widerspruch liefert.

Abbildung 5.10: Skizze: t^i und t^j im Fall 2

Fall 2: Es gilt $\alpha = c^{|\alpha|}$. Da t^i vom Typ L ist, gilt $c > c_1$. Da t^j vom Typ S ist, gilt $c < c_2$. Dies sieht man leicht in der folgenden Abbildung 5.10.

Also muss $c_1 < c < c_2$ gelten. Da aber nach Annahme $t_i > t_j$ gilt, muss $c_1 > c_2$ sein. Dies liefert den gewünschten Widerspruch. ■

Als Folgerung dieses Lemmas erhalten wir das folgende Korollar, dass uns Auskunft über die Verteilung der Suffixe innerhalb eines Buckets liefert, wenn wir alle Suffixe zunächst mit einem Bucketsort nach dem ersten Zeichen sortiert haben.

Korollar 5.8 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei $S_c = \{t^i : t_i = c\}$ die Menge aller Suffixe von $t\$$, die mit $c \in \Sigma$ beginnen. Im Suffix-Array von $t\$$ erscheinen alle Suffixe aus S_c vom Typ L vor denjenigen vom Typ S.

5.2.2 Sortieren der Suffixe mit sortierten Typ S/L Suffixen

Der Algorithmus von Ko und Aluru benötigt im Wesentlichen die folgenden drei Felder:

- Ein Feld A der Länge $n + 1$, das die Startpositionen aller Suffixe von $t\$$ in beliebiger Reihenfolge enthält.
- Ein Feld R der Länge $n + 1$, mit $A[R[j]] = j$.
- Ein Feld C der Länge maximal $n + 1$, das die Startpositionen aller Suffixe von $t\$$ vom Typ S enthält und bereits sortiert ist.

Dabei ist das Feld A mehr oder weniger gegeben, das Feld R lässt sich leicht aus dem Feld A in Linearzeit berechnen. Wie wir das Feld C erhalten, werden wir uns später noch genauer überlegen.

Wir werden jetzt zeigen, wie man mithilfe dieser drei Felder das Feld A sortieren kann. Dazu werden im Wesentlichen die folgenden drei Schritte ausgeführt:

1. Führe einen Bucket-Sort von A nach dem ersten Zeichen der Suffixe aus.

Wie wir schon gesehen haben, kann dies in linearer Zeit durchgeführt werden.

2. Durchlaufe das Feld C von rechts nach links (vom längsten zum kürzesten Suffix) und verschiebe den aktuell betrachteten Suffix $c \in C$ an das aktuelle Ende seines Buckets, indem die Startposition c mit dem Suffix unmittelbar vor dem aktuellen Ende vertauscht wird. Dies ist in Abbildung 5.11 illustriert. Die schraffierten Bereiche bezeichnen hierbei bereits abgearbeitete Bereiche des Feldes C bzw. die bereits korrekt sortierten Suffixe vom Typ S im aktuellen Bucket.

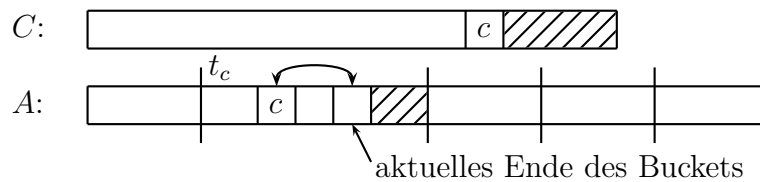


Abbildung 5.11: Skizze: Vertauschen des aktuell betrachteten Suffixes an den Anfang des Ende des Buckets

Offensichtlich ist die Laufzeit dieses Schrittes durch die Länge des Feldes C beschränkt, also ebenfalls in linearer Zeit ausführbar.

3. Durchlaufe das Feld A von links nach rechts (also aufsteigend). Für jedes $A[i]$ mit $t^{A[i]-1}$ vom Typ L verschiebe $A[i] - 1$ an den aktuellen Beginn seines Buckets. Dies ist in Abbildung 5.12 illustriert. Hierbei stellen die schraffierten Bereiche den bereits abgearbeiteten Teil des Feldes A bzw. die bereits korrekt sortierten Suffix vom Typ S und vom Typ L im aktuell betrachteten Bucket dar.

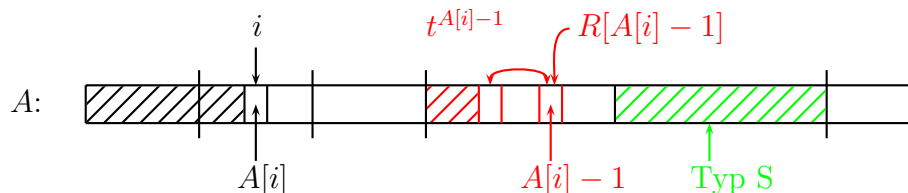


Abbildung 5.12: Skizze: Verschiebung der Typ L Suffixe an den Anfang ihrer Buckets

Auch dieser Schritt ist offensichtlich in linearer Zeit durchführbar.

In Abbildung 5.13 ist der Algorithmus anhand eines Beispiels illustriert. Der gesamte Algorithmus hat also offensichtlich eine lineare Laufzeit. Wir müssen jetzt noch zeigen, dass dieser Algorithmus auch korrekt ist.

Lemma 5.9 *Nach Schritt 3, d.h. wenn die Position i im Feld A erreicht wird, ist der Suffix $t^{A[i]}$ an der korrekten (sortierten) Position in A .*

	1	2	3	4	5	6	7	8	9	10	11	12				
	M	I	S	S	I	S	S	I	P	P	I	\$				
Typ		S			S			S				S				
C:	<table><tr><td>12</td><td>8</td><td>5</td><td>2</td></tr></table>												12	8	5	2
12	8	5	2													
Schritt 1																
A:	12	2	5	8	11	1	9	10	3	4	6	7				
Schritt 2		11	5	8	2											
		11	8	5	2											
		11	8	5	2											
	12															
Schritt 3	12	11	8	5	2	1	9	10	3	4	6	7				
		11		8	5	2										
			8				10	9								
				5					7	4	6	3				
					2				7	4	6	3				
						1	10	9	7	4	6	3				
									7	4	6	3				
Ergebnis:	\$	I	I	I	I	M	P	P	S	S	S	S				

Abbildung 5.13: Beispiel: Der Algorithmus von Ko und Aluru für MISSISSIPPI und gegebenem Feld C

Beweis: Wir führen den Beweis durch Induktion nach i .

Induktionsanfang ($i = 0$): Das ist korrekt, da nach dem Bucket-Sort nach dem ersten Zeichen $A[0] = |t\$|$ gilt, d.h. $A[0]$ entspricht dem Suffix \$, und dieser ist nach Definition der kleinste.

Induktionsschritt ($i-1 \rightarrow i$): Nach Induktionsvoraussetzung ist $A[0], \dots, A[i-1]$ bereits korrekt sortiert. Wir führen den Beweis für die Position i jetzt durch einen Widerspruchsbeweis. Dazu nehmen wir an, dass ein $k > i$ mit $t^{A[i-1]} < t^{A[k]} < t^{A[i]}$ existiert. Dabei sei k so gewählt, dass es das i -te Element in der wirklich sortierten Folge ist.

Nach Schritt 2 sind alle Suffixe vom Typ S an der korrekten Position innerhalb des Buckets und damit auch innerhalb des Feldes. Dies folgt aus der Tatsache, dass beim Verschieben an das Ende des Buckets die relative Reihenfolge des sortierten Feldes C in jedem Bucket erhalten bleibt. Somit müssen also die Suffixe $t^{A[i]}$ und $t^{A[k]}$ beide vom Typ L sein.

Da im ersten Schritt das Feld A mittels eines Bucket-Sort nach dem ersten Zeichen sortiert wurde, muss $t_{A[k]} = t_{A[i]} = c$ für ein $c \in \Sigma$ gelten. Sei also $t^{A[i]} = c \cdot \alpha$ und $t^{A[k]} = c \cdot \beta$ mit $\alpha, \beta \in \Sigma^*\$$.

Da $t^{A[k]}$ vom Typ L ist, muss $t^{A[k]} > \beta$ sein. Aus $\beta < t^{A[k]}$ folgt, dass der Suffix β bereits korrekt sortiert worden sein muss, d.h. $\beta \in \{t^{A[1]}, \dots, t^{A[i-1]}\}$.

Da nach Annahme $t^{A[k]} < t^{A[i]}$ gilt, muss weiter $\beta < \alpha$ gelten. Da $\beta < \alpha$ gilt, muss nach Algorithmenvorschrift $t^{A[k]}$ vor $t^{A[i]}$ an den aktuellen Beginn seines Buckets vertauscht worden sein. Dann würde aber $t^{A[i]}$ nicht vor $t^{A[k]}$ in seinem Bucket stehen. Das liefert den gewünschten Widerspruch. ■

Analog können wir das Feld A auch sortieren, wenn wir die sortierte Reihenfolge der Suffixe vom Typ L kennen, wie der folgenden Algorithmus zeigt:

1. Bucket-Sort von A nach dem ersten Zeichen der Suffixe.
2. Durchlaufe das Feld C von links nach rechts (vom kürzesten zum längsten Suffix) und verschiebe den aktuell betrachteten Suffix $c \in C$ an den aktuellen Beginn seines Buckets, in dem c mit dem Suffix unmittelbar nach dem aktuellen Beginn vertauscht wird.
3. Durchlaufe das Feld A von rechts nach links (also absteigend). Für jedes $A[i]$ mit $t^{A[i]-1}$ vom Typ S verschiebe $A[i] - 1$ an das aktuelle Ende seines Buckets.

Der Korrektheitsbeweis ist analog zu dem vom Lemma 5.9. Damit erhalten wir insgesamt den folgenden Satz.

Theorem 5.10 *Sind alle Typ S oder alle Typ L Suffixe von $t\$$ bereits sortiert, dann können alle Suffixe von $t\$$ in Zeit und Platz $O(|t|)$ sortiert werden.*

5.2.3 Sortierung der Typ S Suffixe

Wir müssen jetzt also nur noch zeigen, wie man alle Suffixe vom Typ S (bzw. vom Typ L) in linearer Zeit sortieren kann.

Definition 5.11 Sei $t = t_1 \cdots t_n \in \Sigma^*$. Eine Position $j \in [1 : n + 1]$ ist vom Typ S (bzw. L), wenn t^j vom Typ S (bzw. L) ist. Ein Teilwort $t_i \cdots t_j \sqsubseteq t\$$ ist vom Typ S (bzw. L), wenn i und j Positionen vom Typ S (bzw. L) sind und jede Position $k \in (i : j)$ vom Typ L (bzw. S) ist.

Betrachten wir als Beispiel das Wort MISSISSIPPI\$. Die roten Positionen sind vom Typ S. Teilwörter vom Typ S sind dann ISSI und IPPI\$. Dies ist in Abbildung 5.14 noch einmal illustriert.

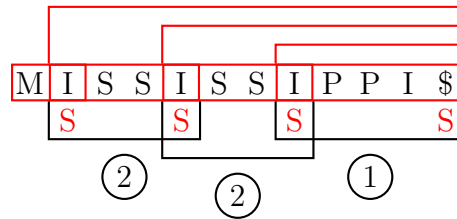


Abbildung 5.14: Beispiel: Typ S Teilwörter von MISSISSIPPI

Im Folgenden nehmen wir ohne Beschränkung der Allgemeinheit an, dass $|\mathcal{S}| \leq |\mathcal{L}|$ gilt. Wie können wir nun die Menge $\mathcal{S}$ sortieren? Wir sortieren dafür zunächst die Menge der Typ S Teilwörter. In unserem Beispiel ist $\text{IPPI\$} < \text{ISSI}$. Um mit dieser sortierten Menge sinnvoll weiterarbeiten zu können, geben wir jedem Typ S Teilwort einen neuen Namen, nämlich ein Zeichen eines neuen Alphabets, so dass diese Zuordnung die Sortierung des Typ S Teilwörter respektiert.

Notation 5.12 Sei $t \in \Sigma^*$ und sei $\mathcal{S}$ die Menge der Typ S Teilwörter von $t\$$, dann bezeichne $\pi(s)$ für $s \in \mathcal{S}$ die relative Ordnung in der sortierten Menge, d.h.

$$\pi(s) := |\{s' \in \mathcal{S} : s' \leq s\}|.$$

Basierend auf dieser Notation können wir jetzt das Wort t zu einem neuen Wort $\tilde{t}$ wie folgt umbauen, das im Wesentlichen als Suffixe die Typ S Suffixe von t besitzt.

Notation 5.13 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und seien $j_1 < \cdots < j_k$ die Typ S Positionen in $t\$$. Dann ist

$$\tilde{t} := \pi(t_{j_1} \cdots t_{j_2}) \cdot \pi(t_{j_2} \cdots t_{j_3}) \cdots \pi(t_{j_{k-1}} \cdots t_{j_k}) \in [1 : n + 1]^*.$$

Für unser Beispielwort MISSISSIPPI\$ ergibt sich dann $\tilde{t} := 2 \cdot 2 \cdot 1$

Jeder Typ S Suffix in $t\$$ korrespondiert eindeutig zu einem Suffix von $\tilde{t}$. Somit kann das Sortieren von Typ S Suffixen auf das Sortieren von Suffixen in $\tilde{t}$ zurückgeführt werden. Damit auch hier wieder das Wort mit einem kleinsten Symbol endet, wird an dieses Wort noch die Zahl 0 angehängt. Damit können wir die Sortierung der Typ S Suffixe auf das Sortieren der Suffixe eines anderen Wortes zurückführen.

Wie ermittelt man aber $\pi(s)$ für alle $s \in \mathcal{S}$? Die entspricht offensichtlich dem Sortieren der Menge $\mathcal{S}$. Dazu benötigen wir zuerst noch die folgende Definition.

Definition 5.14 Sei $t \in \Sigma^*$. Die S-Distanz $\sigma(t^i)$ für ein Suffix t^i von $t\$$ ist definiert durch:

$$\sigma(t^i) := \min \{j \in [1 : i - 1] : t^{i-j} \text{ ist vom Typ } S\},$$

wobei $\min\{\} := 0$ gilt.

Anschaulich gibt die S-Distanz an, um wie viele Zeichen man nach links im Wort gehen muss, um die vorherige Typ S Position aufzufinden, sofern es eine solche gibt. Die S-Distanz ist in der folgenden Abbildung 5.15 illustriert.

$$\begin{array}{ccccccccccccccc} & & \text{M} & \text{I} & \text{S} & \text{S} & \text{I} & \text{S} & \text{S} & \text{I} & \text{P} & \text{P} & \text{I} & \$ \\ & & & \text{S} & & & \text{S} & & & \text{S} & & & \text{S} & \\ \sigma : & 0 & 0 & 1 & 2 & 3 & 1 & 2 & 3 & 1 & 2 & 3 & 4 \end{array}$$

Abbildung 5.15: Beispiel: S-Distanz für MISSISSIPPI

Wenn man sich die Typ S Positionen in einem Booleschen Feld merkt, lassen sich die S-Distanzen mit einem Durchlauf über das Feld A in linearer Zeit ohne zusätzlichen Platz berechnen. Die Details bleiben dem Leser zur Übung überlassen.

Jetzt können wir den Algorithmus zum Sortieren der Typ S Teilwörter angeben. Dabei ist A ein Feld mit $n + 1$ Elementen, in dem die Suffixe von $t\$$ nach ihrem ersten Zeichen sortiert sind.

1. Sei $m := \max \{\sigma(t^i) : i \in [1 : n + 1]\}$. Erzeuge m Listen $L_j = \{i : \sigma(t^i) = j\}$ für $j \in [1 : m]$, wobei die Reihenfolge innerhalb der Listen der Reihenfolge innerhalb des nach dem ersten Zeichen sortierten Feldes A entspricht.

Diese lassen sich mit einem einfachen Scan über das Feld für das Wort t erstellen und mit Hilfe des Feldes R den entsprechenden Positionen in A zuordnen. Also kann dieser Schritt in linearer Zeit bewerkstelligt werden.

2. Sortiere alle Typ S Teilwörter wie folgt: Führe quasi wiederholte Bucket-Sorts nach dem j -ten Zeichen der Teilwörter aus: Für jedes $j \in [1 : m]$ und jedes

$i \in L_j$ bewege des Typ S Teilwort t^{i-j} (respektive die zugehörige Anfangsposition) an den aktuellen Anfang seines Buckets.

Nach jeder Abarbeitung einer Liste L_j müssen die neuen Bucketgrenzen angepasst werden, da jetzt das Feld nach den ersten j Zeichen der Typ S Teilwörter sortiert ist. Da die Summe aller Listenlängen durch $O(n)$ beschränkt ist, kann auch dieser Schritt in Linearzeit ausgeführt werden.

Warum werden mit diesem Algorithmus die Typ S Teilwörter korrekt sortiert? In der Liste L_j befinden sich die Typ S Suffixe in der Reihenfolge, wenn sie nach dem j -ten Zeichen sortiert sind. Mit der Umordnung der Startposition $i - j$ für $i \in L_j$ wird das jeweils kleinste Typ S Teilwort innerhalb seines Buckets nach vorne bewegt, also wird innerhalb der Buckets nach dem j -ten Zeichen sortiert. In Abbildung 5.16 ist dieser Sortierschritt noch einmal illustriert.

In diesem Schritt beachten wir allerdings folgende Modifikation des Bucket-Sorts bzw. der lexikographischen Ordnung: Ist α ein Präfix von β , dann definieren wir $\alpha > \beta$ (im Gegensatz zur üblichen lexikographischen Ordnung). Im Bucket-Sort kann diese Modifikation der Ordnung leicht implementiert werden. Warum wir diese Modifikation benötigen, wird im folgenden Lemma klar. Wir weisen allerdings darauf

$i :$	1	2	3	4	5	6	7	8	9	10	11	12
$t_i :$	M	I	S	S	I	S	S	I	P	P	I	\$
Type		S			S			S				S
$\sigma(t^i) :$	0	0	1	2	3	1	2	3	1	2	3	4

$i :$	1	2	3	4	5	6	7	8	9	10	11	12
$A[i] :$	12	2	5	8	11	1	9	10	3	4	6	7
$t_{A[i]} :$	\$	I	I	I	I	M	P	P	S	S	S	S
$\sigma(t^{A[i]}) :$	4	0	3	3	3	0	1	2	1	2	1	2

$$L_1 = [9, 3, 6]$$

$$L_2 = [10, 4, 7]$$

$$L_3 = [5, 8, 11]$$

$$L_4 = [12]$$

12:\$	2: ISSI	5: ISSI	8: IPPI\$
12:\$	8: IPPI\$	5: ISSI	2: ISSI
12:\$	8: IPPI\$	2: ISSI	5: ISSI

Abbildung 5.16: Beispiel: Sortieren des Typ S Teilwörter von MISSISSIPPI

hin, dass diese Modifikation beim Sortieren von Typ L Teilwörtern nicht erforderlich ist.

Lemma 5.15 *Sei $t \in \Sigma^*$. Seien t^i und t^j zwei Typ S Suffixe von $t\$$ und seien $\tilde{t}^{i'}$ und $\tilde{t}^{j'}$ zwei Suffixe von $\tilde{t} \in \mathbb{N}_0^+$, die jeweils zueinander korrespondieren. Dann gilt $t^i \leq t^j \Leftrightarrow \tilde{t}^{i'} \leq \tilde{t}^{j'}$.*

Beweis: $\Rightarrow$: Es gilt nach Voraussetzung $t^i \leq t^j$. Sei also $t^i = \alpha a \beta$ und $t^j = \alpha b \gamma$ mit $\alpha, \beta, \gamma \in \Sigma^*$ und $a < b \in \Sigma$. Diese Situation ist mit $\alpha = \alpha' \alpha''$ in Abbildung 5.17 illustriert. Beachte hierbei, dass in der Illustration die Bedingung $i < j$ vollkommen willkürlich gewählt wurde (es könnte auch $i > j$ gelten) und dass sich die beiden Bereiche α auch überlappen können.

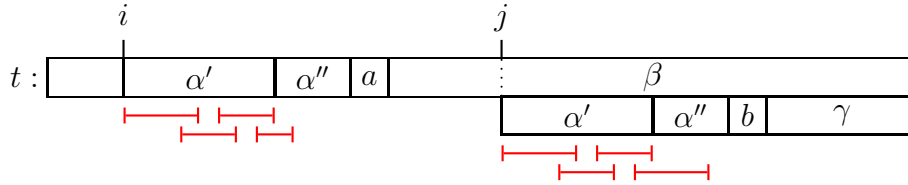


Abbildung 5.17: Skizze: $t^i < t^j$

Wir zerlegen jetzt die Zeichenreihe α in $\alpha = \alpha' \alpha''$, so dass in α' die Typ S Teilwörter in den Suffixen t^i und t^j gleich sind. Sollte α'' nicht das leere Wort sein, so kann dies nur passieren, wenn $\alpha'' = c^{|\alpha''|}$ für ein $c \in \Sigma$ gilt. Der einfache Beweis hierfür sei dem Leser überlassen. In der Abbildung 5.17 sind die Typ S Teilwörter im Bereich α durch die roten Balken schematisch dargestellt.

Ist $\alpha'' = \varepsilon$, dann sind die Typ S Teilwörter in t^i und t^j im Bereich α gleich. Das erste sich unterscheidende Teilwort vom Typ S muss das Zeichen a in t^i bzw. das Zeichen b in t^j enthalten. Nach unsere Sortierung gilt dann offensichtlich auch $\tilde{t}^{i'} < \tilde{t}^{j'}$.

Sei jetzt also $\alpha'' \neq \varepsilon$. Der Zerfall der Positionen in α'' von t^i bzw. t^j in Typ S bzw. Typ L Positionen hängt dann von dem folgenden Zeichen ungleich c in t^i bzw. in t^j und dessen Beziehung zu c ab. Auf jeden Fall sind alle Positionen in α'' vom jeweils gleichen Typ. Sind sowohl in t^i als auch in t^j alle Positionen in α'' vom selben Typ, dann beinhaltet auch α'' sowohl in t^i als auch in t^j dieselbe Zerlegung in Typ S Teilwörter und wir erhalten einen Widerspruch zur Wahl von α'' .

Da $t^i < t^j$ gilt, müssen alle Positionen in α'' in t^i vom Typ L und in t^j vom Typ S sein. Den ebenfalls nicht allzu schweren Beweis überlassen wir dem Leser. Dann sind in $\tilde{t}$ die Zeichen ab Position i' bzw. j' gleich bis es zur Übersetzung der Typ S Teilwörter in α'' kommt. In $\tilde{t}^{i'}$ folgt als nächstes Zeichen die Bucketnummer von

$c^k d$ für ein $1 < k \in \mathbb{N}$ und ein $c \neq d \in \Sigma$. In $\tilde{t}^{j'}$ folgt als nächstes Zeichen die Bucketnummer von cc . In diesem Fall ist also cc ein Präfix von $c^k d$ und nach unserer modifizierten Sortierreihenfolge gilt dann $\tilde{t}^{i'} < \tilde{t}^{j'}$.

$\Leftarrow$: Sei jetzt $\tilde{t}^{i'} < \tilde{t}^{j'}$ mit $\tilde{t}^{i'} = \mu \cdot m \cdot \rho$ und $\tilde{t}^{j'} = \mu \cdot m' \cdot \rho'$ mit $\mu, m, m', \rho, \rho' \in \mathbb{N}$ und $m < m'$.

Sei α das Teilwort von $t\$$ das zu μ korrespondiert. Seien weiter β und β' Teilwörter von t , die zu m und m' korrespondieren. Diese müssen also Typ S Teilwörter von $t\$$ sein. Dies ist in Abbildung 5.18 illustriert

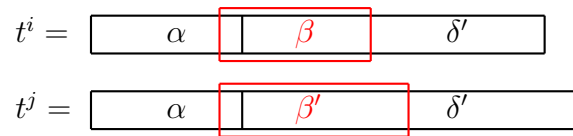


Abbildung 5.18: Skizze: Korrespondierende t^i und t^j zu $\tilde{t}^{i'} < \tilde{t}^{j'}$

Zunächst halten wir fest, dass β kein echtes Präfix von β' sein kann. Dies folgt aus Sortierung der Typ S Teilwörter. Dort gilt: ist β ein echtes Präfix von β' , dann ist $\pi(\beta) > \pi(\beta')$. Wir unterscheiden jetzt zwei Fälle, je nachdem, ob β' ein Präfix von β ist oder nicht.

Fall 1: Es gilt, dass β' kein Präfix von β ist. Dann gilt aufgrund von $m < m'$, dass $\beta < \beta'$ und somit $t^i < t^j$.

Fall 2: Jetzt gilt, dass β' ein echtes Präfix von β ist. Siehe hierzu auch Abbildung 5.19.

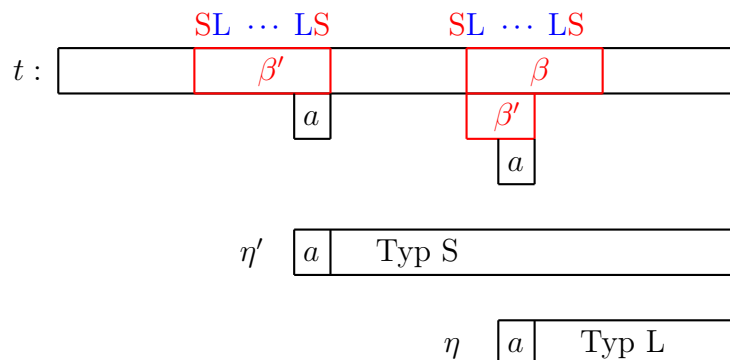


Abbildung 5.19: Skizze: β' ist ein echtes Präfix von β

Sei, wie in Abbildung 5.19 angegeben, a das letzte Zeichen von β' . Im Bereich β' muss sich a an einer Typ S Position befinden. Das weitere Vorkommen als letztes Zeichen

von β' als Präfix von β muss eine Typ L Position sein. Betrachten wir die beiden Suffixe η und η' , die an an diesen Positionen beginnen. Da beide mit demselben Zeichen beginnen, muss nach Lemma 5.7 $\eta < \eta'$ gelten. Damit folgt sofort nach der lexikographischen Ordnung, dass dann auch $t^i < t^j$ gilt. ■

5.2.4 Der Algorithmus von Ko und Aluru und Laufzeitanalyse

Wir wollen jetzt noch den vollständigen Algorithmus von Ko und Aluru angeben:

1. Bestimmung der Typ S-/L Positionen von t . Es gelte im Folgenden $|\mathcal{S}| \leq |\mathcal{L}|$.
2. Sortierung der Typ S Teilwörter (Bestimmung des Feldes C):
 - Bucket-Sort nach dem ersten Zeichen (von allen Typ S Suffixen).
 - Bestimmung der S-Distanzen und Generierung der Listen $L_1, \dots, L_m$.
 - Verfeinerung der Buckets mit Hilfe der Listen $L_1, \dots, L_m$.
3. Konstruktion von $\tilde{t} \in [1 : n + 1]^+$ aus $t \in \Sigma^*$.
 - Ersetzung der Typ S Teilwörter in t durch ihren Rang (mit Hilfe von Schritt 2)
4. Rekursives Sortieren der Suffixe von $\tilde{t}$. Es gilt: $|\tilde{t}| \leq \frac{n+2}{2}$.
5. Aus 4. folgt Sortierung der Typ S Suffixe von t .
6. Sortierung der Suffixe von t mit Hilfe der relativen Ordnung der Suffixe vom Typ S.
 - Bucket-Sort nach dem ersten Zeichen der Suffixe.
 - Verschiebe die Typ S Suffixe innerhalb seines Bucket ans Ende mit Hilfe der bekannten Sortierung der Typ S Suffixe.
 - Verschiebe die Typ L Suffixe innerhalb seines Buckets an den Anfang mit Hilfe des bereits teilweise sortierten Feldes A .

Wir kommen jetzt zur Laufzeitanalyse des Algorithmus von Ko und Aluru. Wie bisher bezeichne:

$$\begin{aligned}\mathcal{S} &:= \{t^j : t^j \text{ ist vom Typ S}\} \\ \mathcal{L} &:= \{t^j : t^j \text{ ist vom Typ L}\}\end{aligned}$$

Gilt nun $|\mathcal{S}| \leq |\mathcal{L}|$ dann sortieren wir alle Suffixe vom Typ S rekursiv (bzw. die neu konstruierten Zeichenreihe $\tilde{t}$), also die Menge $\mathcal{S}$, und anschließend das Feld A

mit Hilfe der sortierten Suffixe vom Typ S. Ansonsten sortieren wir alle Suffixe vom Typ L rekursiv (bzw. die neu konstruierten Zeichenreihe $\tilde{t}$), also die Menge $\mathcal{L}$, und anschließend das Feld A mit Hilfe der sortierten Suffixe vom Typ L. Für die Laufzeit gilt dann:

$$T(n) \leq T\left(\frac{n+2}{2}\right) + O(n).$$

Bekanntermaßen hat diese Rekursionsgleichung die Lösung $T(n) = O(n)$.

Theorem 5.16 *Ein Suffix-Array für $t \in \Sigma^n$ kann mit Hilfe des Algorithmus von Ko und Aluru in Zeit und Platz $O(n)$ konstruiert werden.*

Zum Abschluss noch einige Anmerkungen:

- Bei geschickter Implementierung benötigt man nur zwei zusätzliche Felder der Länge n und einige Felder von Booleschen Werten. Für Details sei der Leser auf die Originalliteratur von Ko und Aluru verwiesen.
- Bei der Rekursion findet ein Alphabetwechsel von $\Sigma \cup \{\$ \}$ auf $[0 : n + 1]$ statt (0 für das neue Endzeichen). Dies stellt jedoch keine Einschränkung dar, da man für ein Wort der Länge $n + 1$ über einem geordneten Alphabet Σ ohne Weiteres das Alphabet Σ selbst als Teilmenge von $[1 : n + 1]$ und $\$$ als 0 interpretieren kann.

5.3 Skew-Algorithmus von Kärkkäinen und Sanders

In diesem Abschnitt wollen wir noch einen einfacheren Algorithmus zur Konstruktion von Suffix-Arrays vorstellen, der jedoch nicht ganz so platzsparend wie der Algorithmus von Ko und Aluru ist.

5.3.1 Tripel und Verkürzung der Zeichenreihe

Sei wiederum $t \in \Sigma^*$ das Wort, für den das Suffix-Array erstellt werden soll. Hier nehmen wir für eine einfachere Darstellung an, dass $t\$$ von 0 mit n indiziert ist, also es gilt $t = t_0 \cdots t_{n-1}$.

Zuerst konstruieren wir alle dreibuchstabigen Teilwörter, so genannte *Tripel* von $t\$$, d.h. $t_i \cdot t_{i+1} \cdot t_{i+2}$. Damit alle Tripel eine Länge von genau 3 Zeichen haben, nehmen wir der Einfachheit halber an, dass das Wort mit 3 Dollarzeichen endet. Für das weitere Vorgehen benötigen wir nur die Tripel ab einer Position $i \in [0 : n]$ mit $i \bmod 3 \neq 0$.

Wir sortieren all diese Tripel mithilfe eines Bucket-Sorts und nummerieren sie nach ihrer Ordnung. Ein Beispiel ist in Abbildung 5.20 angegeben.

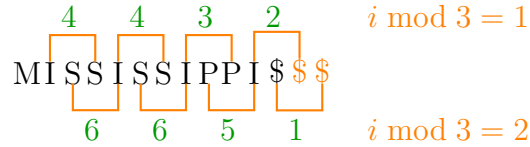


Abbildung 5.20: Beispiel: sortierte Tripel von MISSISSIPPI

Im zweiten Schritt erzeugen wir eine neue Zeichenreihe $\tilde{t}$ wie folgt. Für das Teilwort $t_1 \cdots t_{3\lceil n/3 \rceil}$ ersetzen wir die Tripel $t_i \cdot t_{i+1} \cdot t_{i+2}$ mit $i \bmod 3 = 1$ durch ihre Ordnungsnummer und erhalten somit die Zeichenreihe $t^{(1)}$. Dieselbe Prozedur führen wir für das Teilwort $t_2 \cdots t_{3\lceil n/3 \rceil+1}$ durch. Auch hier ersetzen wir die Tripel $t_i \cdot t_{i+1} \cdot t_{i+2}$ mit $i \bmod 3 = 2$ durch ihre Ordnungsnummer und erhalten somit die Zeichenreihe $t^{(2)}$. Nun setzen wir $\tilde{t} := t^{(1)} \cdot t^{(2)}$.

Im Beispiel MISSISSIPPI erhalten wir dann $\tilde{t} = 4 \cdot 4 \cdot 3 \cdot 2 \cdot 6 \cdot 6 \cdot 5 \cdot 1$, wobei die Zeichenreihe $t^{(1)}$ bzw. $t^{(2)}$ rot bzw. blau dargestellt ist. Für die Länge von $\tilde{t}$ gilt offensichtlich $|\tilde{t}| = 2\lceil \frac{n}{3} \rceil$.

5.3.2 Rekursives Sortieren der verkürzten Zeichenreihe

Da wir nun eine neue, kürzere Zeichenreihe erhalten haben, können wir unseren Sortieralgorithmus rekursiv aufrufen. Dies liefert uns das Feld A' mit den sortierten Suffixen (respektive ihrer Startpositionen) für $\tilde{t}$. Für unser Beispiel MISSISSIPPI ist das Ergebnis dies in Abbildung 5.21 dargestellt.

$A'[0]$	$=$	8	$\hat{=}$	0	$\hat{=}$	ε
$A'[1]$	$=$	7	$\hat{=}$	10	$\hat{=}$	\$\$\$
$A'[2]$	$=$	3	$\hat{=}$	266510	$\hat{=}$	I\$\$...
$A'[3]$	$=$	2	$\hat{=}$	3266510	$\hat{=}$	IPP I\$\$...
$A'[4]$	$=$	1	$\hat{=}$	43266510	$\hat{=}$	ISS IPP I\$\$...
$A'[5]$	$=$	0	$\hat{=}$	443266510	$\hat{=}$	ISS ISS IPP I\$\$...
$A'[6]$	$=$	6	$\hat{=}$	510	$\hat{=}$	PPI \$\$\$
$A'[7]$	$=$	5	$\hat{=}$	6510	$\hat{=}$	SSI PPI \$\$\$
$A'[8]$	$=$	4	$\hat{=}$	66510	$\hat{=}$	SSI SSI PPI \$\$\$

Abbildung 5.21: Beispiel: Das Feld A' für MISSISSIPPI

Wie man leicht sieht, liefert dieser Schritt bereits eine Sortierung aller Suffixe t^i mit $i \bmod 3 \neq 0$. Die Suffixe, die an Position $i \bmod 3 = 2$ beginnen, werden dabei ja so

berücksichtigt, wie wir es wollen. Die Suffixe, die an einer Position $i \bmod 3 = 1$ beginnen, werden noch durch einige Zeichen verlängert. Die Verlängerung wird jedoch durch eine Zeichenreihe aus einigen Dollarzeichen getrennt. Da das Dollarzeichen jedoch das kleinste Zeichen des Alphabets hat, spielt es quasi die Rolle eines Leerzeichens und die Präfixe werden korrekt sortiert. Man muss dabei nur beachten, dass sich nach Voraussetzung die Suffixe spätestens beim ersten Auftreten eines Dollarzeichens unterscheiden müssen.

19.12.06

Wie erhalten wir nun eine Indizierung der Elemente von t anstelle von $\tilde{t}$ (also der Triplets in t)? Zuerst beachten wir, dass $A'[0] = 2\lceil \frac{n}{3} \rceil + 1$ und somit das Sonderzeichen 0 in $\tilde{t}$ beschreibt, was der leeren Zeichenreihe in t entspricht. Dieser Array-Eintrag kann also ignoriert werden bzw. muss aus dem Array entfernt werden. Wir müssen nur berücksichtigen, dass die Array-Werte aus $[0 : \lceil \frac{n}{3} \rceil - 1]$ bzw. $[\lceil \frac{n}{3} \rceil : 2\lceil \frac{n}{3} \rceil - 1]$ in $\tilde{t}$ die Triplets aus $t^{(1)}$ bzw. $t^{(2)}$ beschreibt. Also rechnen wir die Information wie folgt um:

$$A_{12}[i - 1] = \begin{cases} 1 + 3 \cdot (A'[i]) & \text{falls } A'[i] < \lceil \frac{n}{3} \rceil, \\ 2 + 3 \cdot (A'[i] - \lceil \frac{n}{3} \rceil) & \text{sonst.} \end{cases}$$

Hierbei ist $i \in [1 : 2\lceil \frac{n}{3} \rceil]$. Somit entspricht A_{12} dem Array, dass die Startpositionen der sortierten Suffixe aus $\{t^i : i \bmod 3 \neq 0\}$ enthält.

Für unser Beispiel MISSISSIPPI ist das Ergebnis dies in Abbildung 5.22 dargestellt ($n = 11$).

$$\begin{array}{llllll} A_{12}[0] & = & 2 + 3 \cdot (\textcolor{red}{7} - 4) & = & 11 & \hat{=} \$ \\ A_{12}[1] & = & 1 + 3 \cdot (\textcolor{red}{3}) & = & 10 & \hat{=} I\$ \\ A_{12}[2] & = & 1 + 3 \cdot (\textcolor{red}{2}) & = & 7 & \hat{=} IPPI\$ \\ A_{12}[3] & = & 1 + 3 \cdot (\textcolor{red}{1}) & = & 4 & \hat{=} ISSIPPI\$ \\ A_{12}[4] & = & 1 + 3 \cdot (\textcolor{red}{0}) & = & 1 & \hat{=} ISSISSIPPI\$ \\ A_{12}[5] & = & 2 + 3 \cdot (\textcolor{red}{6} - 4) & = & 8 & \hat{=} PPI\$ \\ A_{12}[6] & = & 2 + 3 \cdot (\textcolor{red}{5} - 4) & = & 5 & \hat{=} SSIPPI\$ \\ A_{12}[7] & = & 2 + 3 \cdot (\textcolor{red}{4} - 4) & = & 2 & \hat{=} SSISSIPPI\$ \end{array}$$

Abbildung 5.22: Beispiel: Das Feld A_{12} für MISSISSIPPI

5.3.3 Sortieren der restlichen Suffixe

Im nächsten Schritt sortieren wir alle Suffixe, die an einer durch 3 teilbaren Position beginnen, d.h. wir betrachten nur Suffixe t^i mit $i \bmod 3 = 0$. Es gilt offensichtlich: $t^i = t_i \cdot t^{i+1}$ mit $(i + 1) \bmod 3 = 1$. Da uns ja die relative Ordnung der Suffixe t^{i+1} mit $(i + 1) \bmod 3 = 1$ durch das Feld A_{12} bereits bekannt ist, kann uns das Feld A_{12} für das Sortieren helfen.

Wir führen zuerst einen einfachen Bucket-Sort in der Menge $\{t^i : i \bmod 3 = 0\}$ nach dem erstem Zeichen durch. Mittels eines Scans über A_{12} von links nach rechts, verschieben wir beim Betrachten von $A_{12}[i]$ den Suffix $t^{A_{12}[i]-1}$ an den aktuellen Beginn seines Buckets, sofern $A_{12}[i] - 1 \bmod 3 = 0$ bzw. $A_{12}[i] \bmod 3 = 1$ gilt. Dies liefert uns das sortierte Feld A_0 für $\{t^i : i \bmod 3 = 0\}$.

5.3.4 Mischen von A_{12} und A_0

Jetzt müssen wir nur noch A_{12} und A_0 zusammenmischen. Damit die Vergleiche beim Mischen nicht zu aufwendig werden, beachten wir noch Folgendes. Wenn wir t^i mit t^j mit $i \bmod 3 = 0$ und $j \bmod 3 \in [1 : 2]$ vergleichen, gilt:

Fall 1: Sei $j \bmod 3 = 1$. Dann gilt:

$$\begin{aligned} t^i &= t_i \cdot t^{i+1}, \text{ wobei } (i+1) \bmod 3 = 1, \\ t^j &= t_j \cdot t^{j+1}, \text{ wobei } (j+1) \bmod 3 = 2. \end{aligned}$$

Wir vergleichen also nur t_i mit t_j . Bei Gleichheit greifen wir auf das Resultat in A_{12} zurück. Andernfalls kennen wir das Ergebnis des Vergleichs ja schon.

Fall 2: Sei $j \bmod 3 = 2$. Dann gilt:

$$\begin{aligned} t^i &= t_i \cdot t_{i+1} \cdot t^{i+2}, \text{ wobei } (i+2) \bmod 3 = 2, \\ t^j &= t_j \cdot t_{j+1} \cdot t^{j+2}, \text{ wobei } (j+2) \bmod 3 = 1. \end{aligned}$$

Wir müssen also nur zuerst t_i mit t_j vergleichen, bei Gleichheit dann auch noch t_{i+1} mit t_{j+1} . Ergeben beide Vergleiche eine Gleichheit, dann können wir das Ergebnis mittels eines Vergleichs von t^{i+2} mit t^{j+2} ebenfalls wieder aus A_{12} ablesen.

Für den Vergleich der beiden Werte aus A_{12} benötigen wir in Wirklichkeit wieder das inverse Feld R_{12} mit $A_{12}[R_{12}[i]] = i$. Der vollständige Algorithmus von Kärkkäinen und Sanders ist in Abbildung 5.23 angegeben.

5.3.5 Laufzeitanalyse

Zum Abschluss untersuchen wir noch die Laufzeit des Algorithmus von Kärkkäinen und Sanders. Der Algorithmus gehorcht offensichtlich der folgenden Rekursionsgleichung:

$$T(n) = T\left(2 \left\lceil \frac{n}{3} \right\rceil\right) + O(n)$$

Bekanntermaßen hat diese Rekursionsgleichung die Lösung $T(n) = O(n)$.

Theorem 5.17 *Sei $t = t_0 \cdots t_{n-1} \in \Sigma^*$. der Algorithmus von Kärkkäinen und Sanders kann in Zeit und Platz von $O(n)$ das zugehörige Suffix-Array konstruieren.*

1. Sortiere $\mathcal{T} = \{t_i \cdot t_{i+1} \cdot t_{i+1} : i \bmod 3 \neq 0\}$ mithilfe eines Bucket-Sorts und sei $\pi(i) := |\{t_j \cdot t_{j+1} \cdot t_{j+1} \leq t_i \cdot t_{i+1} \cdot t_{i+1} : j \bmod 3 \neq 0\}|$.
2. Setze $t^{(1)} := \pi(1) \cdot \pi(4) \cdots \pi(3\lceil n/3 \rceil - 2)$ und $t^{(2)} := \pi(2) \cdot \pi(5) \cdots \pi(3\lceil n/3 \rceil - 1)$.
Setze $\tilde{t} := t^{(1)} \cdot t^{(2)}$.
3. Sortiere die Suffixe von $\tilde{t} \cdot 0$ rekursiv. Sei A_{12} das zugehörige Suffix-Array.
4. Sortiere $\{t^i : i \bmod 3 = 0\}$ nach dem ersten Zeichen mithilfe eines Bucketsorts.
5. Sortiere $\{t^i : i \bmod 3 = 0\}$ mithilfe des Feldes A_{12} in das Feld A_0 .
6. Mische A_{12} mit A_0 .

Abbildung 5.23: Algorithmus: Konstruktion vom Suffix-Arrays nach Kärkkäinen und Sanders

5.4 Suchen in Suffix-Arrays

Eine einfache Anwendung von Suffix-Bäumen ist das schnelle Auffinden von Teilwörtern in einem vorverarbeiteten Text. Wie sieht diese Aufgabe nun in Suffix-Arrays aus? In diesem Abschnitt wollen wir zeigen, dass man auch in Suffix-Arrays relativ effizient nach Teilwörtern suchen kann. In diesem Abschnitt nehmen wir an, dass das Suffix-Array zu $t = t_1 \cdots t_n \in \Sigma^*$ gehört und wir nach $s = s_1 \cdots s_m \in \Sigma^*$ suchen wollen.

5.4.1 Binäre Suche

Ein erster simpler Ansatz stellt die Implementierung mithilfe einer binären Suche dar. Wie sieht die Laufzeit dafür aus? Wie bei der binären Suche in einem Feld der Länge $n + 1$ üblich, müssen wir insgesamt $O(\log(n))$ Vergleiche auf Zeichenreihen vornehmen. Man beachte jedoch, dass die Zeichenreihen in einem Suffix-Array im Mittel sehr lang sein können. Im schlimmsten Fall müssen wir jeweils bis zu $O(m)$ Zeichen Vergleichen, um einen Vergleich auf Zeichenreihen beantworten zu können (wenn auch im Mittel deutlich weniger).

Theorem 5.18 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und $s = s_1 \cdots s_m \in \Sigma^*$. Eine Suche nach s in t mithilfe der binären Suche im Suffix-Array von t lässt sich in Zeit $O(m \log(n))$ erledigen.*

Das ist im Vergleich zum Suffix-Baum mit einer Laufzeit von $O(|\Sigma| \cdot m)$ asymptotisch deutlich langsamer. Man beachte aber, das $\log(n)$ für viele Zwecke kleiner gleich 40 ist.

5.4.2 Verbesserte binäre Suche

Man kann die Laufzeit bei der binären Suche mit dem folgenden Hilfsmittel jedoch noch beschleunigen.

Definition 5.19 Sei $t = t_1 \cdots t_n \in \Sigma$ und A das zugehörige Suffix-Array. Dann ist der längste gemeinsame Präfix (engl. longest common prefix) von den Positionen $i < j \in [0 : n]$ der Wert $\text{lcp}(i, j) = k$, wenn $t_{A[i]} \cdots t_{A[i]+k-1} = t_{A[j]} \cdots t_{A[j]+k-1}$ und wenn $t_{A[i]+k} < t_{A[j]+k}$.

Wir nehmen im Folgenden an, dass wir die längsten gemeinsamen Präfixe für das Positionspaar $i < j \in [1 : n]$ in einem Feld $\text{LCP}[i, j]$ gespeichert hätten.

Bezeichne $[L : R]$ das aktuelle Intervall der binären Suche im Suffix-Array A . Für dieses Intervall und die zu suchende Zeichenreihe s gilt dann $t^{A[L]} < s \leq t^{A[R]}$. Mit ℓ und r bezeichnen wir die Anzahl von Zeichenvergleiche, die wir mit $t^{A[L]}$ bzw. $t^{A[R]}$ und s erfolgreich ausgeführt haben, d.h. es gilt

$$\begin{aligned} s_1 \cdots s_\ell &= t_{A[L]} \cdots t_{A[L]+\ell-1} \quad \wedge \quad s_{\ell+1} \neq t_{A[L]+\ell}, \\ s_1 \cdots s_r &= t_{A[R]} \cdots t_{A[R]+r-1} \quad \wedge \quad s_{r+1} \neq t_{A[R]+r}. \end{aligned}$$

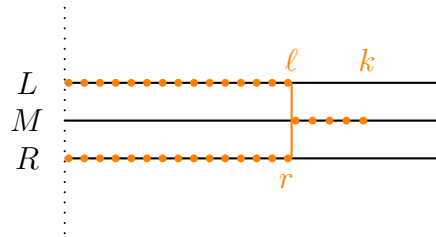
Zu Beginn testen wir, ob $s > t^{A[R]}$ ist (wobei wir annehmen, dass genau k Zeichenvergleiche ausgeführt werden). Falls ja, ist s kein Teilwort von t , andernfalls setzen wir $L = 0$ und $\ell = 0$ sowie $R = n$ und $r = k - 1$. Beachte, dass nach Definition für jedes Wort $s \in \Sigma^+$ die Beziehung $s > \$$ gilt.

Wir wählen dann gemäß der Strategie der binären Suche $M := \lceil \frac{L+R}{2} \rceil$ und versuchen effizient zu entscheiden, ob sich s im Bereich $[L : M]$ oder $[M : R]$ befindet. Wir überlegen uns jetzt, wie wir bei einer binären Suche geschickter vorgehen können. Dazu unterscheiden wir einige Fälle.

Fall 1: Es gilt $\ell = r$. Wir vergleichen s mit $t^{A[M]}$ ab Position $\ell + 1 = r + 1$ bis wir einen erfolglosen Zeichenvergleich an Position k durchgeführt haben (oder s gefunden haben).

Anschließend aktualisieren wir entweder L und ℓ oder R und r . War

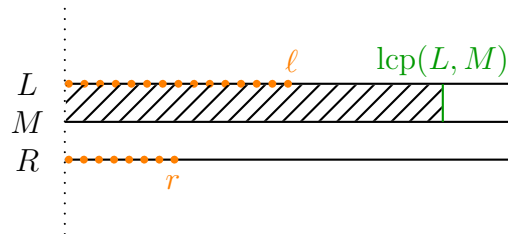
$$s_{\ell+1} \cdots s_k > t_{A[M]+\ell} \cdots t_{A[M]+k-1},$$

Abbildung 5.24: Skizze: Binäre Suche in $[L : R]$ bei $\ell = r$

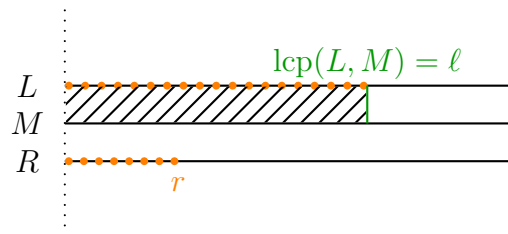
dann setzen wir $L = M$ und $\ell = k - 1$ und sonst $R = M$ und $r = k - 1$. Dies ist in Abbildung 5.24 noch einmal illustriert.

Fall 2: Es gelte jetzt $\ell > r$. Wir unterscheiden jetzt noch drei Unterfälle.

Fall 2a: Wir nehmen an, dass zusätzlich $\text{lcp}(L, M) > \ell$ gilt. Nach dem Prinzip der binären Suche gilt $s > t^{A[L]}$ und somit aufgrund des längsten gemeinsamen Präfixes auch $s > t^{A[M]}$. Wir setzen dann $L := M$ und lassen ℓ , R sowie r unverändert. Dies ist auch in Abbildung 5.25 illustriert. Der schraffierte Bereich zeigt dabei die Übereinstimmung von $t^{A[L]}$ mit $t^{A[M]}$

Abbildung 5.25: Skizze: Binäre Suche in $[L : R]$ mit $\ell > r$ und $\text{lcp}(L, M) > \ell$

Fall 2b: Es gilt $\ell = \text{lcp}(L, M)$. Jetzt vergleichen wir ganz normal $t^{A[M]}$ mit s ab Position $\ell + 1$ und aktualisieren L und ℓ bzw. R und r . Dies ist auch in Abbildung 5.26 illustriert. Der schraffierte Bereich zeigt dabei die Übereinstimmung von $t^{A[L]}$ mit $t^{A[M]}$.

Abbildung 5.26: Skizze: Binäre Suche in $[L : R]$ mit $\ell > r$ und $\text{lcp}(L, M) = \ell$

Fall 2c: Es gilt $\ell > \text{lcp}(L, M)$. Zunächst halten wir fest, dass $\text{lcp}(L, M) \geq r$ gelten muss. Da $t^{A[L]}$ und $t^{A[R]}$ aufgrund der bisherigen Suche nach s in den ersten $\min(\ell, r)$ Positionen übereinstimmen müssen, muss aufgrund der Sortierung des Suffix-Arrays auch $\text{lcp}(L, M) \geq \min(\ell, r)$ gelten. Nach Voraussetzung des Falles 2c gilt dann auch $\text{lcp}(L, M) \geq \min(\ell, r) = r$.

Weiter gilt nach Voraussetzung $s > t^{A[L]}$. Da das Suffix-Array die Suffixe in sortierter Reihenfolge enthält, muss $t^{A[M]} > t^{A[L]}$ und somit muss nach Definition des längsten gemeinsamen Präfixes gelten, dass $t_{A[M]+\text{lcp}(L,M)} > t_{A[L]+\text{lcp}(L,M)} = s_{\text{lcp}(L,M)}$. Somit liegt s im Intervall $[L : M]$ und wir aktualisieren $R := M$ und $r := \text{lcp}(L, M)$ und lassen L und ℓ unverändert. Dies ist auch in Abbildung 5.27 illustriert. Der schraffierte Bereich zeigt dabei die Übereinstimmung von $t^{A[L]}$ mit $t^{A[M]}$.

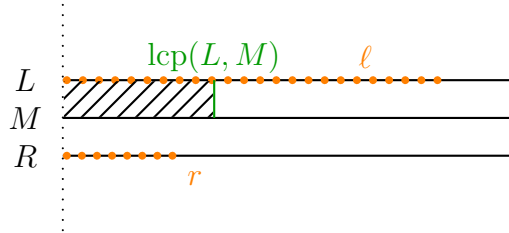


Abbildung 5.27: Skizze: Binäre Suche in $[L : R]$ mit $\ell > r$ und $\text{lcp}(L, M) < \ell$

Fall 3: Nun gilt $\ell < r$. Dieser Fall ist jedoch symmetrisch zum Fall 2.

Damit haben wir den folgenden Satz bewiesen.

Theorem 5.20 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei $s = s_1 \cdots s_m \in \Sigma^*$. Mit Hilfe des Suffix-Arrays und den LCP-Werten von t kann s in t mittels der verfeinerten binären Suche in Zeit $O(m + \log(n))$ gefunden werden.

Beweis: Die Korrektheit haben wir bereits gezeigt, wir müssen uns nur noch um die Laufzeit kümmern. Wir zählen hierzu die erfolgreichen und erfolglosen Zeichenvergleiche getrennt.

Für jeden erfolgreichen Zeichenvergleich wird r oder ℓ erhöht. Da r und ℓ nie erniedrigt werden und da $\ell, r \in [0 : m]$ gilt, kann es maximal $O(m)$ erfolgreichen Zeichenvergleiche geben.

Jeder erfolglose Zeichenvergleich liefert eine Intervallhalbierung. Daher kann es maximal $O(\log(n))$ erfolglose Zeichenvergleiche geben.

Da der Gesamtzeitbedarf proportional zur Anzahl der Zeichenvergleiche ist, beträgt der Gesamtzeitbedarf $O(m + \log(n))$. ■

Wir müssen uns nur noch überlegen, woher wir die LCP-Werte bekommen. Im schlimmsten Falle können $\Theta(n^2)$ verschiedene LCP-Anfragen gestellt werden. Man überlegt sich jedoch leicht, dass bei einer binärer Suche für einen festen Text t (aber variables Suchwort s) nicht alle LCP-Anfragen gestellt werden können. Insbesondere kann man sich überlegen, dass es unabhängig von Anfragezeichenreihe s nur $O(n)$ verschiedene mögliche LCP-Anfragen geben kann. Somit können diese $O(n)$ Anfragen vorab in einem Feld in linearer Zeit vorberechnet werden. Dazu wird insbesondere auch noch das folgende Lemma benötigt. Die Details überlassen wir dem Leser als Übungsaufgabe.

Lemma 5.21 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Dann gilt für $i < j \in [0 : n]$, dass*

$$\text{lcp}(i, j) = \min \{ \text{lcp}(k - 1, k) : k \in [i + 1 : j] \}.$$

Beweis: Übungsaufgabe. ■

21.12.06

5.4.3 Effiziente Berechnung der LCP-Tabelle

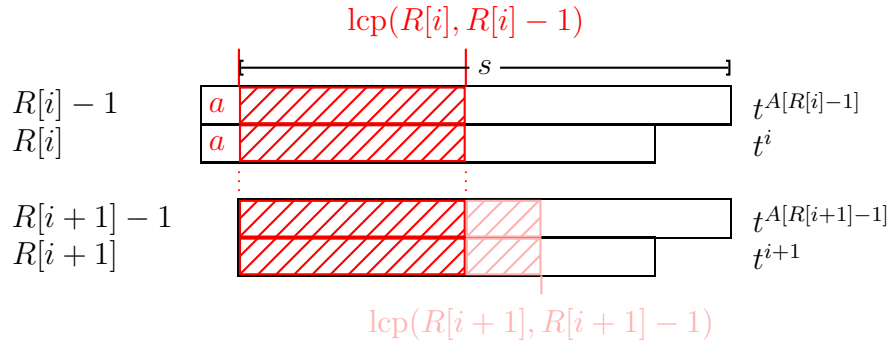
Für die effiziente Berechnung der allgemeinen LCP-Tabelle, wie sie für die binäre Suche benötigt wird, hilft eine spezielle LCP-Tabelle (siehe Lemma 5.21), für die wir jetzt ein effizientes Konstruktionsverfahren angeben wollen. Da wir diese LCP-Tabelle auch im nächsten Abschnitt noch benötigen werden, stellt sich diese Tabelle als relativ universell im Gebrauch mit Suffix-Arrays heraus.

Definition 5.22 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Dann ist das Feld L der Länge n die zu A gehörige LCP-Tabelle, wenn $L[i] = \text{lcp}(i - 1, i)$ für $i \in [1 : n]$.*

Zur Konstruktion arbeiten wir die Suffixe von t in absteigender Länge ab. Sei wiederum A das Suffix-Array von t und R das zugehörige inverse Suffix-Array, d.h. $A[R[i]] = i$. Dabei beachten wir, dass die folgende elementare Beziehung gilt

$$\text{lcp}(R[i + 1], R[i + 1] - 1) \geq \text{lcp}(R[i], R[i] - 1) - 1.$$

Zuerst halten wir fest, dass $t^{i+1} = t^{A[R[i+1]]}$ ein um ein Zeichen kürzeres Suffix von $t^i = t^{A[R[i]]}$ ist. Weiterhin ist $s = t^{A[R[i]-1]+1}$ ein Suffix von t , das zudem auch noch an den ersten $\text{lcp}(R[i], R[i] - 1) - 1$ Positionen mit t^{i+1} übereinstimmt und für das $s < t^{i+1}$ gilt. Somit muss entweder $t^{A[R[i+1]-1]} = s$ sein oder es müssen zumindest

Abbildung 5.28: Skizze: Beziehung zwischen $L[i]$ und $L[i+1]$

s und $t^{A[R[i+1]-1]}$ an den ersten $\text{lcp}(R[i], R[i-1]) - 1$ Positionen übereinstimmen (andernfalls wäre A falsch sortiert). Diese Beziehung ist auch in Abbildung 5.28 noch einmal schematisch dargestellt.

Daraus können wir sofort den sehr einfachen, in Abbildung 5.29 dargestellten Algorithmus ableiten.

Theorem 5.23 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Die zugehörige LCP-Tabelle L kann in Zeit $O(n)$ konstruiert werden.

Beweis: Die Korrektheit haben wir schon gezeigt. Wir müssen uns nur noch die Laufzeit überlegen. Die äußere for-Schleife ohne die innere while-Schleife wird genau n Mal durchlaufen. Somit ist Laufzeit ohne die Inkrementierungen von k und die Anzahl der Zeichenvergleiche innerhalb von t durch $O(n)$ beschränkt. Wie oft kann k

```

LCP (int A[]; int R[]; char t[]; int n)
{
    k := 0;
    for (i = 1; i ≤ n; i++)
        if (R[i] > 0) /* always true, since R[i] = 0 ⇔ i = n + 1 */
        {
            j := A[R[i] - 1]; /* t^j is in front of t^i */
            while (t_{i+k} = t_{j+k}) k++;
            /* Es gilt jetzt t_{i+k} ≠ t_{j+k} */
            L[R[i]] = k;
            k = max(0, k - 1);
        }
}

```

Abbildung 5.29: Algorithmus: Berechnung der LCP-Tabelle L

überhaupt inkrementiert werden? Da $k \in [0 : n]$ gilt und k maximal n mal dekrementiert werden kann (wenn k nicht sowieso schon 0 ist), kann k maximal $2n$ mal inkrementiert werden. Also ist die Gesamtlaufzeit $O(n)$. ■

Dieser Algorithmus benötigt neben dem Platz für den eigentlichen Text t , das Suffix-Array A und die LCP-Tabelle L noch den Platz für das inverse Suffix-Array R . Bei einer normalen Implementierung in der Praxis bedeutet dies für einen Text der Länge $n < 2^{32}$ einen Platzbedarf von $13n$ Bytes.

Mit einem Trick kann man auf das Feld für das inverse Suffix-Array R verzichten und kommt mit nur $9n$ Bytes aus, was bei normalen Anwendungen optimal ist. Für Details verweisen wir den Leser auf die Originalliteratur von Manzini. Dort sind auch weitere Platzreduktionen angegeben, wenn beispielsweise anschließend das Suffix-Array selbst auch nicht mehr benötigt wird.

5.5 Enhanced Suffix-Arrays (*)

In diesem Abschnitt wollen wir jetzt noch zeigen, wie man die Algorithmen für Suffix-Bäume auf Suffix-Arrays übertragen kann. Dabei wird die LCP-Tabelle eine bedeutende Rolle spielen. In diesem Abschnitt nehmen wir jetzt jedoch an, dass $\$ \notin \Sigma$ jetzt das größte Zeichen ist, d.h. es gilt $a < \$$ für alle $a \in \Sigma$.

5.5.1 LCP-Intervalle

Zunächst definieren die so genannten LCP-Intervalle, die grundlegend für das Folgende sind.

Definition 5.24 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Ein Intervall $[i : j]$ mit $i < j \in [0 : n]$ heißt ein LCP-Intervall vom Typ ℓ oder kurz ein ℓ -Intervall, wenn

1. $L[i] < \ell$,
2. $L[k] \geq \ell$ für alle $k \in [i + 1 : j]$,
3. $L[k] = \ell$, für mindestens ein $k \in [i + 1 : j]$ und
4. $L[j + 1] < \ell$.

Hierbei gelte $L[0] = L[n + 1] = -1$.

Wir werden später auch Intervalle der Länge 1 zulassen.

Notation 5.25 ℓ - $[i : j]$ bezeichnet das ℓ -Intervall $[i : j]$.

In Abbildung 5.30 sind die ℓ -Intervalle für unser Beispielwort MISSISSIPPI noch einmal illustriert. Das 0-Intervall $[0 : n]$ ist hierbei nicht dargestellt.

i	$A[i]$	$L[i]$	$C[i]$			$t^{A[i]}$
			down	up	next	
0	8	-1	4			IPPI\$
1	5	1	2		3	ISSIPPI\$
2	2	4				ISSISSIPPI\$
3	11	1		2		I\$
4	1	0		1	5	MISSISSIPPI\$
5	10	0	6		7	PI\$
6	9	1				PPI\$
7	7	0	9	6	11	SIPPI\$
8	4	2				SISSIPPI\$
9	6	1	10	8		SSIPPI\$
10	3	3				SSISSIPPI\$
11	12	0		9		\$

Abbildung 5.30: Beispiel: Suffix-Array, LCP-Tabelle, Child-Tabelle und alle ℓ -Intervalle für MISSISSIPPI (außer dem 0-Intervall)

Lemma 5.26 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Für jedes $k \in [1 : n]$ mit $L[k] = \ell$ existiert ein ℓ -Intervall $[i : j]$ mit $k \in [i + 1 : j]$

Beweis: Wir definieren zunächst einmal $i := \max \{p < k : L[p] < L[k] = \ell\}$ und $j := \min \{p \geq k : L[p + 1] < L[k] = \ell\}$. Diese existieren nach der erweiterten Definition von L immer ($L[0] = L[n + 1] = -1$). Wir behaupten jetzt, dass $[i : j]$ das gesuchte ℓ -Intervall ist. Nach Definition von i gilt $L[i] < \ell$, also gilt Bedingung 1 der Definition eines LCP-Intervalls. Analog gilt nach Definition von j gilt $L[j + 1] < \ell$, also gilt Bedingung 4 der Definition eines LCP-Intervalls. Weiterhin gilt $k \in [i + 1 : j] \neq \emptyset$ und $L[k] = \ell$, also Bedingung 3 der Definition. Da wir i maximal und j minimal unter der Bedingung $L[p] < L[k] = \ell$ wählen, gilt $L[r] \geq L[k] = \ell$ für alle $r \in [i + 1 : j]$ und somit die Bedingung 2. ■

Definition 5.27 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Wenn $[i : j]$ ein ℓ -Intervall ist, dann heißt $k \in [i + 1 : j]$ ein ℓ -Index, wenn $L[k] = \ell$. Die Menge aller ℓ -Indizes eines ℓ -Intervalls $[i : j]$ wird mit $I_\ell[i : j]$ bezeichnet.

Definition 5.28 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Ein m -Intervall $[i' : j']$ ist in einem ℓ -Intervall $[i : j]$ enthalten, wenn $i \leq i' < j' \leq j$ und $m > \ell$ gilt.

Wenn $[i' : j']$ ein in $[i : j]$ enthaltenes Intervall ist und es kein anderes ℓ' -Intervall gibt, das in $[i : j]$ enthalten ist und das $[i' : j']$ enthält, dann heißt $[i' : j']$ ein Kind-Intervall von $[i : j]$. Umgekehrt heißt $[i : j]$ das Elter-Intervall von $[i' : j']$.

Mit Hilfe dieser Enthaltensein-Relation können wir zu den LCP-Intervallen auch eine Baumstruktur generieren. Betrachten wir zunächst noch einmal den Suffix-Baum für MISSISSIPPI in Abbildung 5.31. Konstruieren wir für das zugehörige Suffix-Array und den zugehörigen LCP-Werten den so genannten *LCP-Intervall-Baum*, der auf der Enthaltensein-Relation der LCP-Intervall definiert ist. Beachte dabei, dass $[0 : n]$ immer ein 0-Intervall ist. Weiter ergänzen wir den so konstruierten LCP-Intervall-Baum um die so genannten *Singletons*, das sind die einelementigen Mengen von $[0 : n]$. In der Regel werden wir auch Singletons als Kind-Intervalle zulassen. In Abbildung 5.32 ist der LCP-Intervall-Baum für unser Beispielwort MISSISSIPPI angegeben.

Man sieht sofort die Ähnlichkeit zum entsprechenden Suffix-Baum. In den LCP-Intervallen sind also bereits alle wesentlichen Informationen des Suffix-Baumes (ohne die Blätter) enthalten. In Abbildung 5.32 entsprechen die schwarzen Kanten der Enthaltensein-Relation der Suffix-Bäume und die türkisen Kanten den Kanten zu den Blätter, die nach Definition ja keine ℓ -Intervalle sind. Die roten Kanten entsprechen den Kanten auf das älteste Kind bzw. auf das nächstältere Geschwister gemäß der Definition in der Darstellung eines Suffix-Baumes. Im Folgenden müssen wir nur versuchen diese roten Kanten aus den LCP-Intervallen zu rekonstruieren.

Lemma 5.29 Sei $[i : j]$ ein ℓ -Intervall und seien $i_1 < i_2 < \cdots < i_k$ die ℓ -Indizes des ℓ -Intervalls $[i : j]$, dann sind die Kind-Intervalle von $[i : j]$ gerade $[i : i_1 - 1]$, $[i_1 : i_2 - 1]$, $\dots$, $[i_k : j]$.

Wir merken hier noch an, dass $[i_p : i_{p+1} - 1]$ für $p \in [1 : k - 1]$ durchaus ein Singleton sein kann.

Beweis: Sei $[r : s]$ eines der Intervalle $[i : i_1 - 1]$, $[i_1 : i_2 - 1]$, $\dots$, $[i_k : j]$.

Ist $[r : s]$ ein Singleton (d.h. $r = s$), dann ist $[r : s]$ nach unserer Konvention ein Kind-Intervall.

Gelte also jetzt $r < s$. Wir definieren zunächst $m = \min \{L[x] : x \in [r + 1 : s]\}$ sowie $x' = \operatorname{argmin} \{L[x] : x \in [r + 1 : s]\}$. Nach Definition des ℓ -Intervalls gilt $m > \ell$.

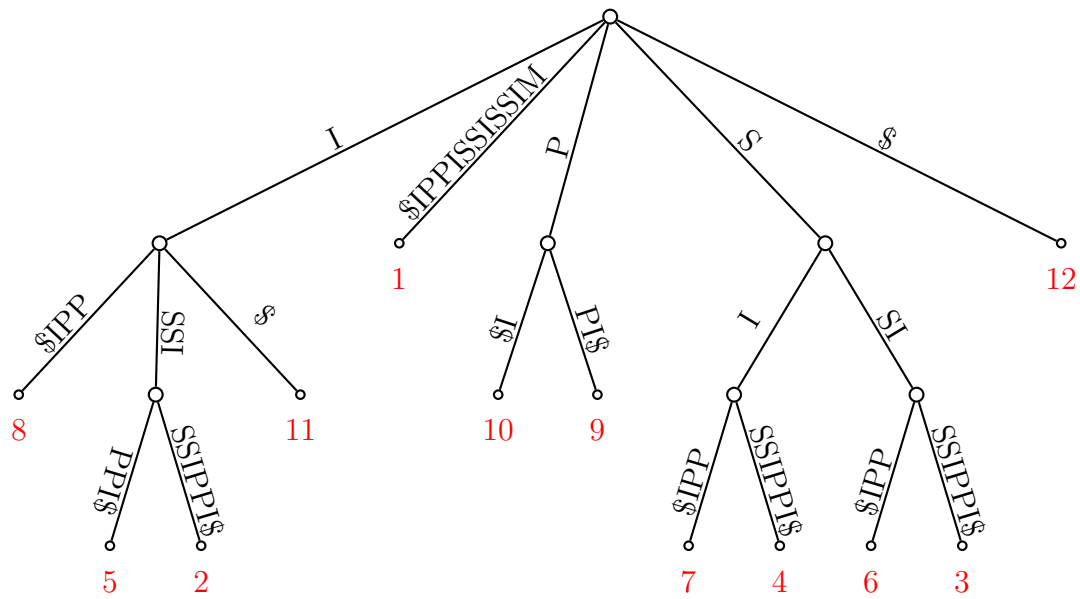


Abbildung 5.31: Beispiel: Suffix-Baum für MISSISSIPPI

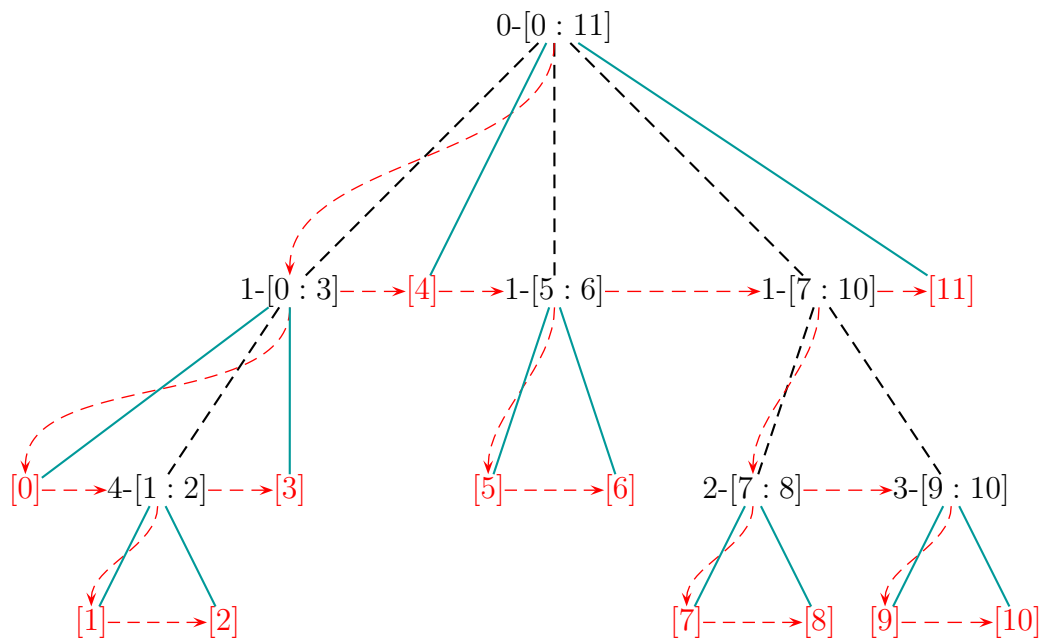


Abbildung 5.32: Beispiel: LCP-Intervall-Baum für MISSISSIPPI

Nach Lemma 5.26 existiert ein m -Intervall I , dass die Position x' umfasst. Nach Definition der ℓ -Indizes des ℓ -Intervalls $[i : j]$ und nach Wahl von m muss diese Intervall I gerade gleich $[r : s]$ sein.

Weiterhin ist $[r : s]$ in $[i : j]$ enthalten. Da weiterhin nach Definition der ℓ -Indizes $L[i_1] = L[i_2] = \dots = L[i_k] = \ell$ gilt, kann kein anderes Intervall in $[i : j]$ enthalten sein, welches $[r : s]$ enthält.

Offensichtlich sind die Intervalle $[i : i_1 - 1]$, $[i_1 : i_2 - 1]$, $\dots$, $[i_k : j]$ alle möglichen Kind-Intervalle, die $[i : j]$ besitzen kann. ■

5.5.2 Die Child-Tabelle und ihre Eigenschaften

Um jetzt das Suffix-Array genauso verwenden zu können wie den entsprechenden Suffix-Baum, müssen wir die Menge der ℓ -Indizes eines LCP-Intervalls in linearer Zeit berechnen können. Dazu dient die folgende Notation der Child-Tabelle.

Definition 5.30 Sei $t = t_1 \dots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Dann ist die Child-Tabelle wie folgt definiert:

$$\begin{aligned} C[i].up &= \min \{q < i : L[q] > L[i] \wedge \forall k \in [q+1 : i-1] : L[k] \geq L[q]\}; \\ C[i].down &= \max \{q > i : L[q] > L[i] \wedge \forall k \in [i+1 : q-1] : L[k] > L[q]\}; \\ C[i].next &= \min \{q > i : L[q] = L[i] \wedge \forall k \in [i+1 : q-1] : L[k] > L[q]\}. \end{aligned}$$

Hierbei gilt $\min \emptyset = \max \emptyset = \perp$, wobei $\perp$ für undefiniert steht.

Im Wesentlichen ist das Suffix-Array zusammen mit der LCP-Tabelle und der Child-Tabelle das *Enhanced-Suffix-Array*. In Abbildung 5.33 sind die LCP-Werte zur Definition der Child-Tabelle noch einmal illustriert. Der grüne Bereich gibt dabei den Bereich des LCP-Wertes von $L[q+1]$ bzw. $L[q-1]$ an.

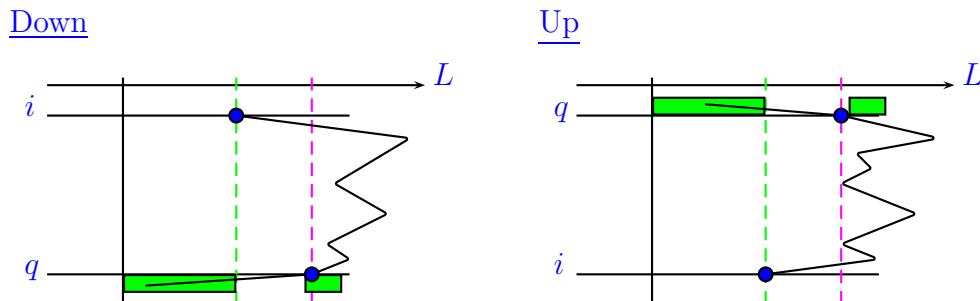


Abbildung 5.33: Skizze: Die L -Gebirge in der Definition von $C[i].down$ und $C[i].up$

Lemma 5.31 Für jedes ℓ -Intervall $[i : j]$ gelten die folgenden Aussagen:

- 1) $C[j+1].up \in [i+1 : j]$ oder $C[i].down \in [i+1 : j]$.
- 2) $C[j+1].up$ speichert den ersten ℓ -Index von $[i : j]$, wenn $C[j+1].up \in [i+1 : j]$.
- 3) $C[i].down$ speichert den ersten ℓ -Index von $[i : j]$, wenn $C[i].down \in [i+1 : j]$.

Beweis: zu 1.) Sei $\ell' := L[j+1]$. Da $[i : j]$ ein ℓ -Intervall ist, gilt $\ell' < \ell$. Weiter gilt $L[k] \geq \ell$ für alle $k \in [i+1 : j]$.

Gilt $L[i] < \ell'$, dann ist $C[j+1].up \in [i+1 : j]$ nach Definition von $C[j+1].up$. Das L -Gebirge im Intervall $[i : j+1]$ ist für diesen Fall Abbildung 5.34 a) noch einmal illustriert.

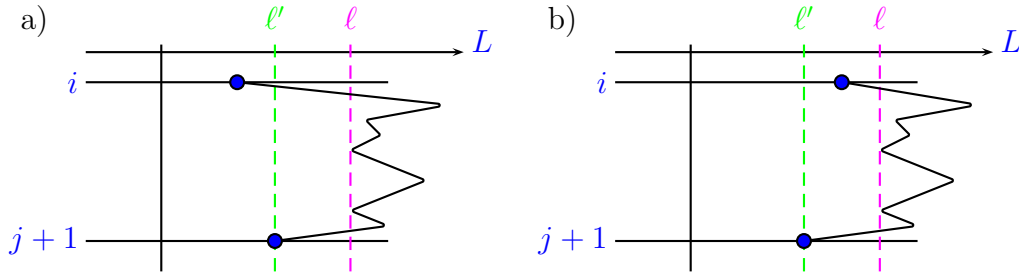


Abbildung 5.34: Skizze: Die L -Gebirge in Behauptung 1

Gilt andererseits $L[i] \geq \ell'$, dann ist $C[i].down \in [i+1 : j]$ nach Definition von $C[i].down$. Beachte, dass nach Definition des ℓ -Intervalls $[i : j]$ in jedem Falle $L[i] < \ell$ gelten muss. Das L -Gebirge im Intervall $[i : j+1]$ ist für diesen Fall Abbildung 5.34 b) noch einmal illustriert.

zu 2.) Wenn $C[j+1].up \in [i+1 : j]$, dann gilt:

$$\begin{aligned} & C[j+1].up \\ &= \min \{q \in [i+1 : j] : L[q] > L[j+1] \wedge \forall k \in [q+1 : j] : L[k] \geq L[q]\} \end{aligned}$$

Da $[i : j]$ ein ℓ -Intervall ist und da $L[q] \geq \ell > L[j+1]$ für $q \in [i+1 : j]$ sein muss, erhalten wir weiter

$$\begin{aligned} &= \min \{q \in [i+1 : j] : \forall k \in [q+1 : j] : L[k] \geq L[q]\} \\ &= \min I_\ell[i : j]. \end{aligned}$$

zu 3.) Sei i_1 der erste ℓ -Index des ℓ -Intervalls $[i : j]$. Dann gilt $L[i_1] = \ell > L[i]$ und für alle $k \in [i+1 : i_1 - 1]$ gilt $L[k] > \ell = L[i_1]$. Weiter gilt für alle $q \in [i_1 + 1 : j]$, dass

$L[q] \geq \ell > L[i]$, aber nicht $L[i_1] > L[q]$. Somit wird das Maximum an der Position i_1 angenommen. ■

Korollar 5.32 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A , L bzw. C das zugehörige Suffix-Array, die zugehörige LCP-Tabelle bzw. die zugehörige Child-Tabelle. Für jedes LCP-Intervall $[i : j]$ gilt:

$$\text{lcp}(i, j) = \begin{cases} L[C[j+1].\text{up}] & \text{falls } C[j+1].\text{up} \in [i+1 : j], \\ L[C[i].\text{down}] & \text{sonst.} \end{cases}$$

Beweis: Die Korrektheit folgt unmittelbar aus dem vorherigen Lemma zusammen mit Lemma 5.21. ■

Mit Hilfe dieser fundamentalen Eigenschaften der Child-Tabelle können wir den in Abbildung 5.35 angegebenen Algorithmus GETNEXTCHILDINTERVAL zum Auffinden eines Kind-Intervalls entwerfen. Als Eingabe erhält der Algorithmus ein ℓ -Intervall $[i : j]$ und einen Parameter $k \in \{i\} \cup I_\ell[i : j]$. Ist $k = i$ so liefert der Algorithmus den ersten ℓ Index von $[i : j]$, ansonsten den kleinsten ℓ -Index, der größer als k ist.

Somit können wir zu einem ℓ -Intervall alle Kinder-Intervalle in konstanter Zeit pro Intervall ermitteln. Dazu werden wir später noch zeigen, dass wir die Child-Tabelle in linearer Zeit konstruieren können.

Lemma 5.33 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. C das zugehörige Suffix-Array bzw. die zugehörige Child-Tabelle. Für ein gegebenes ℓ -Intervall $[i : j]$ von A und einen Index $k \in \{i\} \cup I_\ell[i : j]$ kann der kleinste ℓ -Index von $\ell\text{-}[i : j]$, der größer als k ist, in konstanter Zeit bestimmt werden.

5.5.3 Optimale Suche in Enhanced Suffix-Arrays

Basierend auf dem Algorithmus zum Auffinden aller Kind-Intervalle eines ℓ -Intervalls geben wir in diesem Abschnitt einen optimalen Algorithmus zum Suchen nach Teilwörtern basierend auf Suffix-Arrays an. Der Algorithmus ist in Abbildung 5.35 angegeben.

Die Invariante im Algorithmus ist die folgende: Während der Suche nach einem Wort $s = s_1 \cdots s_m$ gibt die Boolesche Variable **prefix** an, ob ein Präfix der Länge p von s ein Teilwort t enthalten ist oder nicht. Weiterhin ist $[i : j]$ ein p -Intervall (außer im

FIND (string s)

```
{
    interval  $[i : j] = [0 : n]$ ;
    int  $p = 0$ ; bool prefix = TRUE; /* prefix is TRUE iff  $s_1 \cdots s_p$  is a prefix of  $t^i$  */
    while ( $([i : j] \neq \perp) \ \&\& \ (p < m) \ \&\& \ (\text{prefix})$ )
    {
        if ( $i < j$ )          /* a Child-Interval */
        {
             $k = \min(\text{LCP}(i, j), m)$ ;
            prefix =  $(t_{A[i]+p} \cdots t_{A[i]+k-1} = s_{p+1} \cdots s_k)$ ;
             $p = k$ ;
        }
        elseif ( $i = j$ )      /* a Singleton */
        {
            prefix =  $(t_{A[i]+p} \cdots t_{A[i]+m-1} = s_{p+1} \cdots s_m)$ ;
             $p = m$ ;
        }
         $[i : j] = \text{getChildIntervalByChar}(i, j, s_{p+1})$ ;
    }
    output (prefix)? $[i : j] : \perp$ ;
}
```

GETCHILDINTERVALBYCHAR (int i, j ; char c)

```
{
    interval  $[i' : j'] = \text{getNextChildInterval}(i, j, i)$ ;
    while ( $(t_{A[i']} \neq c) \ \&\& \ (j' < j)$ )
         $[i' : j'] = \text{getNextChildInterval}(i, j, j' + 1)$ ;
    return  $(t_{A[i']} = c) ? ([i' : j']) : (\perp)$ ;
}
```

GETNEXTCHILDINTERVAL (int i, j, k)

```
{
    if ( $k = i$ ) return  $(C[j+1].\text{up} \in [i+1 : j]) ? [i : C[j+1].\text{up} - 1] : [i : C[i].\text{down} - 1]$ ;
    else return  $(C[k].\text{next} \neq \perp) ? [k : C[k].\text{next} - 1] : [k : j]$ ;
}
```

LCP(int i, j)

```
{
    return  $(C[j+1].\text{up} \in [i+1 : j]) ? L[C[j+1].\text{up}] : L[C[i].\text{down}]$ ;
}
```

Abbildung 5.35: Algorithmus: Suche s im Enhanced Suffix-Array von t

Falle $p = m$), dass nur Suffixe von t enthält, die mit dem Präfix $s_1 \cdots s_p$ beginnen. Halten wir das Ergebnis im folgenden Satz fest.

Theorem 5.34 *In einem Suffix-Array zusammen mit den Tabellen L und C kann in Zeit $O(|\Sigma| \cdot m)$ festgestellt werden, ob $s \in \Sigma^m$ in $t \in \Sigma^*$ enthalten ist.*

Beweis: Die Korrektheit folgt aus den Zeichenvergleichen. Die Laufzeit folgt aus der Tatsache, dass alle anderen Funktionen in konstanter Zeit durchgeführt werden können und die Laufzeit somit proportional zur Anzahl der Zeichenvergleiche ist, also $O(m)$ beträgt. ■

5.5.4 Berechnung der Child-Tabelle

Wir müssen uns jetzt nur noch überlegen, wie wir die Child-Tabelle in linearer Zeit konstruieren können. Die Algorithmen hierzu sind in der Abbildung 5.36 angegeben.

Sowohl im Algorithmus COMPUTEUPDOWNTABLE als auch im Algorithmus COMPUTENEXTTABLE gilt die folgende Invariante. Sind $0, i_1, \dots, i_k$ die Element auf dem Stack S , dann gilt zum einen $0 < i_1 < \dots < i_k$ und zum anderen $L[i_1] \leq \dots \leq L[i_k]$. Weiterhin gilt für alle $k \in [i_j + 1 : i_{j+1} - 1]$ für zwei aufeinander folgende Stackelemente $i_j < i_{j+1}$ mit $L[i_j] < L[i_{j+1}]$, dass $L[k] > L[i_{j+1}]$.

Lemma 5.35 *Der Algorithmus ComputeUpDownTable ermittelt die korrekten Up- und Down-Werte der Child-Tabelle in linearer Zeit.*

Beweis: Wenn $C[S.top()].down$ gesetzt wird, gilt $L[k] \leq L[S.top()] < L[lastIdx]$ sowie $S.top() < lastIdx < k$. Es gilt also $lastIdx < k$ und $L[lastIdx] > L[S.top()]$ sowie aufgrund der Invariante $L[p] > L[lastIdx]$ für alle $p \in [S.top() + 1 : lastIdx - 1]$. Somit befindet sich $lastIdx$ in der Menge, deren Maximum gerade $C[S.top()].down$ ist. Angenommen $lastIdx$ wäre nicht das Maximum, sondern $q' \in [lastIdx + 1 : k - 1]$. Nach Definition von $C[S.top()].down$ muss $L[lastIdx] > L[q']$ sein. Somit muss $lastIdx$ vom Stack entfernt worden sein, als der Index q' betrachtet wurde, was den gewünschten Widerspruch liefert.

Wenn $C[k].up$ gesetzt wird, dann gilt insbesondere $L[S.top()] \leq L[k] < L[lastIdx]$ und $S.top() < lastIdx < k$. Es gilt also $lastIdx < k$ und $L[lastIdx] > L[k]$ sowie aufgrund der Invariante $L[p] \geq L[lastIdx]$ für alle $p \in [lastIdx + 1 : i - 1]$. Somit befindet sich $lastIdx$ in der Menge, deren Minimum gerade $C[k].up$ ist. Angenommen $lastIdx$ wäre nicht das Minimum, sondern $q' \in [S.top() + 1 : lastIdx - 1]$. Nach

COMPUTEUPDOWNTABLE ()

```

{
    int lastIdx = -1;
    stack S = empty();
    S.push(0);
    for (k = 1; k ≤ n; k++)
    {
        while (L[k] < L[S.top()])
        {
            lastIdx := S.pop();
            if ((L[k] ≤ L[S.top()]) && (L[S.top()] < L[lastIdx]))
                C[S.top()].down := lastIdx;
        }
        /* Now L[k] ≥ L[S.top()] */
        if (lastIdx ≠ -1)
        {
            C[k].up = lastIdx;
            lastIdx = -1;
        }
        S.push(k);
    }
}

```

COMPUTENEXTTABLE ()

```

{
    stack S = empty();
    S.push(0);
    for (k = 1; k ≤ n; k++)
    {
        while (L[k] < L[S.top()])
            S.pop();
        if (L[k] = L[S.top()])
        {
            lastIdx = S.pop();
            C[lastIdx].next = k;
        }
        S.push(k);
    }
}

```

Abbildung 5.36: Algorithmen: Berechnung der Up- und Down- sowie Next-Werte

Definition von $C[k].up$, muss $L[\text{lastIdx}] \geq L[q'] > L[k] \geq L[S.top()]$ gelten. Somit gilt $q \in [S.top() + 1 : \text{lastIdx} - 1]$, was den gewünschten Widerspruch liefern. ■

Die Korrektheit des folgenden Lemmas folgt direkt aus der Inspektion des Algorithmus in Abbildung 5.36.

Lemma 5.36 *Der Algorithmus `ComputeNextTable` ermittelt die korrekten Next-Werte der Child-Tabelle in linearer Zeit.*

5.5.5 Komprimierte Darstellung der Child-Tabelle

Wir wollen jetzt noch zeigen, dass die Child Tabelle in einem Feld mit n statt $3n$ Einträgen gespeichert werden kann.

Lemma 5.37 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. C das zugehörige Suffix-Array bzw. die zugehörige Child-Tabelle. Für jedes $i \in [0 : n]$ gilt, dass $C[i].next = \perp$, wenn $C[i+1].up \neq \perp$ gilt.*

Beweis: Ist $C[i+1].up \neq \perp$, dann ist insbesondere $L[i] > L[i+1]$. Dann kann aber $C[i].next$ nicht definiert sein. ■

Somit kann das Up-Feld im entsprechenden Next-Feld gespeichert werden. Gilt bei einer Anfrage $C[i].next > i$, dann wird tatsächlich das Next-Feld gespeichert, andernfalls das Up-Feld an Position $i+1$.

Wir werden jetzt noch zeigen, dass auch das Down-Feld im Next-Feld gespeichert werden kann. Wir bemerken zuerst, dass wir für ein ℓ -Intervall $[i : j]$ das Feld $C[i].down$ nur dann gespeichert werden muss, wenn $C[j+1].up \notin [i+1 : j]$.

Lemma 5.38 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. C das zugehörige Suffix-Array bzw. die zugehörige Child-Tabelle. Für jedes ℓ -Intervall $[i : j]$, für das $C[j+1].up \notin [i+1 : j]$ gilt, gilt $C[i].next = \perp$ und $C[i+1].up = \perp$.*

Beweis: Wenn $[i : j]$ ein ℓ -Intervall ist, gilt $L[i] < \ell$, $L[j+1] < \ell$ sowie $L[k] \geq \ell$ für $k \in [i+1 : j]$. Wenn weiter $C[j+1].up \notin [i+1 : j]$ gilt, dann muss $L[i] > L[j+1]$. Dann gilt $L[k] \geq \ell > L[i]$ für alle $k \in [i+1 : j]$ und $L[j+1] < L[i]$. Also muss $C[i].next = \perp$ gelten.

Da $L[i] < \ell$ und $L[i+1] \geq \ell$ gilt, muss $C[i+1].up = \perp$ sein. ■

Somit können wir auch den Wert $C[i].\text{down}$, wenn er denn nicht bereits durch $C[j+1].\text{up}$ gegeben ist, in $C[i].\text{next}$ speichern. Wir müssen jetzt nur noch unterscheiden, ob $C[i].\text{next}$ den Next- oder ein Up-Eintrag enthält, wenn $C[i].\text{next} > i$ ist. Gilt $L[i] = L[C[i].\text{next}]$, dann wird der Next-Wert gespeichert, gilt $L[i] < L[C[i].\text{next}]$, dann wird der zugehörige Down-Wert gespeichert.

i	$A[i]$	$L[i]$	$C[i]$			$t^{A[i]}$
			down	up	next	
0	8	-1	4			IPPI\$
1	5	1	2		3	ISSIPPI\$
2	2	4				ISSISSIPPI\$
3	11	1		2		I\$
4	1	0		1	5	MISSISSIPPI\$
5	10	0	6		7	PI\$
6	9	1				PPI\$
7	7	0	9	6	11	SIPPI\$
8	4	2				SISSIPPI\$
9	6	1	10	8		SSIPPI\$
10	3	3				SISSIPPI\$
11	12	0		9		\$

Abbildung 5.37: Beispiel: Suffix-Array, LCP-Tabelle, komprimierter Child-Tabelle und alle ℓ -Intervalle für MISSISSIPPI (außer dem 0-Intervall)

Definition 5.39 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A , L bzw. C das zugehörige Suffix-Array, die zugehörige LCP-Tabelle bzw. die zugehörige Child-Tabelle. Dann ist (A, L, C) das zu t gehörige Enhanced-Suffix-Array.

Wie man leicht sieht, kann das Enhanced Suffix-Array für $n < 2^{32}$ mit $12n$ Bytes realisiert werden. Da die LCP-Tabelle hauptsächlich kleine Werte enthält, kann diese auch durch ein Feld mit n Bytes realisiert werden. Analog kann auch die Child-Tabelle mit einem Feld der Größe von n Bytes realisiert werden, wenn statt der tatsächlichen Indexpositionen jeweils die relativen Offsets zur Position i gespeichert werden.

Für Werte außerhalb des Intervalls $[0 : 254]$ bzw. $[-127 : 127]$, wird im Feld 255 bzw. -128 gespeichert und in einem Extrafeld das Paar aus Indexposition und Wert gespeichert, auf die dann mit einer binären Suche im Ausnahmefall schnell zugegriffen werden kann. Dann benötigt das Enhanced-Suffix-Array etwa $6n$ Bytes plus den Text t sowie Platz für die übergroßen Werte (außerhalb des Intervalls $[0 : 254]$ bzw. $[-127 : 127]$).

Die Suche in der Tabelle mit den übergroßen Werten kann mithilfe einer binären Suche relativ effizient ausgeführt werden. Die worst-case Laufzeit sinkt zwar bei dieser speicherplatzsparenden Darstellung, macht in der Praxis aber nicht allzuviel aus, wenn die zugehörigen Texte nur wenige übergroße Werte implizieren.

5.5.6 Simulation von Suffix-Baum-Algorithmen auf Suffix-Arrays

Somit lassen sich leicht DFS-Traversierungen auf einem Suffix-Array vornehmen, wie beispielsweise in Abbildung 5.38 angegeben. Auch BFS-Traversierung können ähnlich durchgeführt werden, die Details seien dem Leser zur Übung überlassen.

```

TRAVERSE (interval  $[i : j]$ )
{
    /* traverse suffix-array recursively in a depth-first-search manner */

     $[i' : j'] = \text{getNextChildInterval}(i, j, i)$ ;
    while ( $[i' : j'] \neq \perp$ )
    {
        traverse( $[i' : j']$ )
         $[i' : j'] = \text{getNextChildInterval}(i, j, j' + 1)$ ;
    }
}

```

Abbildung 5.38: Algorithmus: Traversierung des Suffix-Arrays

Es bleibt noch zu überlegen, wie man Suffix-Links in die Enhanced-Suffix-Arrays integrieren kann.

Definition 5.40 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Für $t^{A[i]} = aw\$$ mit $a \in \Sigma$ und $w \in \Sigma^*$ ist $\sigma(i) = j$ mit $t^{A[j]} = w\$$ der Suffix-Link an Position i .

Wie man leicht sieht, gilt die folgende Beobachtung.

Beobachtung 5.41 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. R das zugehörige Suffix-Array bzw. das zugehörige inverse Suffix-Array. Dann gilt $\sigma(i) = R[A[i] + 1]$, sofern $A[i] \leq n$

Wir haben somit nur Suffix-Links für Singletons definiert. Diese zeigen selbst jedoch wieder auf Singletons, namentlich das Intervall $[R[A[i] + 1] : R[A[i] + 1]]$, sofern

$A[i] \neq n + 1$ ist. Ist $A[i] = n + 1$, so beschreibt das Singleton das Suffix $\$,$ dessen Suffix-Link das Wurzel-Intervall $0-[0 : n]$ ist.

Für LCP-Intervalle können wir das folgende Lemma beweisen.

Lemma 5.42 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Zu jedem ℓ -Intervall $[i : j]$ mit $\ell > 0$ existiert genau ein kleinstes $(\ell - 1)$ -Intervall $[r : s]$ mit $r \leq \sigma(i) < \sigma(j) \leq s$.*

Beweis: Da $[i : j]$ ein ℓ -Intervall ist, gilt offensichtlich für $k \in [\sigma(i) + 1 : \sigma(j)]$, dass $L[k] \geq \ell - 1$. Da $[i : j]$ ein ℓ -Intervall ist, muss es weiter einen ℓ -Index $k \in [i + 1 : j]$ mit $L[k] = \ell$ geben. Deshalb muss es auch einen Index $k' \in [\sigma(i) + 1 : \sigma(j)]$ mit $L[k'] = \ell - 1$ geben. Nach Lemma 5.26 existiert dann ein $(\ell - 1)$ -Intervall $[r : s]$, das nach Konstruktion $[\sigma(i) : \sigma(j)]$ umfasst. ■

Mit diesem Existenz-Satz können wir nun die folgende Definition für Suffix-Links von LCP-Intervallen wagen.

Definition 5.43 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Zu jedem ℓ -Intervall $[i : j]$ mit $\ell > 0$ ist das kleinste $(\ell - 1)$ -Intervall $[r : s]$ mit $r \leq \sigma(i) < \sigma(j) \leq s$ das Suffix-Link-Intervall von $[i : j]$.*

Der Beweis des folgenden Korollars ergibt sich sofort aus einer genauen Inspektion des Beweises von Lemma 5.42.

Korollar 5.44 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Sei $[i : j]$ ein ℓ -Intervall mit $\ell > 0$ und sei $[r : s]$ sein Suffix-Link-Intervall. Dann existiert ein $(\ell - 1)$ -Index $k \in I_{\ell-1}[r : s]$ und $k \in [\sigma(i) + 1 : \sigma(j)]$.*

Sei nun $[i : j]$ ein ℓ -Intervall und $[r : s]$ das zugehörige Suffix-Link-Intervall. Dann gilt $r \leq \sigma(i) < \sigma(j) \leq s$ und es existiert ein $k \in [\sigma(i) + 1 : \sigma(j)]$ mit $L[k] = \ell - 1$. Dieses k lässt sich mithilfe einer Range Minimum Query Anfrage $\text{RMQ}_L(\sigma(i) + 1, \sigma(j))$ auf dem Feld L in konstanter Zeit ermitteln (nach einer linearen Vorverarbeitungszeit). Wenn wir an jedem ℓ -Index die Intervallgrenzen des zugehörigen ℓ -Intervalls notiert haben, können wir diese in konstanter Zeit nachschlagen.

Wie speichern wir uns nun für jedes ℓ -Intervall seine Intervallgrenzen an seinen ℓ -Indizes? Dazu traversieren wir den LCP-Intervall-Baum mit einer Breitensuche (die Implementierung einer Breitensuche sei dem Leser zur Übung überlassen). Jedes Mal wenn wir ein neues ℓ -Intervall bearbeiten, tragen wir an seinen ℓ -Indizes die

Intervall grenzen ein. Die Laufzeit ist also proportional zur Anzahl aller ℓ -Indizes (für alle ℓ -Intervalle und für alle ℓ) und ist somit linear in $|t|$.

Theorem 5.45 *Sei E ein Enhanced-Suffix-Array, dann kann dieses in linearer Zeit mit linearem Platz so erweitert werden, dass für jedes LCP-Intervall von E dessen Suffix-Link-Intervall in konstanter Zeit aufgefunden werden kann.*

5.6 Extended Suffix Arrays

In diesem Abschnitt wollen wir jetzt noch eine einfachere Variante vorstellen, wie man die Algorithmen für Suffix-Bäume auf Suffix-Arrays übertragen kann. Dabei wird die LCP-Tabelle wiederum eine bedeutende Rolle spielen. In diesem Abschnitt nehmen wir jetzt wieder an, dass $\$ \notin \Sigma$ jetzt das größte Zeichen ist, d.h. es gilt $\$ < a$ für alle $a \in \Sigma$.

5.6.1 LCP-Intervalle

Zunächst definieren die so genannten LCP-Intervalle, die grundlegend für das Folgende sind.

Definition 5.46 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Ein Intervall $[i : j]$ mit $i < j \in [0 : n]$ heißt ein LCP-Intervall vom Typ ℓ oder kurz ein ℓ -Intervall, wenn*

1. $L[i] < \ell$,
2. $L[k] \geq \ell$ für alle $k \in [i + 1 : j]$,
3. $L[k] = \ell$, für mindestens ein $k \in [i + 1 : j]$ und
4. $L[j + 1] < \ell$.

Hierbei gelte $L[0] = L[n + 1] = -1$.

Wir werden später auch Intervalle der Länge 1 zulassen.

Notation 5.47 ℓ - $[i : j]$ bezeichnet das ℓ -Intervall $[i : j]$.

i	$A[i]$	$L[i]$	$t^{A[i]}$
0	12	-1	\$
1	11	0	I\$
2	8	1	IPPI\$
3	5	1	ISSIPPI\$
4	2	4	ISSISSIPPI\$
5	1	0	MISSISSIPPI\$
6	10	0	PI\$
7	9	1	PPI\$
8	7	0	SIPPI\$
9	4	2	SISSIPPI\$
10	6	1	SSIPPI\$
11	3	3	SSISSIPPI\$

Abbildung 5.39: Beispiel: Suffix-Array, LCP-Tabelle, Child-Tabelle und alle ℓ -Intervalle für MISSISSIPPI (außer dem 0-Intervall)

In Abbildung 5.39 sind die ℓ -Intervalle für unser Beispielwort MISSISSIPPI noch einmal illustriert. Das 0-Intervall $[0 : n]$ ist hierbei nicht dargestellt.

Lemma 5.48 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Für jedes $k \in [1 : n]$ mit $L[k] = \ell$ existiert ein ℓ -Intervall $[i : j]$ mit $k \in [i + 1 : j]$

Beweis: Wir definieren zunächst einmal $i := \max \{p < k : L[p] < L[k] = \ell\}$ und $j := \min \{p \geq k : L[p + 1] < L[k] = \ell\}$. Diese existieren nach der erweiterten Definition von L immer ($L[0] = L[n + 1] = -1$). Wir behaupten jetzt, dass $[i : j]$ das gesuchte ℓ -Intervall ist. Nach Definition von i gilt $L[i] < \ell$, also gilt Bedingung 1 der Definition eines LCP-Intervalls. Analog gilt nach Definition von j gilt $L[j + 1] < \ell$, also gilt Bedingung 4 der Definition eines LCP-Intervalls. Weiterhin gilt $k \in [i + 1 : j] \neq \emptyset$ und $L[k] = \ell$, also Bedingung 3 der Definition. Da wir i maximal und j minimal unter der Bedingung $L[p] < L[k] = \ell$ wählen, gilt $L[r] \geq L[k] = \ell$ für alle $r \in [i + 1 : j]$ und somit die Bedingung 2. ■

Jetzt definieren wir noch so genannte ℓ -Indizes

Definition 5.49 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Wenn $[i : j]$ ein ℓ -Intervall ist, dann heißt $k \in [i + 1 : j]$ ein ℓ -Index, wenn $L[k] = \ell$. Die Menge aller ℓ -Indizes eines ℓ -Intervalls $[i : j]$ wird mit $I_\ell[i : j]$ bezeichnet.

Definition 5.50 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. L das zugehörige Suffix-Array bzw. die zugehörige LCP-Tabelle. Ein m -Intervall $[i' : j']$ ist in einem ℓ -Intervall $[i : j]$ enthalten, wenn $i \leq i' < j' \leq j$ und $m > \ell$ gilt.

Wenn $[i' : j']$ ein in $[i : j]$ enthaltenes Intervall ist und es kein anderes ℓ'' -Intervall gibt, das in $[i : j]$ enthalten ist und das $[i' : j']$ enthält, dann heißt $[i' : j']$ ein Kind-Intervall von $[i : j]$. Umgekehrt heißt $[i : j]$ das Elter-Intervall von $[i' : j']$.

Mit Hilfe dieser Enthaltensein-Relation können wir zu den LCP-Intervallen auch eine Baumstruktur generieren. Betrachten wir zunächst noch einmal den Suffix-Baum für MISSISSIPPI in Abbildung 5.40. Konstruieren wir für das zugehörige Suffix-Array und den zugehörigen LCP-Werten den so genannten *LCP-Intervall-Baum*, der auf der Enthaltensein-Relation der LCP-Intervall definiert ist. Beachte dabei, dass $[0 : n]$ immer ein 0-Intervall ist. Weiter ergänzen wir den so konstruierten LCP-Intervall-Baum um die so genannten *Singletons*, das sind die einelementigen Mengen von $[0 : n]$. In der Regel werden wir auch Singletons als Kind-Intervalle zulassen. In Abbildung 5.41 ist der LCP-Intervall-Baum für unser Beispielwort MISSISSIPPI angegeben.

Man sieht sofort die Ähnlichkeit zum entsprechenden Suffix-Baum. In den LCP-Intervallen sind also bereits alle wesentlichen Informationen des Suffix-Baumes (ohne die Blätter) enthalten. In Abbildung 5.41 entsprechen die schwarzen Kanten der Enthaltensein-Relation der Suffix-Bäume und die türkisen Kanten den Kanten zu den Blätter, die nach Definition ja keine ℓ -Intervalle sind. Die roten Kanten entsprechen den Kanten auf das älteste Kind bzw. auf das nächstältere Geschwister gemäß der Definition in der Darstellung eines Suffix-Baumes. Im Folgenden müssen wir nur versuchen diese roten Kanten aus den LCP-Intervallen zu rekonstruieren.

Lemma 5.51 Sei $[i : j]$ ein ℓ -Intervall und seien $i_1 < i_2 < \cdots < i_k$ die ℓ -Indizes des ℓ -Intervalls $[i : j]$, dann sind die Kind-Intervalle von $[i : j]$ gerade $[i : i_1 - 1]$, $[i_1 : i_2 - 1]$, $\dots$, $[i_k : j]$.

Wir merken hier noch an, dass $[i_p : i_{p+1} - 1]$ für $p \in [1 : k - 1]$ durchaus ein Singleton sein kann.

Beweis: Sei $[r : s]$ eines der Intervalle $[i : i_1 - 1]$, $[i_1 : i_2 - 1]$, $\dots$, $[i_k : j]$.

Ist $[r : s]$ ein Singleton (d.h. $r = s$), dann ist $[r : s]$ nach unserer Konvention ein Kind-Intervall.

Gelte also jetzt $r < s$. Wir definieren zunächst $m = \min \{L[x] : x \in [r + 1 : s]\}$ sowie $x' = \operatorname{argmin} \{L[x] : x \in [r + 1 : s]\}$. Nach Definition des ℓ -Intervalls gilt $m > \ell$.

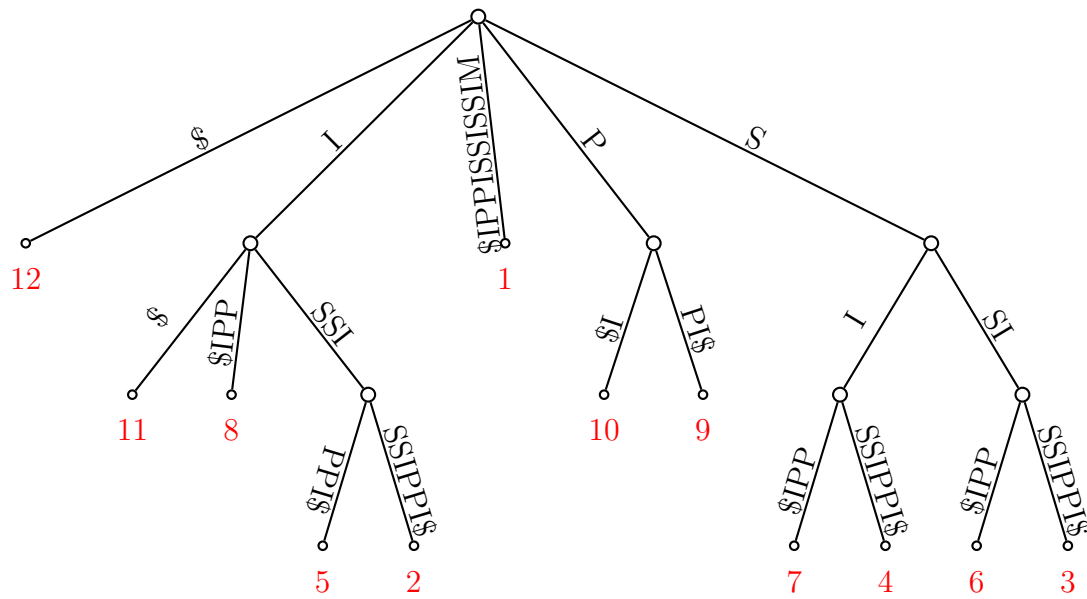


Abbildung 5.40: Beispiel: Suffix-Baum für MISSISSIPPI

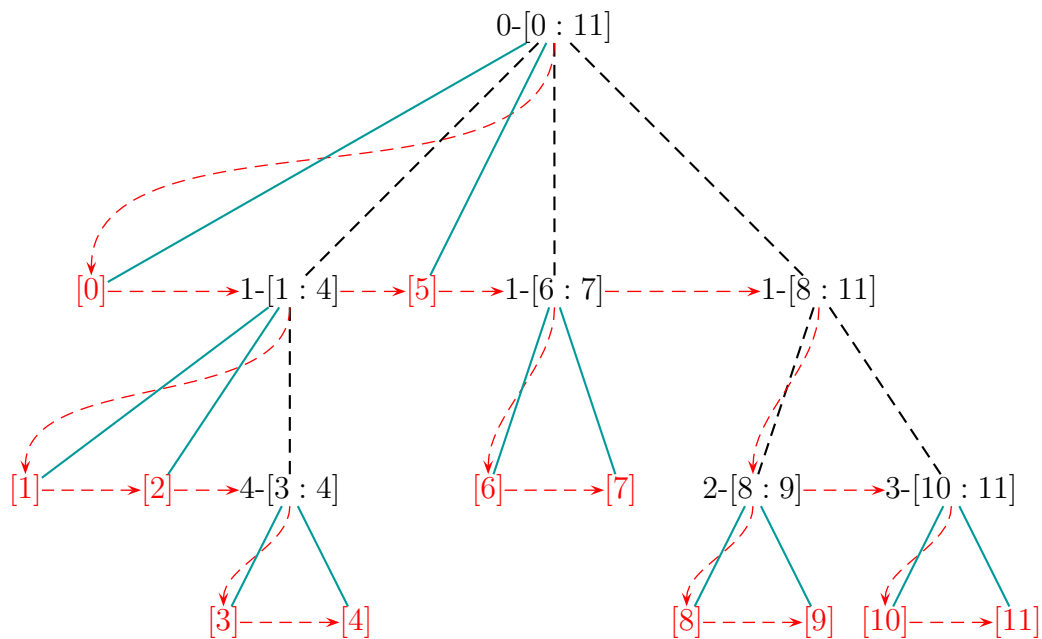


Abbildung 5.41: Beispiel: LCP-Intervall-Baum für MISSISSIPPI

Nach Lemma 5.48 existiert ein m -Intervall I , dass die Position x' umfasst. Nach Definition der ℓ -Indizes des ℓ -Intervalls $[i : j]$ und nach Wahl von m muss diese Intervall I gerade gleich $[r : s]$ sein.

Weiterhin ist $[r : s]$ in $[i : j]$ enthalten. Da weiterhin nach Definition der ℓ -Indizes $L[i_1] = L[i_2] = \dots = L[i_k] = \ell$ gilt, kann kein anderes Intervall in $[i : j]$ enthalten sein, welches $[r : s]$ enthält.

Offensichtlich sind die Intervalle $[i : i_1 - 1]$, $[i_1 : i_2 - 1]$, $\dots$, $[i_k : j]$ alle möglichen Kind-Intervalle, die $[i : j]$ besitzen kann. ■

5.6.2 Navigation im Extended Suffix-Array

Wie findet man nun das älteste Kind eines ℓ -Intervalls? Sei ℓ - $[i : j]$ das betrachtete ℓ -Intervall. Das erste Kind-Intervall muss die Form ℓ' - $[i : k]$ mit $\ell' > \ell$ haben. Dabei gilt offensichtlich, dass der erste ℓ -Index der minimale Wert in der LCP-Tabelle im Bereich $[i + 1 : j]$ sein muss, also gilt für das Ende des ersten Kind-Intervalls:

$$k = \text{RMQ}_L(i + 1, j) - 1.$$

Hierbei nutzen wir aus, dass die RMQ-Anfrage bei mehrdeutigen Ergebnissen immer den linkensten Index zurückliefert.

Wie findet man nun den Wert ℓ' , sofern $i < k$? Wir suchen einfach im ersten gefundenen Kind-Intervall nach dem ersten ℓ' -Index respektive nach dem zugehörigen LCP-Wert:

$$\ell' = L[\text{RMQ}_L(i + 1, k)].$$

Und wie findet man nun das nächstjüngere Geschwister eines ℓ -Intervalls? Dies lässt sich ebenso leicht mit einer RMQ-Anfrage erledigen, vorausgesetzt, das Eltern-Intervall ist bekannt. Sei also ℓ' - $[i' : j' - 1]$ ein Kind von ℓ - $[i : j]$. Dann gilt offensichtlich für das benachbarte Geschwister ℓ' - $[j' : k]$:

$$k = \text{RMQ}_L(j' + 1, j) - 1,$$

sofern $L[k] = \ell'$. Andernfalls haben wir das letzte Kind des Eltern-Intervall ℓ - $[i : j]$ gefunden und setzen $k = j$.

Damit können wir alle Algorithmen auf Suffix-Bäumen, die mit einer Tiefen- oder Breitensuche implementiert sind, auf Extended Suffix-Arrays übertragen. Wir werden später noch sehen, welche zusätzlichen Datenstrukturen man braucht, um den Elter eines ℓ -Intervalls auffinden zu können.

Wir können also die Navigation durch die ℓ -Intervalle mit Hilfe einer Datenstruktur für RMQ auf der LCP-Tabelle L effizient erledigen. Wie im Kapitel 4 angemerkt, kann eine solche Datenstruktur für n Werte mit $2n + o(n)$ zusätzlichen Bits auskommen.

09.01.07

5.6.3 Optimale Suche in Extended Suffix-Arrays

Basierend auf dem Algorithmus zum Auffinden aller Kind-Intervalle eines ℓ -Intervalls geben wir in diesem Abschnitt einen optimalen Algorithmus zum Suchen nach Teilwörtern basierend auf Suffix-Arrays an. Der Algorithmus ist in Abbildung 5.42 angegeben.

Die Invariante im Algorithmus ist die folgende: Während der Suche nach einem Wort $s = s_1 \cdots s_m$ gibt die Boolesche Variable **prefix** an, ob ein Präfix der Länge p von s ein Teilwort t enthalten ist oder nicht. Weiterhin ist $[i : j]$ ein p -Intervall (außer im Falle $p = m$), dass nur Suffixe von t enthält, die mit dem Präfix $s_1 \cdots s_p$ beginnen. Halten wir das Ergebnis im folgenden Satz fest.

Theorem 5.52 *In einem Extended Suffix-Array kann in Zeit $O(|\Sigma| \cdot m)$ festgestellt werden, ob $s \in \Sigma^m$ in $t \in \Sigma^*$ enthalten ist.*

Beweis: Die Korrektheit folgt aus den Zeichenvergleichen. Die Laufzeit folgt aus der Tatsache, dass alle anderen Funktionen in konstanter Zeit durchgeführt werden können und die Laufzeit somit proportional zur Anzahl der Zeichenvergleiche ist, also $O(m)$ beträgt. ■

5.6.4 Auffinden des Elters

Überlegen wir uns zunächst, wie man das Elter-Intervall eines ℓ -Intervalls ℓ - $[i : j]$ ermitteln kann. Zunächst überlegt man sich leicht, dass das Elter-Intervall ein k -Intervall mit $k = \max\{L[i], L[j+1]\}$ ist. Weiterhin definieren wir zwei Felder P (wie previous) und N (wie next) wie folgt:

$$\begin{aligned} P[i] &= \max \{j \in [0 : i] : L[j] < L[i]\}, \\ N[i] &= \min \{j \in [i : n+1] : L[j] < L[i]\}. \end{aligned}$$

Wir halten zunächst das folgende Lemma fest, wobei der Beweis ähnlich zu dem Beweis zur Konstruktion des Vektors L^B bei der Vorverarbeitung zu RMQ ist (siehe Seite 131).

FIND (string s)

```
{
    interval  $[i : j] = [0 : n]$ ;
    int  $p = 0$ ; bool prefix = TRUE; /* prefix is TRUE iff  $s_1 \cdots s_p$  is a prefix of  $t^i$  */
    while ( $([i : j] \neq \perp) \ \&\& \ (p < m) \ \&\& \ (\text{prefix})$ )
    {
        if ( $i < j$ )          /* a Child-Interval */
        {
             $k = \min(\text{LCP}(i, j), m)$ ;
            prefix =  $(t_{A[i]+p} \cdots t_{A[i]+k-1} = s_{p+1} \cdots s_k)$ ;
             $p = k$ ;
        }
        elseif ( $i = j$ )      /* a Singleton */
        {
            prefix =  $(t_{A[i]+p} \cdots t_{A[i]+m-1} = s_{p+1} \cdots s_m)$ ;
             $p = m$ ;
        }
         $[i : j] = \text{getChildIntervalByChar}(i, j, s_{p+1})$ ;
    }
    output (prefix)? $[i : j] : \perp$ ;
}
```

GETCHILDINTERVALBYCHAR (int i, j ; char c)

```
{
    interval  $[i' : j'] = \text{getNextChildInterval}(i, j, i)$ ;
    while ( $(t_{A[i']} \neq c) \ \&\& \ (j' < j)$ )
         $[i' : j'] = \text{getNextChildInterval}(i, j, j' + 1)$ ;
    return  $(t_{A[i']} = c) ? ([i' : j']) : (\perp)$ ;
}
```

GETNEXTCHILDINTERVAL (int i, j, k)

```
{
     $\ell = \text{RMQ}_L(i + 1, j)$ ;
     $j' = \text{RMQ}_L(k + 1, j) - 1$ ;
    return  $(L[j'] > L[\ell]) ? ([k : j]) : ([k : j'])$ ;
}
```

LCP(int i, j)

```
{
    return  $L[\text{RMQ}_L(i, j)]$ ;
}
```

Abbildung 5.42: Algorithmus: Suche s im Extended Suffix-Array von t

Lemma 5.53 *Sei $t \in \Sigma^n$ und S bzw. L das zugehörige Suffix- bzw. LCP-Array. Die Felder P und N können in linearer Zeit erstellt werden.*

Beweis: Übungsaufgabe. ■

Mit Hilfe der Felder P und N kann dann das Elter-Intervall $\ell'-[i' : j']$ von $\ell-[i : j]$ leicht wie folgt angegeben werden:

$$\begin{aligned}\ell' &= \max\{L[i], L[j+1]\}, \\ k &= \operatorname{argmax}\{L[i], L[j+1]\}, \\ i' &= P[k], \\ j' &= N[k] - 1.\end{aligned}$$

Der Beweis hierfür orientiert sich im Wesentlichen an Lemma 5.48 bzw. an dessen Beweis.

5.6.5 Suffix-Links

Es bleibt noch zu überlegen, wie man Suffix-Links in die Enhanced-Suffix-Arrays integrieren kann.

Definition 5.54 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Für $t^{A[i]} = aw\$$ mit $a \in \Sigma$ und $w \in \Sigma^*$ ist $\sigma(i) = j$ mit $t^{A[j]} = w\$$ der Suffix-Link an Position i .*

Wie man leicht sieht, gilt die folgende Beobachtung.

Beobachtung 5.55 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A bzw. R das zugehörige Suffix-Array bzw. das zugehörige inverse Suffix-Array. Dann gilt $\sigma(i) = R[A[i] + 1]$, sofern $A[i] \leq n$*

Wir haben somit nur Suffix-Links für Singletons definiert. Diese zeigen selbst jedoch wieder auf Singletons, namentlich das Intervall $[R[A[i] + 1] : R[A[i] + 1]]$, sofern $A[i] \neq n + 1$ ist. Ist $A[i] = n + 1$, so beschreibt das Singleton das Suffix \$, dessen Suffix-Link das Wurzel-Intervall $0-[0 : n]$ ist.

Für LCP-Intervalle können wir das folgende Lemma beweisen.

Lemma 5.56 *Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Zu jedem ℓ -Intervall $[i : j]$ mit $\ell > 0$ existiert genau ein kleinstes $(\ell - 1)$ -Intervall $[r : s]$ mit $r \leq \sigma(i) < \sigma(j) \leq s$.*

Beweis: Da $[i : j]$ ein ℓ -Intervall ist, gilt offensichtlich für $k \in [\sigma(i) + 1 : \sigma(j)]$, dass $L[k] \geq \ell - 1$. Da $[i : j]$ ein ℓ -Intervall ist, muss es weiter einen ℓ -Index $k \in [i + 1 : j]$ mit $L[k] = \ell$ geben. Deshalb muss es auch einen Index $k' \in [\sigma(i) + 1 : \sigma(j)]$ mit $L[k'] = \ell - 1$ geben. Nach Lemma 5.26 existiert dann ein $(\ell - 1)$ -Intervall $[r : s]$, das nach Konstruktion $[\sigma(i) : \sigma(j)]$ umfasst. ■

Mit diesem Existenz-Satz können wir nun die folgende Definition für Suffix-Links von LCP-Intervallen wagen.

Definition 5.57 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Zu jedem ℓ -Intervall $[i : j]$ mit $\ell > 0$ ist das kleinste $(\ell - 1)$ -Intervall $[r : s]$ mit $r \leq \sigma(i) < \sigma(j) \leq s$ das Suffix-Link-Intervall von $[i : j]$.

Der Beweis des folgenden Korollars ergibt sich sofort aus einer genauen Inspektion des Beweises von Lemma 5.42.

Korollar 5.58 Sei $t = t_1 \cdots t_n \in \Sigma^*$ und sei A das zugehörige Suffix-Array. Sei $[i : j]$ ein ℓ -Intervall mit $\ell > 0$ und sei $[r : s]$ sein Suffix-Link-Intervall. Dann existiert ein $(\ell - 1)$ -Index $k \in I_{\ell-1}[r : s]$ und $k \in [\sigma(i) + 1 : \sigma(j)]$.

Sei nun $[i : j]$ ein ℓ -Intervall und $[r : s]$ das zugehörige Suffix-Link-Intervall. Dann gilt $r \leq \sigma(i) < \sigma(j) \leq s$ und es existiert ein $k \in [\sigma(i) + 1 : \sigma(j)]$ mit $L[k] = \ell - 1$. Dieses k lässt sich mithilfe einer Range Minimum Query Anfrage $k = \text{RMQ}_L(\sigma(i) + 1, \sigma(j))$ auf dem Feld L in konstanter Zeit ermitteln (nach einer linearen Vorverarbeitungszeit). Dann kann auch das Suffix-Link-Intervall $[r : s]$ wie folgt bestimmt werden:

$$\begin{aligned} r &= P[k] = P[\text{RMQ}_L(\sigma(i) + 1, \sigma(j))], \\ s &= N[k] - 1 = N[\text{RMQ}_L(\sigma(i) + 1, \sigma(j))] - 1. \end{aligned}$$

Theorem 5.59 Sei E ein Extended Suffix Array, dann kann dieses in linearer Zeit mit linearem Platz so erweitert werden, dass für jedes LCP-Intervall von E dessen Suffix-Link-Intervall in konstanter Zeit aufgefunden werden kann.

5.6.6 LCA-Queries in Extended Suffix-Arrays

Wir wollen nun für zwei LCP-Intervalle $[i : j]$ und $[i' : j']$ mit $i < i'$ das zum kleinsten gemeinsamen Vorfahren gehörige LCP-Intervall bestimmen. Gilt $j' \leq j$, dann ist offensichtlich $[i : j]$ ein Vorfahre von $[i' : j']$ und somit ist $[i : j]$ das

gesuchte Elter-Intervall. Gelte nun also $j < j'$, dann muss auch $j < i$ gelten, da sich LCP-Intervalle nicht überlappen können.

Für das niedrigsten Elter-Intervall ℓ - $[x : y]$ von $[i : j]$ und $[i' : j']$ gilt, dass ein ℓ -Index in der LCP-Tabelle im Intervall $[j + 1 : i']$ stehen muss. Dieser ℓ -Index lässt sich dann wie folgt bestimmen:

$$k := \text{RMQ}_L(j + 1, i').$$

Das zugehörige LCP-Intervall ist dann durch $[r : s]$ mit

$$\begin{aligned} r &= P[k], \\ s &= N[k] - 1. \end{aligned}$$

gegeben. Der Beweis hierfür orientiert sich wiederum an Lemma 5.48 bzw. an dessen Beweis.

5.6.7 Speicherplatzbedarf

In diesem Abschnitt fassen wir noch einmal kurz den Speicherplatzbedarf für die Varianten des Extended Suffix-Arrays zusammen. Siehe hierzu auch die Tabelle 5.43. Dort ist unter Name die verwendete Datenstruktur genannt, unter Platz der theoretische Platzbedarf in Bits und unter Bytes der praktische Platzbedarf in Bytes bei Verwendung einer 32-Bit Architektur (Maschinenwort gleich 4 Byte) aufgelistet. Unter Verwendung wird erwähnt, für welche Funktionalität die entsprechende Datenstruktur benötigt wird.

Name	Feld	Platz	Bytes	Verwendung
Text	t	$O(n \log(\Sigma))$	n	Immer
Suffix-Array	A	$O(n \log(n))$	$4n$	Immer
LCP-Tabelle	L	$O(n \log(n))$	$n + x$	Immer
RMQ-Info.	Q	$O(n)$	$n/4 + x$	Immer
Prev/Next-Array	P, N	$O(n \log(n))$	$8n$	Parent, Suffix-Link, LCA
Inv. Suffix-Array	R	$O(n \log(n))$	$4n$	Suffix-Link

Abbildung 5.43: Tabelle: Platzbedarf Extended Suffix-Array

Für den gegebenen Text wird die Zeichenreihe t abgespeichert, dabei wird in der Praxis angenommen, dass jedes Zeichen durch ein Byte dargestellt wird. Das Suffix-Array (bzw. das inverse Suffix-Array) speichert ja eine Permutation von $[1 : n + 1]$ (bzw. $[0 : n]$) und benötigt daher pro Eintrag n Bits, in der Praxis 4 Bytes ($n < 2^{32}$). Die LCP-Tabelle beträgt normalerweise (außer bei starken Wiederholungen) nur

kleine Werte, so dass Werte aus $[1 : 254]$ in einem Byte in einem Feld gespeichert werden. Sind die Werte größer, wird in der LCP-Tabelle der Wert 255 gespeichert und der korrekte Wert einer Hash-Tabelle außerhalb des Feldes. Sind die LCP größtenteils klein, so bleibt auch die Hash-Tabelle vernachlässigbar klein. Wir schon erwähnt kann die zusätzlich Information für einen konstante RMQ-Anfrage auf die LCP-Tabelle in Platz $2n + o(n)$ Bits gespeichert werden. Für zusätzliche Funktionalitäten wie explizite Elter-Information (die bei normalen Suffix-Baum-Traversierungen gar nicht benötigt wird) werden noch die Previous- und Next-Tabelle sowie das inverse Suffix-Array benötigt. Speichert man in den Previous- und Next-Tabelle auch nicht die absoluten Indexwerte, sondern relative Offsets, so sind auch diese of klein und es kann derselbe Trick wie beim Speichern der LCP-Tabelle angewendet werden. Es gibt bereits Verfahren wie diese Information wesentlich platzsparender gespeichert werden kann. Dazu verweisen wir allerdings auf die Originalliteratur.

6.1 Modellbildung

In diesem Abschnitt wollen wir uns mit so genannten Genome Rearrangements beschäftigen. Es hat sich herausgestellt, dass die DNA-Sequenzen von Genen nah verwandter Spezies sehr ähnlich sind, aber sich die Anordnung der Gene auf dem Genom doch sehr erheblich unterscheiden kann. Aus diesem Grunde betrachten wir nun die DNA nicht als Folge von Basenpaaren, sondern als lineare Folge von Genen bzw. Bereiche von Genen und wollen die Änderungen auf diesem Niveau studieren.

6.1.1 Rearrangements und zugehörige Mutationen

Wir verstehen also im Folgenden unter einem Genom die lineare Abfolge der kodierten Gene auf diesem. In Abbildung 6.1 ist ein künstliches Beispiels von zwei Genomen gezeigt, in der die Gene oder auch zusammengefasste Bereiche von Genen mit den Buchstaben von A bis J markiert sind.

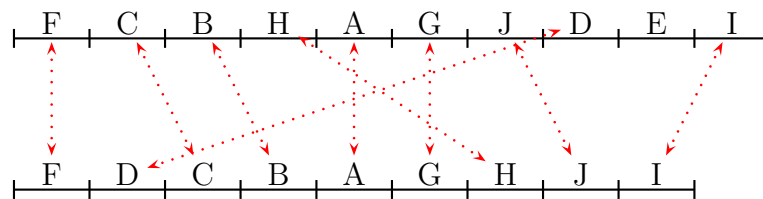


Abbildung 6.1: Beispiel: Genome Rearrangement

Wir wollen uns zunächst einmal überlegen, welche biologischen Mutationen solche Umordnungen der Gene hervorrufen können.

Deletion, Insertion: Wie auf Basen-Ebene können hier jetzt jedoch ganze Gen-Gruppen eliminiert oder eingefügt werden. Siehe hierzu auch Abbildung 6.2.



Abbildung 6.2: Skizze: Deletion bzw. Insertion

Duplikation: Hierbei werden ganze Gruppen von Genen dupliziert. Siehe hierzu auch Abbildung 6.3.



Abbildung 6.3: Skizze: Duplikation

Inversion, Reversion: Die Ordnung der Gene wird in einem Teilbereich umgedreht. Für den biologischen Begriff Inversion verwenden wir den mathematisch korrekten Begriff der Reversion. Siehe hierzu auch Abbildung 6.4.



Abbildung 6.4: Skizze: Reversion

Transposition: Bei einer Transposition werden zwei benachbarte Bereiche von Gengruppen miteinander vertauscht. Siehe hierzu auch Abbildung 6.5.



Abbildung 6.5: Skizze: Transposition

Translokation: Bei einer Translokation werden von jeweils zwei Chromosomen die jeweiligen Enden vertauscht. Siehe hierzu auch Abbildung 6.6.



Abbildung 6.6: Skizze: Translokation

Fission: Hierbei wird ein Chromosom in zwei neue Chromosomen aufgesplittet. Siehe hierzu auch Abbildung 6.7.



Abbildung 6.7: Skizze: Fission

Fusion: Bei der Fusion werden zwei Chromosome zu einem neuen Chromosom verschmolzen. Siehe hierzu auch Abbildung 6.8.



Abbildung 6.8: Skizze: Fusion

Bei Genomen von mehreren Chromosomen kann eine Duplikation nicht nur einen Bereich sondern auch ein ganzes Chromosom duplizieren. In Abbildung 6.9 ist ein Beispiel für ein Genom Rearrangement angegeben. Man sollte sich überlegen, dass es zumindest eine kürzere Folge von Operationen für dieses Beispiel gibt.

F	C	B	H	A	G	J	D	E	I	Deletion
F	C	B	H	A	G	J	D	I		Reversal
F	D	J	G	A	H	B	J	I		Reversal
F	D	C	B	H	A	G	J	I		Transposition
F	D	C	B	A	G	H	C	I		

Abbildung 6.9: Beispiel: Ein Genom Rearrangement

6.1.2 Permutationen

Für die folgenden Untersuchungen betrachten wir folgende Einschränkungen:

Keine Deletionen/Insertionen: Diese sind sehr leicht festzustellen, da man nur überprüfen muss, welche Gruppen nur auf einem Genom vorhanden sind und welche nicht.

Keine Duplikate/Duplikationen: Diese machen momentan die mathematische Analyse noch sehr schwierig, deshalb konzentrieren wir uns zunächst auf den einfacheren Fall ohne Duplikate.

Mono-chromosomale Genome: Dies ist eine vorläufige Einschränkung, um diesen Fall zuerst behandeln zu können.

Aufgrund des Ausschluss von Duplikaten, können wir die Gen-Gruppen in einem der beiden Genome von 1 bis n in der auftretenden Reihenfolge durchnummerieren. Da auch keine Insertionen und Deletionen erlaubt sind, muss das andere Genom dann eine Permutation der Zahlen von 1 bis n darstellen. Siehe hierzu auch das Beispiel in Abbildung 6.10.

1	2	3	4	5	6	7	8	9
F	C	B	H	A	G	J	D	I
F	D	C	B	A	G	H	C	I
1	8	2	3	5	6	4	7	9

Abbildung 6.10: Beispiel: Nummerierung der Gen-Gruppen

Damit kann ein Genom als Permutation von $[1 : n]$ dargestellt werden, das andere Genom besitzt dann ohne Beschränkung der Allgemeinheit die Identität.

Definition 6.1 Sei $n \in \mathbb{N}$, dann heißt

$$S_n = \{(\pi_1, \dots, \pi_n) : \{p_1, \dots, p_n\} = [1 : n]\}$$

die symmetrische Gruppe und seine Elemente Permutationen. Die zugehörige Operation $\circ$ ist die Hintereinanderausführung von Permutationen, die durch $(\pi \circ \sigma)_i = \sigma_{\pi_i}$ definiert ist.

Wie allseits bekannt ist, ist $(S_n, \circ)$ eine Gruppe. Mit $\text{id} = (1, \dots, n)$ bezeichnen wir die identische Permutation.

Damit stellt sich nun die folgende Frage: Wie viele (und welche) Elementaroperationen (wie Reversionen oder Transpositionen) sind nötig, um eine gegebene Permutation $\pi \in S_n$ in die Identität zu überführen (d.h. zu *sortieren*).

REARRANGEMENT DISTANCE

Eingabe: Eine Permutation $\pi \in S_n$.

Gesucht: Eine kürzeste Folge von Elementaroperationen $\rho_1, \dots, \rho_k$ vorgegeben Typs mit $\pi \circ \rho_1 \circ \dots \circ \rho_k = \text{id}$.

Der minimale Abstand ist ein Maß für den evolutionären Abstand. In Wirklichkeit kann der Abstand auch größer sein, aber der tatsächliche Abstand ist logischerweise nur aus den beiden Anordnungen in den Genomen nicht mehr rekonstruierbar. Es ist auch mehr als plausibel, dass der von der Natur gewählte Weg relativ nah am kürzest möglichen Weg liegt, da sinnlose Umwege eher unwahrscheinlich sind.

Reversionen oder Transpositionen sind Generatoren der symmetrischen Gruppe, d.h. für jedes $\pi \in S_n$ existieren $\rho_1, \dots, \rho_k$ mit $\pi = \rho_1 \circ \dots \circ \rho_k$. Auch Reversionen und Transpositionen können wir in Form von Permutationen angeben:

$$\begin{aligned} \text{Reversal}(i, j) &= \begin{pmatrix} 1 \dots i-1 & i & i+1 \dots j-1 & j & j+1 \dots n \\ 1 \dots i-1 & j & j-1 \dots i+1 & i & j+1 \dots n \end{pmatrix} \\ \text{Transposition}(i, j, k) &= \begin{pmatrix} 1 \dots i-1 & i \dots j & j+1 \dots k & k+1 \dots n \\ 1 \dots i-1 & j+1 \dots k & i \dots j & k+1 \dots n \end{pmatrix} \end{aligned}$$

Man beachte hierbei, dass in dieser Darstellung die Überlappung des grünen und blauen Bereiches bei einer Transposition nicht notwendigerweise korrekt dargestellt ist und eine korrekte Darstellung auch immer von den konkreten Werten i , j und k abhängt.

Weiterhin muss man auch beachten, dass man Permutationen auf zwei Arten interpretieren kann: zum einen als aktuelle Umordnung der Elemente und zum anderen als

<u>Permutation als Umordnung</u>	<u>Permutation als Abbildung</u>
$\pi = (2, 3, 1, 5, 4)$	$\pi' = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 1 & 2 & 5 & 4 \end{pmatrix}$
Reversion $\rho_1 = (2, 4)$	$\rho'_1 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 1 & 4 & 3 & 2 & 5 \end{pmatrix}$
$\pi \cdot \rho_1 = (2, 5, 1, 3, 4)$	$\pi' \circ \rho'_1 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 1 & 4 & 5 & 2 \end{pmatrix}$
Reversion $\rho_2 = (1, 4)$	$\rho'_2 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 3 & 2 & 1 & 5 \end{pmatrix}$
$\pi \cdot \rho_1 \cdot \rho_2 = (3, 1, 5, 2, 4)$	$\pi' \circ \rho'_1 \circ \rho'_2 = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 2 & 4 & 1 & 5 & 3 \end{pmatrix}$
Transposition $\tau = (1, 2, 4)$	$\tau' = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 4 & 1 & 2 & 5 \end{pmatrix}$
$\pi \cdot \rho_1 \cdot \rho_2 \cdot \tau = (5, 2, 3, 1, 4)$	$\pi' \circ \rho'_1 \circ \rho'_2 \circ \tau' = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 2 & 3 & 5 & 1 \end{pmatrix}$

Abbildung 6.11: Beispiel: Rechnen mit Permutationen

Abbildungsvorschrift. Der letztere Fall hat den Vorteil, dass man dann auch die Elementaroperationen als Permutationen auffassen kann und die Verknüpfung der symmetrischen Gruppe genau der Ausführung dieser Elementaroperationen entspricht. Um die beiden Darstellungsarten unterscheiden zu können, wird die Umordnung $(2, 3, 1, 5, 4)$, interpretiert als Abbildung, durch $\begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 3 & 1 & 2 & 5 & 4 \end{pmatrix}$ notiert. In der Abbildung bedeutet dies beispielsweise, dass das erste bzw. vierte Element an die dritte bzw. fünfte Stelle kommt. Genau dort steht auch die 1 bzw. die 4 in der Umordnung. In der Abbildung 6.11 ist ein etwas umfangreicheres Beispiel mit Anwendungen von Reversionen und Transpositionen angegeben.

Man beachte, dass wir im Folgenden der Einfachheit halber oft die Schreibweise für Umordnung auch für Abbildungen verwenden werden.

6.2 Sorting by Reversals

In diesem Abschnitt wollen wir uns nun mit dem Genome Rearrangement beschäftigen, wenn nur Reversionen als Elementaroperationen verwendet werden dürfen.

SORTING BY REVERSALS (MIN-SBR)

Eingabe: Eine Permutation $\pi \in S_n$.

Gesucht: Eine kürzeste Folge von Reversionen $\rho_1, \dots, \rho_k$ mit $\pi \circ \rho_1 \circ \dots \circ \rho_k = \text{id}$.

6.2.1 Komplexität von Min-SBR

In diesem Abschnitt wollen wir kurz ohne Beweise auf die Komplexität von Min-SBR eingehen. Man kann zeigen, dass Min-SBR ein $\mathcal{NP}$ -hartes Optimierungsproblem ist. Viel schlimmer noch, Min-SBR ist sogar $\mathcal{APX}$ -hart.

Definition 6.2 Ein Optimierungsproblem P gehört zur Klasse $\mathcal{APX}$, wenn es einen polynomiell zeitbeschränkten Algorithmus A sowie ein $r \geq 1$ gibt, so dass A eine r -Approximation für P ist.

Hierbei bedeutet r -Approximation, dass der Algorithmus eine Lösung liefert, die höchstens um den Faktor r von der optimalen Lösung abweicht.

Ein $\mathcal{APX}$ -hartes Optimierungsproblem gehört somit zu den schwierigsten Problemen der Klasse, die sich bis auf einen konstanten Faktor approximieren lassen. Eine weitere Folgerung ist, dass $\mathcal{APX}$ -harte Probleme kein polynomielles Approximationsschema besitzen können (außer wenn $\mathcal{P} = \mathcal{NP}$).

Ein Optimierungsproblem gehört zur Klasse $\mathcal{PTAS}$ (Polynom Time Approximation Scheme), wenn es ein Polynom p und einen Algorithmus A gibt, der für jedes $\varepsilon > 0$ das Problem P in Zeit $p(x)$ für die Eingabe x mit Approximationsrate $1 + \varepsilon$ löst. Hierbei kann das Polynom p aber durchaus exponentiell in $1/\varepsilon$ sein (und wird es in der Regel auch sein). Eine Folgerung für $\mathcal{APX}$ -harte Probleme ist, dass sie nicht zur Klasse $\mathcal{PTAS}$ gehören können.

Mit Hilfe des PCP-Theorems über probabilistisch verifizierbare Beweise konnten Berman und Karpinski sogar zeigen, dass jeder polynomielle Approximationsalgorithmus für Min-SBR eine Approximationsrate schlechter als 1.008 haben muss, außer wenn $\mathcal{P} = \mathcal{NP}$ gilt.

Theorem 6.3 Jeder polynomielle Approximationsalgorithmus für Min-SBR hat eine Approximationsgüte von 1.008 oder schlechter, außer wenn $\mathcal{P} = \mathcal{NP}$.

6.2.2 2-Approximation für Min-SBR

Nach den schlechten Nachrichten aus dem vorherigen Abschnitt wollen wir jetzt noch zeigen, dass es zumindest einen einfachen Algorithmus gibt, der Min-SBR in polynomieller Zeit bis auf einen Faktor 2 löst.

Definition 6.4 Sei $\pi \in S_n$, dann heißt $(0, \pi_1, \dots, \pi_n, n+1)$ die erweiterte Permutation von $\pi \in S_n$.

Beachte, dass wir bei dieser Definition von erweiterten Permutationen die Indizes ab 0 (statt 1) laufen lassen. Wir werden im Folgenden nur erweiterte Permutationen betrachten.

Definition 6.5 Sei $\pi \in S_n$ eine erweiterte Permutation. Ein Breakpoint von π ist ein Paar $(i, i+1)$ mit $i \in [0 : n]$, so dass $|\pi_i - \pi_{i+1}| \neq 1$. Dann bezeichnet $b(\pi)$ die Anzahl der Breakpoints von π .

In $\pi = 0|432|7|1|56|89$ sind die Breakpoints mit blauen Strichen markiert. Offensichtlich gilt $b(\pi) = 5$.

Lemma 6.6 Sei $\pi \in S_n$ eine erweiterte Permutation, dann sind mindestens $\lceil \frac{b(\pi)}{2} \rceil$ viele Reversionen nötig um π in die Identität zu überführen.

Beweis: Die Identität ist die einzige erweiterte Permutation mit $b(\pi) = 0$. Offensichtlich kann jede Reversion maximal 2 Breakpoints eliminieren. Daraus folgt die Behauptung. ■

Aus diesem Lemma wird klar, warum wir die erweiterten statt der gegebenen Permutationen betrachten. Ansonsten gäbe es neben der Identität auch noch die Permutation $(n, n-1, \dots, 2, 1)$ ohne Breakpoints.

Definition 6.7 Sei $\pi \in S_n$ eine erweiterte Permutation.

- Ein Strip von π ist eine maximal konsekutive Teilfolge $(\pi_i, \dots, \pi_j)$ von π ohne einen Breakpoint.
- Ein Strip heißt steigend bzw. fallend, wenn $\pi_i < \dots < \pi_j$ bzw. $\pi_i > \dots > \pi_j$ gilt.
- Einelementige Strips sind fallend, außer sie sind (0) oder $(n+1)$. Die Strips (0) und $(n+1)$ sind steigend.

In Abbildung 6.12 sind für die Permutation π die steigenden Strips grün und die fallenden rot angezeichnet.

$\pi:$ 0 | 4 | 3 | 2 | 7 | 1 | 5 | 6 | 8 | 9

Abbildung 6.12: Beispiel: Strips in einer erweiterten Permutation

Lemma 6.8 Sei $\pi \in S_n$ eine erweiterte Permutation.

- a) Gehört $k \in [1 : n]$ zu einem fallenden Strip und $k - 1$ zu einem steigenden Strip, dann existiert eine Reversion ρ , die mindestens einen Breakpoint entfernt, d.h. $b(\pi \circ \rho) < b(\pi)$.
- b) Gehört $\ell \in [1 : n]$ zu einem fallenden Strip und $\ell + 1$ zu einem steigenden Strip, dann existiert eine Reversion ρ , die mindestens einen Breakpoint entfernt, d.h. $b(\pi \circ \rho) < b(\pi)$.

Beweis: zu a): Betrachten wir zuerst den Fall, dass k in der Permutation vor $k - 1$ auftaucht, wie in Abbildung 6.13 illustriert. Da k zu einem fallenden Strip gehört, kann $k + 1$ nicht unmittelbar hinter k stehen. Da $k - 1$ zu einem steigenden Strip gehört, muss $k - 2$ unmittelbar vor $k - 1$ stehen (man überlegt sich leicht, dass in diesem Fall $k - 1 \notin \{0, n + 1\}$). Also kann $k - 1$ auch nicht unmittelbar hinter k stehen und nach k bzw. $k - 1$ ist ein Breakpoint. Wie man leicht sieht, sorgt die dort angezeigte Reversion ρ dafür, dass mindestens ein Breakpoint verschwindet.

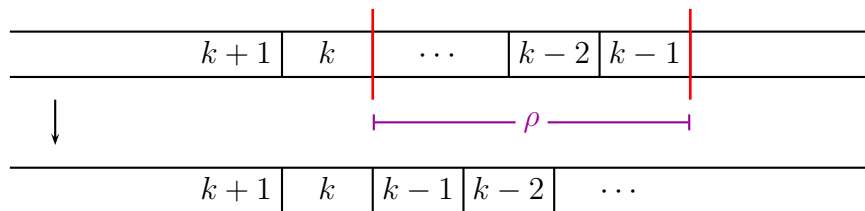


Abbildung 6.13: Skizze: Beweis zu Teil a), 1. Fall

Betrachten wir jetzt den Fall, dass k in der Permutation nach $k - 1$ auftaucht, wie in Abbildung 6.14 illustriert. Da k zu einem fallenden und $k - 1$ zu einem steigenden Strip gehört, kann $k - 1$ nicht unmittelbar hinter k in der Permutation stehen. Da $k - 1$ zu einem steigenden Strip gehört, muss $k - 2$ unmittelbar vor $k - 1$ stehen (außer $k - 1 = 0$). Weiterhin ist unmittelbar nach k bzw. $k - 1$ ein Breakpoint. Wie man leicht sieht, sorgt die dort angezeigte Reversion ρ dafür, dass mindestens ein Breakpoint verschwindet.

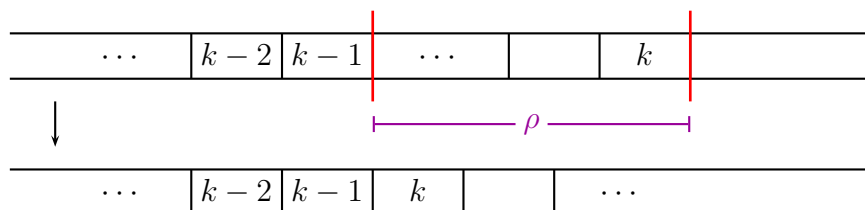


Abbildung 6.14: Skizze: Beweis zu Teil a), Fall 2

zu b): Betrachten wir zuerst den Fall, dass ℓ in der Permutation vor $\ell + 1$ auftaucht, wie in Abbildung 6.15 illustriert. Da ℓ zu einem fallenden und $\ell + 1$ zu einem steigenden Strip gehört, kann ℓ nicht unmittelbar vor $\ell + 1$ in der Permutation stehen. Da $\ell + 1$ zu einem steigenden Strip gehört, muss $\ell + 2$ unmittelbar hinter $\ell + 1$ in der Permutation stehen (außer $\ell = n$). Weiterhin ist unmittelbar vor ℓ bzw. $\ell + 1$ ein Breakpoint. Wie man leicht sieht, sorgt die dort angezeigte Reversion ρ dafür, dass mindestens ein Breakpoint verschwindet.

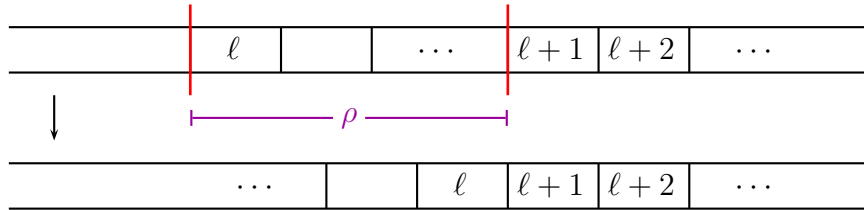


Abbildung 6.15: Skizze: Beweis zu Teil b, Fall 1

Betrachten wir jetzt den Fall, dass ℓ in der Permutation nach $\ell + 1$ auftaucht, wie in Abbildung 6.16 illustriert. Da $\ell + 1$ zu einem steigenden Strip gehört, muss $\ell + 2$ unmittelbar hinter $\ell + 1$ in der Permutation stehen. Weiterhin ist unmittelbar vor ℓ bzw. $\ell + 1$ ein Breakpoint. Wie man leicht sieht, sorgt die dort angezeigte Reversion ρ dafür, dass mindestens ein Breakpoint verschwindet. ■

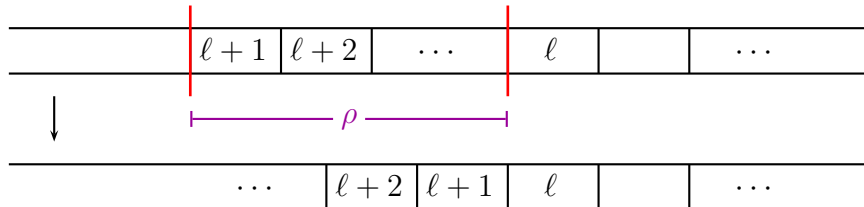


Abbildung 6.16: Skizze: Beweis zu Teil b), Fall 2

Korollar 6.9 Sei $\pi \in S_n$ eine erweiterte Permutation mit einem fallenden Strip. Dann existiert eine Reversion ρ mit $b(\pi \circ \rho) < b(\pi)$.

Beweis: Sei k das kleinste Element, das in einem fallenden Strip liegt. Offensichtlich gilt $k \geq 1$. Dann muss also $k - 1 \geq 0$ in einem steigenden Strip liegen. Mit dem vorhergehenden Lemma folgt die Behauptung.

Sei alternativ ℓ das größte Element, das in einem fallenden Strip liegt. Offensichtlich gilt $\ell \leq n$. Dann muss also $\ell + 1 \leq n + 1$ in einem steigenden Strip liegen. Mit dem vorhergehenden Lemma folgt die Behauptung. ■

Lemma 6.10 Sei $\pi \in S_n$ eine erweiterte Permutation ohne fallenden Strip, dann ist entweder $\pi = \text{id}$ oder es existiert eine Reversion ρ , so dass $\pi \circ \rho$ einen fallenden Strip enthält und $b(\pi \circ \rho) \leq b(\pi)$.

Beweis: Sei $\pi \neq \text{id}$. Sei $(0, \dots, i)$ und $(j, \dots, n+1)$ die beiden steigenden Strips, die die Endsymbole 0 und $n+1$ enthalten. Da $\pi \neq \text{id}$ gilt, muss auch $i+1 < j-1$ gelten. Führt man nun eine Reversion $\rho = (i+1, j-1)$ durch, so kann sich die Anzahl der Breakpoints nicht erhöhen. Nach der Reversion sind also Strips, die nicht $(0, \dots, i)$ und $(j, \dots, n+1)$ enthalten, fallende Strips, da sie vor der Reversion steigend waren. Dies ist in Abbildung 6.17 illustriert. ■

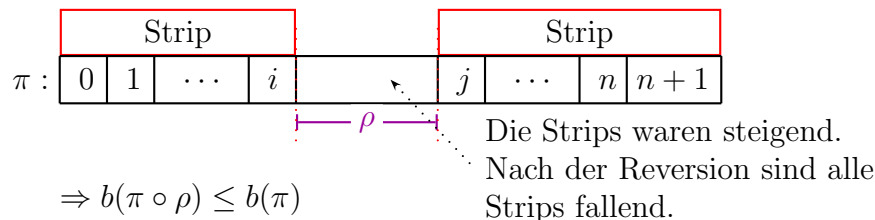


Abbildung 6.17: Skizze: Eine Permutation ohne fallende Strips

Damit erhalten wir eine 4-Approximation, da mit Lemma 6.10 immer ein fallender Strip erzeugt werden kann, ohne die Anzahl der Breakpoints zu erhöhen, und anschließend mit Korollar 6.9 ein Breakpoint entfernt werden kann. Somit kann mit maximal zwei Reversionen ein Breakpoint entfernt werden. Mit der unteren Schranke aus Lemma 6.6 folgt die Behauptung.

16.01.07

Lemma 6.11 Sei $\pi \in S_n$ eine erweiterte Permutation mit mindestens einem fallendem Strip. Sei k bzw. ℓ das kleinste bzw. größte Element in einem fallenden Strip in π . Sei ρ bzw. σ die Reversion, die $k-1$ und k bzw. ℓ und $\ell+1$ zu benachbarten Elementen macht. Wenn weder $\pi \circ \rho$ noch $\pi \circ \sigma$ fallende Strips enthalten, dann ist $\rho = \sigma$ und $b(\pi \circ \rho) = b(\pi) - 2$.

Beweis: Nehmen wir zunächst an, dass k vor $k-1$ in der Permutation π auftritt. Da k das kleinste Element in einem fallenden Strip ist, muss sich $k-1$ in einem aufsteigenden Strip befinden, d.h. $k-2$ steht unmittelbar vor $k-1$ in π . Die zugehörige Reversion ρ erzeugt dann aber einen fallenden Strip, siehe Abbildung 6.18.

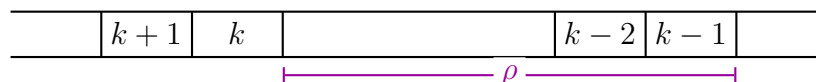


Abbildung 6.18: Skizze: k tritt vor $k-1$ auf

Es gilt also, dass $k - 1$ vor k in der Permutation π auftritt, wie in Abbildung 6.19 illustriert.

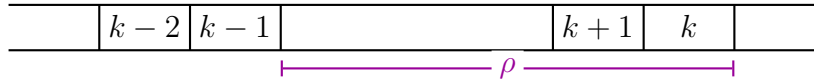


Abbildung 6.19: Skizze: $k - 1$ tritt vor k auf

Analog gilt für die Betrachtung von σ , dass ℓ vor $\ell + 1$ in der Permutation π auftreten muss. Siehe hierzu auch Abbildung 6.20.

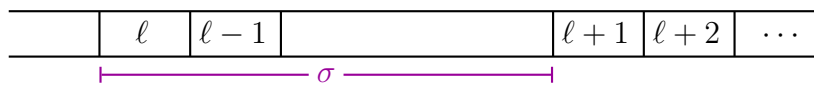


Abbildung 6.20: Skizze: ℓ tritt vor $\ell + 1$ auf

Wir beobachten jetzt Folgendes: Die Reversion σ muss den Strip mit k umdrehen, sonst enthält $\pi \circ \sigma$ den fallenden Strip, der k enthält; Die Reversion π muss den Strip mit ℓ umdrehen, sonst enthält $\pi \circ \rho$ den fallenden Strip, der ℓ enthält.

Für einen Widerspruchsbeweis nehmen wir jetzt an, dass $\rho \neq \sigma$ gilt. Also existiert ein Strip, der nur von σ oder nur von ρ gedreht wird.

Fall 1: Es existiert ein Strip, der nur von ρ gedreht wird. Ist dieser Strip fallend, dann bleibt dieser Strip in $\pi \circ \sigma$ ebenfalls fallend. Dies ist aber ein Widerspruch zu unserer Annahme. Also muss dieser Strip steigend sein. Dann enthält aber $\pi \circ \rho$ einen fallenden Strip und wir erhalten auch in diesem Fall einen Widerspruch zu den Voraussetzungen.

Fall 2: Es existiert ein Strip, der nur von σ gedreht wird. Die Argumentation verläuft analog zum ersten Fall und wir erhalten auch hier einen Widerspruch.

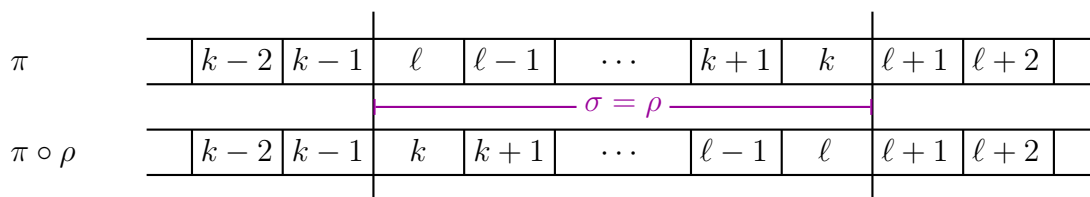


Abbildung 6.21: Skizze: Der Fall $\rho = \sigma$

Also muss $\rho = \sigma$ sein. Diese Situation ist in Abbildung 6.21 illustriert. Wäre $k = \ell$, dann können k und ℓ nicht in einem fallenden Strip liegen. Also gilt $k \neq \ell$ und es werden offensichtlich genau zwei Breakpoints eliminiert. ■

6.2.3 Algorithmus und Laufzeitanalyse

In diesem Abschnitt entwerfen wir basierend auf den Lemmata des letzten Abschnitts einen Approximationsalgorithmus für das Min-SBR Problem. Der Algorithmus selbst ist in Abbildung 6.22 angegeben.

SBR_APPROX (permutation π)

```

{
  while ( $\pi \neq id$ )
  {
    if ( $\pi$  has a decreasing strip)
    {
      let  $k$  be the minimal element of any decreasing strip in  $\pi$ ;
      let  $i$  be the position of  $k$  in  $\pi$ ;
      let  $i'$  be the position of  $k - 1$  in  $\pi$ ;
      if ( $i' < i$ )
        let  $\rho = (i' + 1, i)$ ;
      else
        let  $\rho = (i + 1, i')$ ;
      if ( $\pi \circ \rho$  has no decreasing strip)
      {
        let  $\ell$  be the maximal element of any decreasing strip in  $\pi$ ;
        let  $j$  be the position of  $\ell$  in  $\pi$ ;
        let  $j'$  be the position of  $\ell + 1$  in  $\pi$ ;
        if ( $j' > j$ )
          let  $\rho = (j, j' - 1)$ ;
        else
          let  $\rho = (j', j - 1)$ ;
      }
    }
    else /*  $\pi$  has no decreasing strips */
    {
      let  $(i, i + 1)$  be the leftmost breakpoint in  $\pi$ ;
      let  $(j - 1, j)$  be the rightmost breakpoint in  $\pi$ ;
      let  $\rho = (i + 1, j - 1)$ ;
    }
     $\pi = \pi \circ \rho$ ;
    Reversals = Reversals++ $\rho$ ;
  }
}

```

Abbildung 6.22: Algorithmus: 2-Approximation für Min-SBR

Dabei verwenden wir für Permutationen mit fallenden Strips im Wesentlichen die Aussage von Lemma 6.8 bzw. Korollar 6.9. Hierbei verwenden wir allerdings eine solche Reversion, die entweder eine Permutation mit fallenden Strip generiert oder aber mindestens zwei Breakpoints entfernt, siehe Lemma 6.11. Besitzt die Permutation keinen fallenden Strip, dann verwenden wir die Aussage von Lemma 6.10.

Damit erhalten wir das folgende Theorem.

Theorem 6.12 *Der vorherige Algorithmus liefert eine 2-Approximation für Min-SBR mit Laufzeit $O(n^2)$.*

Beweis: Dass der Algorithmus eine Lösung des Min-SBR Problems liefert, ist klar. Wir müssen nur noch die Aussagen über die Approximationsgüte und die Laufzeit beweisen.

Nach Lemma 6.6 gilt: $d(\pi) \geq \lceil b(\pi)/2 \rceil$, wobei $d(\pi)$ die minimale Anzahl von Reversionen ist, die nötig sind, um π zu sortieren. Wir müssen für eine 2-Approximation also zeigen, dass mit jeder Reversion im Schnitt ein Breakpoint entfernt wird.

Für Permutationen mit fallendem Strip ist das klar, da wir dann jeweils eine Reversion anwenden, die mindestens einen Breakpoint entfernt (siehe Lemma 6.8 bzw. Korollar 6.9).

Wie kann jetzt überhaupt eine Permutation ohne fallende Strips entstehen? Entweder durch eine Reversion, die zuvor zwei Breakpoints eliminiert hat (siehe auch Lemma 6.11), oder in der ursprünglich gegebenen Permutation.

Im ersten Fall können wir die zwei eliminierten Breakpoints mit der folgenden Reversion verrechnen, die keinen Breakpoint entfernt (aber auch keine neuen erzeugt) und einen fallenden Strip generiert.

Außerdem überlegen wir uns leicht, dass die letzte Reversion ebenfalls zwei Breakpoints entfernen muss, da es keine erweiterte Permutation mit genau einem Breakpoint geben kann. Auch in diesem Fall können wir die zwei entfernten Breakpoint der allerletzten Reversion mit dem Nichtentfernen eines Breakpoints der allerersten Reversion errechnen, falls die Eingabepermutation keine fallenden Strips enthalten sollte.

Nun zur Laufzeit. Die Anzahl der Durchläufe der while-Schleife: ist $O(b(\pi)) = O(n)$. Weiterhin beträgt die Arbeit pro Schleifendurchlauf $O(n)$. Somit ist die Gesamtlaufzeit $O(b(\pi)n) = O(n^2)$. ■

6.3 Eine bessere untere Schranke für Min-SBR

In diesem Abschnitt wollen wir eine bessere untere Schranke für die Anzahl benötigter Reversionen angeben.

6.3.1 Breakpoint-Graphen

Zuerst definieren wir den so genannten Breakpoint-Graphen.

Definition 6.13 Sei $\pi \in S_n$ eine erweiterte Permutation. Der Breakpoint-Graph $G(\pi) = (V, E)$ ist definiert durch:

- $V := [0 : n + 1]$;
- $E := R_\pi \cup D_\pi$, wobei

$$\begin{aligned} R_\pi &= \{ \{ \pi_i, \pi_{i+1} \} : |\pi_i - \pi_{i+1}| \neq 1 \}, \\ D_\pi &= \{ \{ \pi_i, \pi_j \} : |\pi_i - \pi_j| = 1 \wedge |i - j| \neq 1 \}. \end{aligned}$$

Die Kanten in R_π bzw. in D_π werden auch als Reality-Edges bzw. Desire-Edges bezeichnet.

Reality-Edges werden im Breakpoint-Graphen als rote Kanten und Desire-Edges als grüne Kanten gezeichnet. Manchmal ist es auch sinnvoll für benachbarte Positionen, die keinen Breakpoint darstellen, im Breakpoint-Graphen die entsprechenden Reality- und Desire-Edges einzuzichnen. Formal handelt es sich hierbei um einen Multi-Graphen (da es zwischen zwei Knoten mehr als eine Kante geben kann).

Definition 6.14 Sei $\pi \in S_n$ eine erweiterte Permutation. Der erweiterte Breakpoint-Graph oder auch das so genannte Reality-Desire-Diagramm ist ein Multi-Graph $G'(\pi) = (V', E')$, der wie folgt definiert ist:

- $V' := [0 : n + 1]$;
- $E' := R'_\pi \cup D'_\pi$, wobei

$$\begin{aligned} R'_\pi &= \{ \{ \pi_i, \pi_{i+1} \} : i \in [0 : n] \}, \\ D'_\pi &= \{ \{ \pi_i, \pi_j \} : |\pi_i - \pi_j| = 1 \}. \end{aligned}$$

Die Kanten in R'_π bzw. in D'_π werden auch als Reality-Edges bzw. Desire-Edges bezeichnet.

Wir halten jetzt noch die folgende fundamentale Beziehung zwischen Breakpoint-Graphen und erweiterten Breakpoint-Graphen fest.

Lemma 6.15 Sei $\pi \in S_n$ eine erweiterte Permutation, dann gilt $G(\pi) \subseteq G'(\pi)$.

In Abbildung 6.23 ist für die Permutation $\pi = (4, 3, 2, 7, 1, 5, 6, 8)$ der zugehörige (erweiterte) Breakpoint-Graph illustriert.

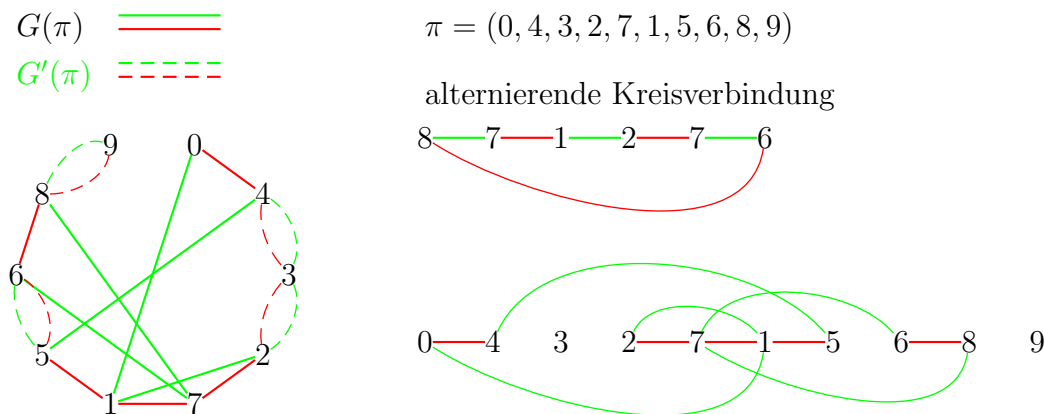


Abbildung 6.23: Beispiel: Ein (erweiterter) Breakpoint-Graph

Wir definieren jetzt noch das Gegenteil von Breakpoints, die so genannten Adjazenzen.

Definition 6.16 Sei $\pi \in S_n$ eine erweiterte Permutation. Eine Adjazenz von π ist ein Paar $(i, i+1)$ mit $i \in [0 : n]$, so dass $|\pi_i - \pi_{i+1}| = 1$. Die Anzahl der Adjazenzen von π wird mit $a(\pi)$ bezeichnet.

Für die folgende Definition einer alternierenden Zyklenzerlegung wiederholen wir zuerst noch den Begriff eines Matchings.

Definition 6.17 Sei $G = (V, E)$ ein ungerichteter Graph. Eine Teilmenge $M \subseteq E$ heißt Matching, wenn für je zwei Kanten $e, e' \in M$ gilt, dass $e \cap e' = \emptyset$.

Für die zu entwickelnde untere Schranke werden Zyklenzerlegungen des Breakpoint-Graphen eine wichtige Rolle spielen.

Definition 6.18 Sei $G = (V, E)$ mit $E = E_1 \cup E_2$ ein 2-kantengefärbter Multi-Graph, in dem jeder Knoten zu gleich vielen Kanten aus E_1 wie zu E_2 inzident ist, d.h. $|\{e \in E_1 : v \in e\}| = |\{e \in E_2 : v \in e\}|$ für jeden Knoten $v \in V$. Eine alternierende Zyklenzerlegung $C(G)$ von G ist eine Menge von Kreisen $C_1, \dots, C_k$ mit:

- $\bigcup_{i=1}^k V(C_i) \subseteq V$,
- $\bigcup_{i=1}^k E(C_i) = E$,
- $E(C_i) \cap E(C_j) = \emptyset$ für alle $i \neq j \in [1 : k]$,
- $E_1 \cap E(C_i)$ und $E_2 \cap E(C_i)$ bilden für alle $i \in [1 : k]$ jeweils ein Matching.

Basierend auf einer alternierenden Zyklenzerlegung definieren wir den wichtigen Parameter $c(\pi)$.

Notation 6.19 Sei $\pi \in S_n$ eine erweiterte Permutation. Dann bezeichnet $c(\pi)$ bzw. $c'(\pi)$ die Anzahl von Kreisen einer maximalen alternierenden Zyklenzerlegung von $G(\pi)$ bzw. $G'(\pi)$.

Wir wollen im Folgenden zeigen, dass dann $d(\pi) \geq b(\pi) - c(\pi)$ gilt.

6.3.2 Elementare Beobachtungen

Bevor wir zum Beweis der unteren Schranke für das Sortieren mit Reversionen kommen, halten wir einige fundamentale Eigenschaften des Breakpoint-Graphen fest.

Beobachtung 6.20 Es gilt $b(\pi) + a(\pi) = n + 1$ für jede erweiterte Permutation $\pi \in S_n$.

Diese Beobachtung entstammt der Tatsache, dass jeder Zwischenraum $(i, i + 1)$ für $i \in [0 : n]$ entweder ein Breakpoint oder eine Adjazenz ist.

Beobachtung 6.21 Sei $\pi \in S_n$ eine erweiterte Permutation und sei $G(\pi)$ (bzw. $G'(\pi)$) der zugehörige (erweiterte) Breakpoint-Graph. $G(\pi)$ besitzt genau dann eine maximale Zyklenzerlegung in k Kreise, wenn $G'(\pi)$ eine maximale Zyklenzerlegung in $k + a(\pi)$ Kreise besitzt.

Den Beweis dieser Beobachtung überlassen wir dem Leser zur Übung.

Beobachtung 6.22 Sei $\pi \in S_n$ eine erweiterte Permutation und sei $G'(\pi)$ der zugehörige erweiterte Breakpoint-Graph. Für jeden Knoten $v \in [1 : n]$ gilt $\deg(v) = 4$ in $G'(\pi)$. Es gilt sogar, dass jeder Knoten zu genau zwei roten und zu genau zwei grünen Kanten inzident ist.

Diese Beobachtung folgt unmittelbar aus der Definition des erweiterten Breakpoint-Graphen.

18.01.07

Lemma 6.23 Sei $\pi \in S_n$ eine erweiterte Permutation und sei $G'(\pi)$ der zugehörige erweiterte Breakpoint-Graph. Eine Reversion ρ erzeugt eine neue alternierende Zyklenzerlegung von $G'(\pi \circ \rho)$, in der sich die Anzahl der Kreise gegenüber $c'(\pi)$ um höchstens eins ändert.

Beweis: Eine Reversion kann nur die Reality-Edges des erweiterten Breakpoint-Graphen verändern. Durch die Reversion werden zwei Reality-Edges entfernt und zwei neue hinzugefügt. Genauer gesagt, werden zwei entfernte Reality-Edges (i, i') und (j, j') durch (i, j) und $(i'j')$ ersetzt. Dabei kann die Anzahl der Kreise sich nur um eins ändern. Dies kann man genauer der Abbildung 6.24 entnehmen. Dort sind

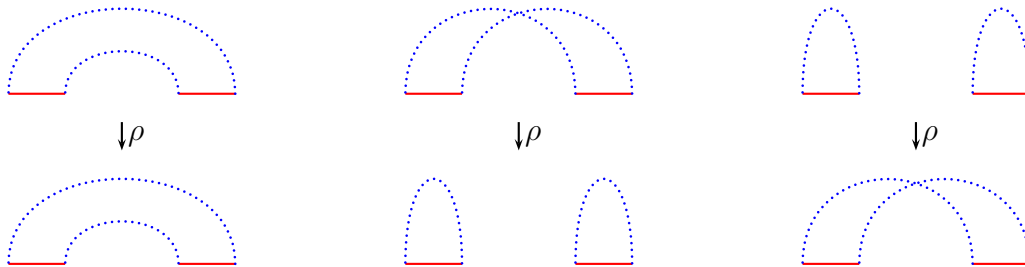


Abbildung 6.24: Skizze: Änderung der alternierenden Zyklenzerlegung durch eine Reversion

nur die Pfade der Zyklen (in blau) skizziert, die die (beiden roten) Breakpoints bzw. Adjazenzen verbinden, auf denen die Reversion operiert. Diese können nur entweder ein oder zwei Zyklen sein. ■

Nun beweisen wir die fundamentale Eigenschaft, dass jede Reversion die Anzahl von Kreisen in einer maximalen alternierenden Zyklenzerlegung um höchstens eins ändern kann.

Lemma 6.24 Sei $\pi \in S_n$ eine erweiterte Permutation und sei ρ ein Reversion. Dann gilt $|c'(\pi \circ \rho) - c'(\pi)| \leq 1$.

Beweis: Betrachten wir eine maximale alternierende Zyklenzerlegung von $G'(\pi)$ mit $c'(\pi)$ Zyklen. Nach dem vorherigen Lemma erzeugt eine Reversion ρ eine alternierende Zyklenzerlegung von $G'(\pi \circ \rho)$, deren Anzahl nur um eins niedriger ist, d.h. $c'(\pi \circ \rho) \geq c'(\pi) - 1$.

Betrachten wir eine maximale alternierende Zyklenzerlegung von $G'(\pi \circ \rho)$ mit $c'(\pi \circ \rho)$ Zyklen. Nach dem vorherigen Lemma erzeugt eine Reversion ρ eine alternierende Zyklenzerlegung von $G'(\pi \circ \rho \circ \rho)$, deren Anzahl nur um eins niedriger ist, d.h. $c'(\pi \circ \rho \circ \rho) \geq c'(\pi \circ \rho) - 1$. Da Reversionen zu sich selbst invers sind, d.h. $\rho^{-1} = \rho$, gilt dann auch $c'(\pi) = c'(\pi \circ \rho \circ \rho^{-1}) = c'(\pi \circ \rho \circ \rho) \geq c'(\pi \circ \rho) - 1$.

Insgesamt gilt also $|c'(\pi) - c'(\pi \circ \rho)| \leq 1$ und somit die Behauptung. ■

Theorem 6.25 Sei $\pi \in S_n$ eine erweiterte Permutation und sei ρ eine Reversion. Es gilt $|b(\pi \circ \rho) - c(\pi \circ \rho) - (b(\pi) - c(\pi))| \leq 1$.

Beweis: Im erweiterten Breakpoint-Graphen $G'(\pi)$ gilt $|c'(\pi \circ \rho) - c'(\pi)| \leq 1$, wobei $c'(\pi)$ die Anzahl von Zyklen einer maximalen Zyklenzerlegung des erweiterten Breakpoint-Graphen $G'(\pi)$ ist. Dann gilt mit der vorherigen Beobachtung 6.21: $c'(\pi) = c(\pi) + a(\pi)$, wobei $c(\pi)$ die Anzahl der Zyklen einer maximalen Zyklenzerlegung des normalen Breakpoint-Graphen $G(\pi)$ ist. Somit gilt

$$|c(\pi \circ \rho) + a(\pi \circ \rho) - c(\pi) - a(\pi)| \leq 1.$$

Weiter gilt mit der vorherigen Beobachtung 6.20, dass $a(\pi) = n + 1 - b(\pi)$. Somit gilt

$$|c(\pi \circ \rho) + n + 1 - b(\pi \circ \rho) - c(\pi) - (n + 1) + b(\pi)| \leq 1.$$

Somit gilt auch

$$|c(\pi \circ \rho) - b(\pi \circ \rho) - (c(\pi) - b(\pi))| \leq 1$$

und somit die Behauptung. ■

6.3.3 Die untere Schranke

Jetzt können wir die angekündigte untere Schranke beweisen.

Theorem 6.26 Für jede erweiterte Permutation $\pi \in S_n$ gilt $d(\pi) \geq b(\pi) - c(\pi)$.

Beweis: Sei $\rho_k, \dots, \rho_1$ eine kürzeste Folge von Reversionen, die π sortiert, und sei $\pi_k = \pi$, $\pi_{i-1} = \pi_i \circ \rho_i$ für $i \in [1 : k]$ sowie $\pi_0 = \text{id}$. Dann gilt für $i \in [1 : k]$ aufgrund von Theorem 6.25 ($1 \geq c(\pi_i \circ \rho_i) - b(\pi_i \circ \rho_i) - (c(\pi_i) - b(\pi_i))$):

$$d(\pi_i) = d(\pi_{i-1}) + 1 \geq d(\pi_{i-1}) + c(\pi_i \circ \rho_i) - b(\pi_i \circ \rho_i) - (c(\pi_i) - b(\pi_i))$$

Daraus folgt sofort:

$$\begin{aligned} d(\pi_i) - (b(\pi_i) - c(\pi_i)) &\geq d(\pi_{i-1}) - (b(\pi_{i-1}) - c(\pi_{i-1})) \\ &\geq d(\pi_{i-2}) - (b(\pi_{i-2}) - c(\pi_{i-2})) \\ &\vdots \\ &\geq d(\pi_0) - (b(\pi_0) - c(\pi_0)) \\ &= 0. \end{aligned}$$

Mit $i = k$ folgt sofort

$$d(\pi) - (b(\pi) - c(\pi)) = d(\pi_k) - (b(\pi_k) - c(\pi_k)) \geq 0.$$

Also gilt $d(\pi) \geq b(\pi) - c(\pi)$. ■

Korollar 6.27 Für jede erweiterte Permutation $\pi \in S_n$ gilt $d(\pi) \geq n + 1 - c'(\pi)$.

Beweis: Nach dem vorherigen Satz gilt $d(\pi) \geq b(\pi) - c(\pi)$. Nach Beobachtung 6.20 gilt $b(\pi) = n + 1 - a(\pi)$ und nach Beobachtung 6.21 gilt $c'(\pi) = c(\pi) + a(\pi)$. Also folgt

$$\begin{aligned} d(\pi) &\geq b(\pi) - c(\pi) \\ &= (n + 1 - a(\pi)) - (c'(\pi) - a(\pi)) \\ &= n + 1 - c'(\pi) \end{aligned}$$

und somit die Behauptung. ■

Wir halten zum Abschluss dieses Abschnittes noch folgende Bemerkungen fest.

- Bereits die Bestimmung einer maximalen Zyklenzerlegung von $G(\pi)$ oder $G'(\pi)$ ist $\mathcal{NP}$ -hart.
- Bei Verwendung von Präfix-Reversionen statt normaler Reversionen heißt das Problem auch *Sorting by Prefix Reversals* bzw. *Pancake Flipping* und wurde zuerst 1979 von Gates und Papadimitriou untersucht. Allerdings wurde dabei nicht der Abstand für ein Paar von Permutationen bestimmt, sondern der

maximale mögliche Abstand für alle Paare von Permutationen (auch Durchmesser genannt). Hierfür gilt

$$d(S_n) := \max \{d(\pi) : \pi \in S_n\} \leq \frac{5}{3}(n+1).$$

- Der angegebene Approximationsalgorithmus verwendet Reversionen, die auf zwei Breakpoints operieren. Es gibt aber Permutationen, wo es besser ist, solche Reversionen zu verwenden, die nicht nur auf Breakpoints operieren. Man kann jedoch zeigen, dass nur das Aufbrechen kurzer Strips hilfreich ist.
- Auch allgemeinere Fragestellungen wurden bereits untersucht. Gegeben sei eine Menge von Generatoren $\rho_1, \dots, \rho_k$ von S_n und eine Permutation $\pi \in S_n$. Gesucht wird die kürzeste Darstellung für π mithilfe der Generatoren der symmetrischen Gruppe. Dieses Problem ist selbst für $k = 2$ bereits $\mathcal{PSPACE}$ -vollständig.

6.4 Sorting by Oriented Reversals

In diesem Abschnitt wollen wir nun so genannte orientierte Permutationen betrachten. Bei einer Reversion wird ja nicht nur eine Gengruppe bzw. ein Gen umgedreht sondern auch die natürliche Ordnung innerhalb dieser Gengruppe bzw. des Gens, das durch die Leserichtung auf der DNA gegeben ist. Von daher erscheint es sinnvoll, sich nicht nur die Gengruppe zu merken, sondern auch deren Leserichtung bzw. auf welchem Strang der DNA sie kodiert sind. Bei einer Reversion werden sie ja dann auf den entsprechenden Komplementärstrang verlagert.

6.4.1 Orientierte Permutationen

Formalisieren wir zunächst solche orientierten Permutationen.

Definition 6.28 Sei $\bar{\pi} = (\bar{\pi}_1, \dots, \bar{\pi}_n) \in [-n : n]^n$ eine Folge ganzer Zahlen mit $|\bar{\pi}| := (|\bar{\pi}_1|, \dots, |\bar{\pi}_n|) \in S_n$ heißt orientierte Permutation. Die Menge aller orientierten Permutationen bezeichnen wir mit $\bar{S}_n$. $\bar{S}_n$ wird auch als orientierte symmetrische Gruppe bezeichnet.

Man kann die Verknüpfung von zwei orientierten Permutationen wie folgt definieren $(\bar{\pi} \circ \bar{\sigma})(i) = \text{sgn}(\bar{\pi}(i)) \cdot \bar{\sigma}(|\bar{\pi}|(i))$, wobei $|\bar{\pi}|(i) = |\bar{\pi}(i)|$ ist. Der Leser möge sich überlegen, dass mit dieser Verknüpfung von Abbildungen auf $\bar{S}_n$ eine Gruppe definiert wird.

Definition 6.29 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation, dann operiert eine orientierte Reversion (oder auch kurz eine Reversion) $\bar{\rho} = (i, j)$ wie folgt:

$$\bar{\pi} \circ \bar{\rho} = (\bar{\pi}_1, \dots, \bar{\pi}_{i-1}, -\bar{\pi}_j, \dots, -\bar{\pi}_i, \bar{\pi}_{i+1}, \dots, \bar{\pi}_n).$$

In Abbildung 6.25 ist ein Beispiel für eine erweiterte orientierte Permutation und eine orientierte Reversion angegeben.

$$\begin{aligned}\bar{\pi} &= (+0, +3, \underbrace{+1, +6, -2}_{\bar{\rho}}, +4, +7) \\ \bar{\pi} \circ \bar{\rho} &= (+0, +3, +2, -6, -1, +4, +7)\end{aligned}$$

Abbildung 6.25: Beispiel: Reversion in einer erweiterten orientierten Permutation

Damit erhalten wir die folgende Problemstellung.

SORTING BY ORIENTED REVERSALS (MIN-SOR)

Eingabe: Eine orientierte Permutation $\bar{\pi} \in \bar{S}_n$.

Gesucht: Eine kürzeste Folge von Reversionen $\bar{\rho}_1, \dots, \bar{\rho}_k$ mit $\bar{\pi} \circ \bar{\rho}_1 \circ \dots \circ \bar{\rho}_k = \text{id}$.

Auch hier definieren wir wieder erweiterte orientierte Permutationen.

Definition 6.30 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation, dann bezeichnet die Folge $(+0, \bar{\pi}_1, \dots, \bar{\pi}_n, +(n+1))$ die zugehörige erweiterte orientierte Permutation.

Auch für orientierte Permutationen können wir Breakpoints und Adjazenzen definieren.

Definition 6.31 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Dann ist $(\bar{\pi}_i, \bar{\pi}_{i+1})$ ein orientierter Breakpoint oder auch kurz ein Breakpoint, wenn kein $j \in [0 : n]$ existiert mit:

- $(\bar{\pi}_i, \bar{\pi}_{i+1}) = (j, j+1)$ oder
- $(\bar{\pi}_i, \bar{\pi}_{i+1}) = (-(j+1), -j)$.

Andernfalls heißt $(\bar{\pi}_i, \bar{\pi}_{i+1})$ eine orientierte Adjazenz oder auch kurz eine Adjazenz. Die Anzahl der orientierten Breakpoints bzw. der Adjazenzen einer orientierten Permutation $\bar{\pi} \in \bar{S}_n$ bezeichnet man mit $\bar{b}(\bar{\pi})$ bzw. $\bar{a}(\bar{\pi})$.

Notation 6.32 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Dann bezeichnet $d(\bar{\pi})$ die maximale Anzahl von orientierten Reversionen, die nötig sind, um $\bar{\pi}$ zu sortieren.

Für eine einfache Weiterbearbeitung von orientierten Permutationen definieren wir jetzt die zugehörige unorientierte Permutation.

23.01.07

Definition 6.33 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation, dann ist $\pi \in S_{2n}$ die zugehörige unorientierte Permutation, wobei $\pi = (\pi_1, \dots, \pi_{2n})$ mit:

- $(\pi_{2i-1}, \pi_{2i}) = (2\bar{\pi}_i - 1, 2\bar{\pi}_i)$, wenn $\bar{\pi}_i > 0$,
- $(\pi_{2i-1}, \pi_{2i}) = (-2\bar{\pi}_i, -(2\bar{\pi}_i - 1))$, wenn $\bar{\pi}_i < 0$.

In Abbildung 6.26 ist ein Beispiel einer orientierten Permutation mit ihrer zugehörigen unorientierten Permutation angegeben.

$$\begin{aligned}\bar{\pi} &= (+0, +3, -2, -1, +4, -5, +6) \\ \pi &= (0, \overset{+3}{\underbrace{5, 6}}, \overset{-2}{\underbrace{4, 3}}, \overset{-1}{\underbrace{2, 1}}, \overset{+4}{\underbrace{7, 8}}, \overset{-5}{\underbrace{10, 9}}, 11)\end{aligned}$$

Abbildung 6.26: Beispiel: zugehörige unorientierte Permutation einer orientierten Permutation

Im Folgenden werden wir immer die zugehörigen erweiterten unorientierten Permutationen und darauf operierende Reversionen betrachten. Dabei operieren Reversionen aber immer nur auf Breakpoints in der zugehörigen erweiterten unorientierten Permutation.

6.4.2 Reality-Desire-Diagramm

Wir definieren nun das Analogon von Breakpoint-Graphen für orientierte Permutationen, so genannte Reality-Desire-Diagrams.

Definition 6.34 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Dann ist das (erweiterte) Reality-Desire-Diagramm (kurz RDD) $G(\bar{\pi})$ bzw. $G'(\bar{\pi})$ der (erweiterte) Breakpoint-Graph $G(\pi)$ bzw. $G'(\pi)$, wobei π die zu $\bar{\pi}$ gehörige erweiterte unorientierte Permutation ist, mit der Einschränkung, dass es für jedes $i \in [1 : n]$ keine Reality- bzw. Desire-Edges zwischen $(2i - 1)$ und $(2i)$ gibt.

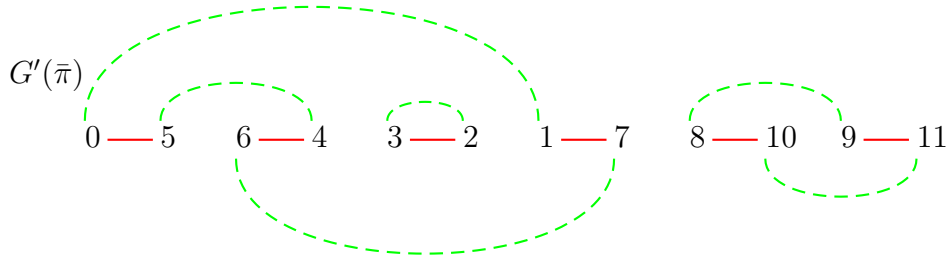


Abbildung 6.27: Beispiel: Erweitertes Reality-Desire-Diagramm einer orientierten Permutation

In Abbildung 6.27 ist für die Permutation $\bar{\pi} = (+3, -2, -1, +4, -5)$ das zugehörige erweiterte Reality-Desire-Diagramm angegeben.

Beobachtung 6.35 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Dann ist jeder Knoten in $G'(\bar{\pi})$ zu genau einer roten und genau einer grünen Kante inzident. Damit gibt es genau eine Zyklenzerlegung in alternierende Kreise und diese kann in Linearzeit ermittelt werden.

Notation 6.36 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Mit $\bar{c}(\bar{\pi})$ bzw. $\bar{c}'(\bar{\pi})$ bezeichnet man die Anzahl von alternierenden Kreisen in $G(\bar{\pi})$ bzw. $G'(\bar{\pi})$.

Mit Hilfe dieser Beobachtungen, folgt das folgende Lemma

Lemma 6.37 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation, dann gilt $d(\bar{\pi}) \geq \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi})$ bzw. $d(\bar{\pi}) \geq n + 1 - \bar{c}'(\bar{\pi})$.

Beweis: Sei $\pi \in S_{2n}$ die zu $\bar{\pi} \in \bar{S}_n$ gehörige unorientierte Permutation. Dann gilt $\bar{b}(\bar{\pi}) = b(\pi)$ sowie $\bar{c}(\bar{\pi}) = c(\pi)$. Offensichtlich kann jede Folge von Reversionen von $\bar{\pi}$ in eine Folge von Reversionen in π verwandelt werden und umgekehrt, sofern für keine Reversion (i, j) (mit $i < j$ gilt) in π gilt, dass i gerade oder j ungerade ist. Denn genau dann werden die künstlich erzeugten Zweier-Strips der zugehörigen unorientierten Permutation nie aufgebrochen.

Insbesondere kann eine optimale Folge von Reversionen für die orientierte Permutation nicht kürzer sein als eine optimale Folge von Reversionen für die zugehörige unorientierte Permutation und die untere Schranke überträgt sich, d.h.

$$d(\bar{\pi}) \geq d(\pi) = b(\pi) - c(\pi) = \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}).$$

Für das erweiterte Reality-Desire-Diagramm gilt $\bar{c}'(\bar{\pi}) = c'(\pi) - n$, da die Adjazenzen innerhalb der neu erzeugten Zweierstrips im erweiterten Reality-Desire-Diagramm nicht gezählt werden. Also gilt:

$$d(\bar{\pi}) \geq d(\pi) \geq 2n + 1 - c'(\pi) = 2n + 1 - (\bar{c}'(\bar{\pi}) + n) = n + 1 - \bar{c}'(\bar{\pi}).$$

Damit sind die Behauptungen gezeigt. ■

6.4.3 Der Overlap-Graph

In diesem Abschnitt betrachten wir noch ein weiteres Hilfsmittel, den so genannten Overlap-Graphen. Zuerst definieren wir noch orientierte und unorientierte Desire-Edges im Reality-Desire-Diagramm.

Definition 6.38 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G(\bar{\pi})$ das zugehörige Reality-Desire-Diagramm. Eine Desire-Edge ist orientiert, wenn die beiden inzidenten Reality-Edges in $G(\bar{\pi})$ in gegenläufiger Richtung durchlaufen werden. Ansonsten heißt die Desire-Edge unorientiert.

In Abbildung 6.28 sind noch einmal alle vier möglichen Fälle einer Desire-Edge mit ihrer Einordnung angegeben.

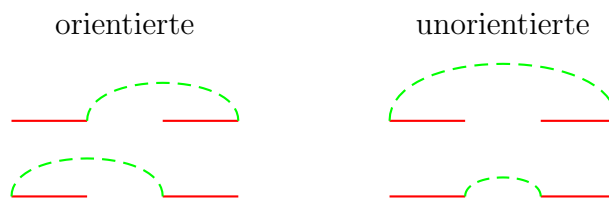
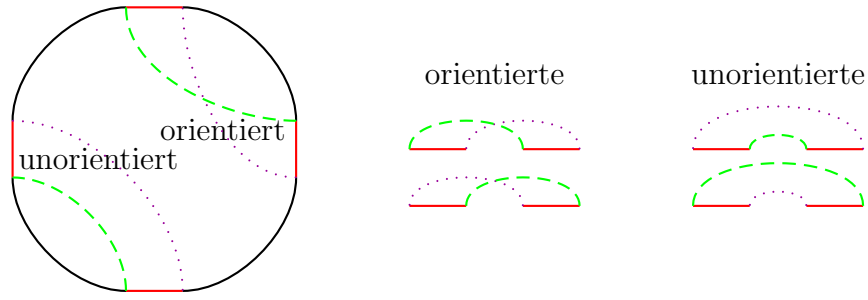


Abbildung 6.28: Skizze: Orientierte und unorientierte Desire-Edges

Nun können wir die Orientierung von Desire-Edges auf Kreise in der eindeutigen alternierenden Zyklenzerlegung im Reality-Desire-Diagramm übertragen.

Definition 6.39 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G(\bar{\pi})$ das zugehörige Reality-Desire Diagramm. Ein Kreis in $G(\bar{\pi})$ heißt orientiert, wenn für eine Desire-Kante die beiden inzidenten Reality-Edges in unterschiedlicher Richtung durchlaufen werden. Die anderen Kreise bezeichnen wir als unorientiert.



Abbildungung 6.29: Skizze: Orientierte und unorientierte Kreise in der Kreis-Darstellung

In Abbildung 6.29 sind orientierte und unorientierte Kreise schematisch in der Kreis- und in der Intervall-Darstellung illustriert.

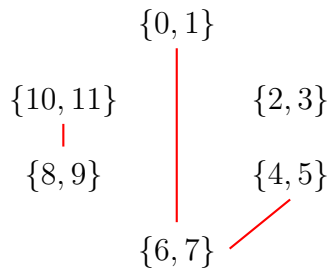
Kommen wir nun zur Definition eines Overlap-Graphen.

Definition 6.40 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation und sei weiter $G(\bar{\pi}) = (V', R_\pi \cup D_\pi)$ das zugehörige Reality-Desire-Diagramm. Der Overlap-Graph $OV(\bar{\pi}) = (V, E)$ ist dann wie folgt definiert:

- $V := D_\pi$,
- $E := \{\{e, e'\} : e, e' \in D_\pi \text{ und } e \text{ schneidet } e' \text{ in } G(\bar{\pi})\}$.

In der obigen Definition bedeutet, dass sich zwei Desire-Edges schneiden, wenn sich ihre Darstellung als Sehnen im Reality-Desire-Diagramm schneiden.

In Abbildung 6.30 ist für das Beispiel $\pi = (+3, -2, -1 + 4, -5)$ der zugehörige Overlap-Graph angegeben.



Abbildungung 6.30: Beispiel: Ein Overlap-Graph für $\pi = (+3, -2, -1 + 4, -5)$

Der Begriff der Orientiertheit lässt sich nun auch auf Zusammenhangskomponenten des Overlap-Graphen ausdehnen (die, wir später sehen werden, Mengen von Kreisen entsprechen).

Definition 6.41 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Ein Knoten von $OV(\bar{\pi})$ heißt orientiert, wenn er zu einer orientierter Desire-Edge korrespondiert. Sonst heißt er unorientiert. Eine Zusammenhangskomponente von $OV(\bar{\pi})$ heißt orientierte Komponente, wenn sie einen orientierte Knoten enthält. Sonst heißt sie unorientierte Komponente.

Nach diesen Definitionen können wir das folgende Lemma zeigen, das uns beim Beweis der Behauptung hilft, dass jede Zusammenhangskomponente im Overlap-Graphen eine Menge von Kreise ist.

Lemma 6.42 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Sei weiter M eine Menge von Desire-Edges im erweiterten Reality-Desire-Diagramm $G(\bar{\pi})$, die zu einer Zusammenhangskomponente von $OV(\bar{\pi})$ korrespondieren. Dann ist $\min(M)$ gerade und $\max(M)$ ungerade, wobei

$$\begin{aligned}\min(M) &:= \min \{i, j : \{\pi_i, \pi_j\} \in M\}, \\ \max(M) &:= \max \{i, j : \{\pi_i, \pi_j\} \in M\}.\end{aligned}$$

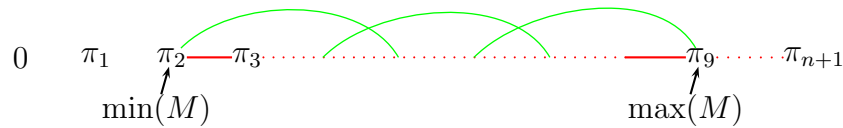


Abbildung 6.31: Skizze: Illustration der Aussage von Lemma 6.42

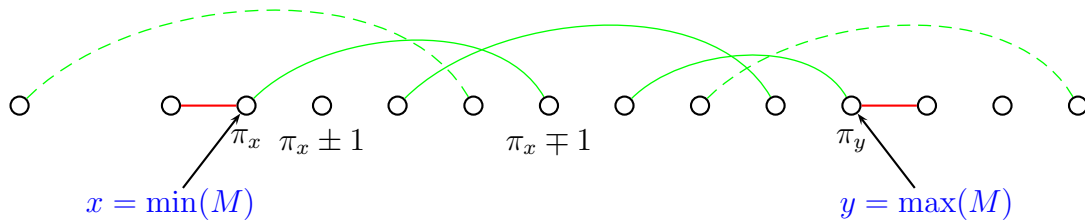
Die Aussage des Lemmas ist in Abbildung 6.31 noch einmal illustriert. Kommen wir jetzt zum Beweis dieser Aussage.

Beweis: Für einen Widerspruchsbeweis nehmen wir an, dass $\min(M)$ ungerade ist. Weiter sei $x = \min(M)$ und $y = \max(M)$ und $\pi \in S_{2n}$ die zu $\bar{\pi}$ gehörige unorientierte Permutation. Betrachten wir nun die zu π_x benachbarten Werte $\pi_x + 1$ bzw. $\pi_x - 1$. Einer der beiden muss an Position $x + 1$ stehen, da x ungerade ist. Der andere Wert ist von π_x über eine Desire-Edge erreichbar. Also gilt, dass sich die Elemente $\pi_x - 1$ und $\pi_x + 1$ im Intervall $[x + 1 : y] = [\min(M) + 1 : \max(M)]$ befinden. Beachte, dass $1 \leq \min(M) \leq \max(M) \leq 2n$ gilt. Diese Situation für das Reality-Desire-Diagramm ist in Abbildung 6.32 illustriert.

Wir betrachten jetzt den Knoten $\pi_x - 1$. Von dort aus folgen wir dem Pfad

$$(\pi_x - 1, \pi_x - 2, \pi_x - 3, \dots, 1, 0),$$

der abwechselnd aus nebeneinander stehenden Knoten der Form $\{2i - 1, 2i\}$ und über Desire-Edges verbundene Knoten besteht. Dieser Pfad muss irgendwann aus

Abbildung 6.32: Skizze: Die Menge M von Desire-Edges

dem Intervall $[\min(M) : \max(M)]$ ausbrechen, da die Kanten aus M nach Definition einer Zusammenhangskomponente das gesamte Intervall $[\min(M) : \max(M)]$ überspannen. Bricht der Pfad mit einer Desire-Edge (siehe die gestrichelte Linie in Abbildung 6.32) aus, dann gehört diese Desire-Edge ebenfalls zu M und es muss $\min(M) < x$ oder $\max(M) > y$ gelten, was den gewünschten Widerspruch liefert. Bricht er über die ein Paar benachbarter Knoten der Form $\{2i - 1, 2i\}$ über $(\pi_y, \pi_{y+1} = \pi_y - 1)$ aus, dann betrachten wir den Pfad

$$(\pi_x + 1, \pi_x + 2, \pi_x + 3, \dots, 2n, 2n + 1),$$

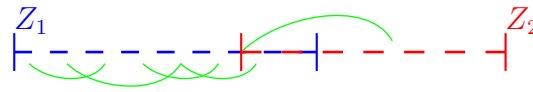
der abwechselnd aus nebeneinander stehenden Knoten der Form $\{2i - 1, 2i\}$ und über Desire-Edges verbundene Knoten besteht. Dieser Pfad muss irgendwann aus dem Intervall $[\min(M) : \max(M)]$ ausbrechen, da die Kanten aus M nach Definition einer Zusammenhangskomponente das gesamte Intervall $[\min(M) : \max(M)]$ überspannen. Bricht der Pfad mit einer Desire-Edge aus, dann gehört diese Desire-Edge ebenfalls zu M und es muss $\min(M) < x$ oder $\max(M) > y$ gelten, was den gewünschten Widerspruch liefert. Ein anderer Ausbruch ist nicht möglich, da $(\pi_y, \pi_{y+1} = \pi_y + 1)$ bereits aufgrund des ersten Falles durch die Werte $(\pi_y, \pi_{y+1} = \pi_y - 1)$ belegt ist.

Die Beweisführung für $\max(M)$ ist analog. ■

25.01.07

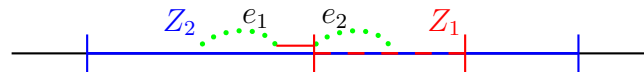
Lemma 6.43 *Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Jede Zusammenhangskomponente von $OV(\bar{\pi})$ korrespondiert zu einer Menge von Desire-Edges, die eine Vereinigung von alternierenden Zyklen in $G'(\bar{\pi})$ ist.*

Beweis: Für einen Widerspruchsbeweis nehmen wir an, dass es einen alternierenden Kreis C in $G'(\bar{\pi})$ gibt, dessen Desire-Edges zu mindestens 2 Zusammenhangskomponenten in $OV(\bar{\pi})$ gehören. Seien Z_1 und Z_2 zwei Zusammenhangskomponenten und $e_1 \in Z_1$ sowie $e_2 \in Z_2$ zwei Desire-Edges, die mit einer Reality-Edge verbunden sind. Man überlegt sich leicht, dass es solche Kanten e_1 und e_2 geben muss.

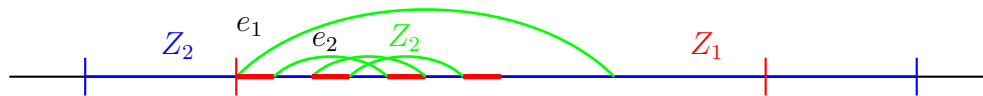
Abbildung 6.33: Skizze: Z_1 und Z_2 überlappen sich

Fall 1: Wir betrachten zuerst den Fall, dass sich Z_1 und Z_2 überlappen. Dies ist in Abbildung 6.33 illustriert. Da Z_1 und Z_2 Zusammenhangskomponenten sind, kann dieser Fall nicht auftreten.

Fall 2: Wir nehmen jetzt an, dass eine Zusammenhangskomponente in der anderen enthalten ist. Betrachten wir zuerst den Unterfall, dass e_1 nicht in Z_2 eingebettet ist. Dies ist in Abbildung 6.34 illustriert. Damit muss aber $\min(Z_1)$ gerade sein und wir erhalten einen Widerspruch zum vorherigen Lemma.

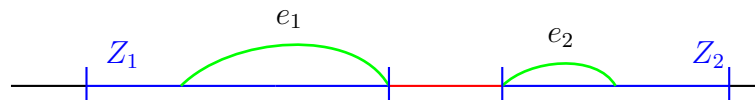
Abbildung 6.34: Skizze: Z_1 ist in Z_2 total in enthalten

Alternativ könnte ein Teil von Z_2 sich innerhalb von Z_1 befinden. Dies ist in Abbildung 6.35 illustriert. Da die Kantengruppe mit e_2 innerhalb von Z_1 zur Zusam-

Abbildung 6.35: Skizze: Z_1 ist in Z_2 total in enthalten

menhangskomponente von Z_2 gehört, muss es eine Folge von sich schneidenden Desire-Edges aus Z_2 geben, die e_2 mit einer Desire-Edge im äußeren Bereich von Z_2 verbindet. Offensichtlich muss eine davon e_1 schneiden und somit wären Z_1 und Z_2 keine verschiedenen Zusammenhangskomponenten.

Fall 3: Nun nehmen wir an, dass die beiden Zusammenhangskomponenten sich in disjunkten Intervallen befinden. Dies ist in Abbildung 6.36 illustriert.

Abbildung 6.36: Skizze: Z_1 und Z_2 sind disjunkt

Auch in diesem Falle muss $\min(Z_2)$ gerade sein und wir erhalten ebenfalls einen Widerspruch zum vorherigen Lemma. ■

Definition 6.44 Eine Zusammenhangskomponente heißt gut, wenn sie einen orientierten Kreis enthält. Ansonsten heißt sie schlecht.

In Abbildung 6.37 sind für ein Beispiel im Reality-Desire-Diagramm die Zusammenhangskomponenten in gute und schlechte Komponenten aufgeteilt.

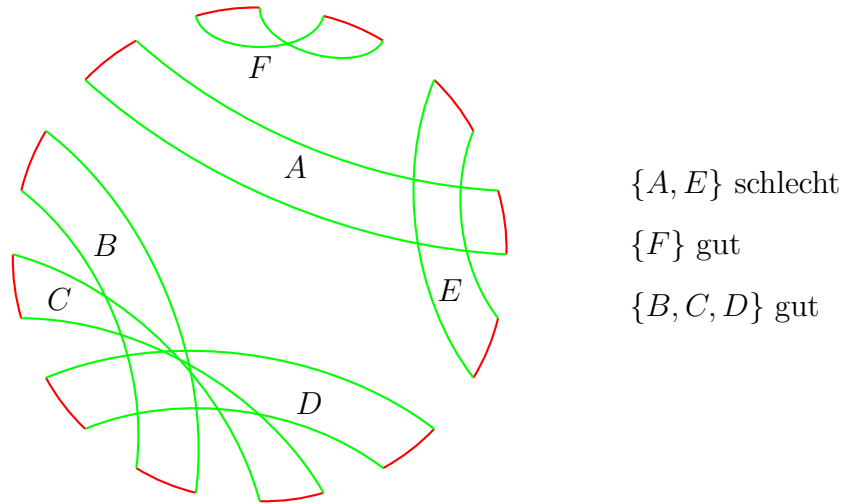


Abbildung 6.37: Skizze: Kreise und Komponenten in Reality-Desire-Diagramm

6.4.4 Hurdles and Fortresses

Nach Definition können wir in guten Komponenten immer eine Reversion auf den beiden unterschiedlich durchlaufenen Reality-Edges, die adjazent zu einer orientierten Desire-Edge sind, anwenden und wir erhöhen die Anzahl der Kreise im erweiterten Reality-Desire-Diagramm. Wie wir später noch sehen werden, lässt sich bei Vorhandensein orientierter Komponenten immer ein solches Paar von Reality-Edges auswählen, so dass man eine kürzesten Folge von Reversionen zur Sortierung der gegebenen Permutation erhält.

Problematisch sind also die schlechten Komponenten, die nur unorientierte Desire-Edges enthalten und somit jede Reversion auf zwei dazu adjazenten Reality-Edges nie die Anzahl der Kreise im erweiterten Reality-Desire-Diagramm erhöhen kann. Weiterhin bleibt eine schlechte Komponente erhalten, bis eine Reversion innerhalb dieser Komponente ausgeführt wird oder auf zwei Komponenten ausgeführt wird.

Definition 6.45 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation und $OV(\bar{\pi})$ der zugehörige Overlap-Graph. Eine schlechte Komponente von $OV(\bar{\pi})$ heißt Hurdle, wenn sie keine anderen schlechten Komponenten trennt. $\bar{h}(\bar{\pi})$ bezeichne die Anzahl von Hurdles in einer gegebenen orientierten Permutation $\bar{\pi}$.

Im Beispiel in Abbildung 6.38 sind die Komponenten A , C , D und F Hurdles, die anderen beiden sind nur einfache schlechte Komponenten. Hierbei sind die Komponenten statt im Overlap-Graphen der besseren Übersichtlichkeit wegen im Reality-Desire-Diagramm angegeben.

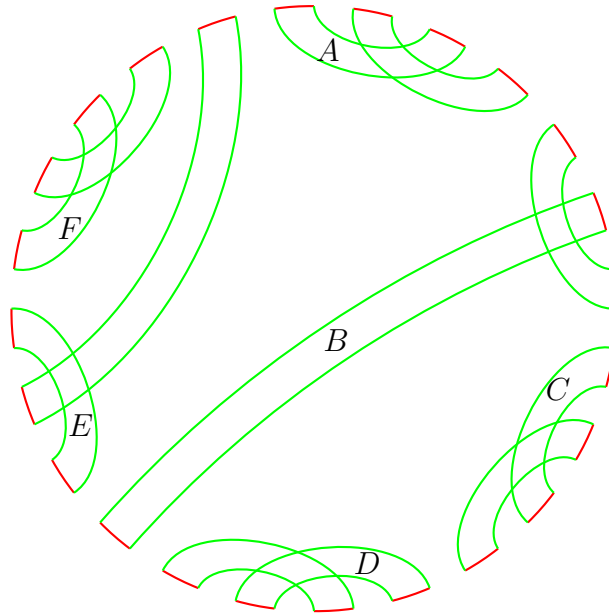


Abbildung 6.38: Skizze: Hurdles und schlechte Komponenten

Hurdles haben die Eigenschaft, dass man auf ihnen Reversionen anwenden muss, um darin orientierte Kreise zu generieren. In Komponenten, die keine Hurdles sind, kann man eine Reversion auf die beiden anderen schlechten Komponente anwenden und damit auch in der darin enthaltenen schlechten Komponente einen orientierten Kreis erzeugen. Hurdles sind somit Komponenten die über die untere Schranke von $\bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi})$ hinaus zusätzliche Reversionen erfordern.

Definition 6.46 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Eine Hurdle, deren Löschung eine Nicht-Hurdle in eine Hurdle transformiert, heißt Super-Hurdle. Andernfalls nennen wir sie eine einfache Hurdle.

Im Beispiel in Abbildung 6.38 ist die Hurdle F eine Super-Hurdle, die Hurdles A , C und D hingegen nicht. Bei Super-Hurdles muss man also aufpassen, da deren Eliminierung neue Hurdles erzeugt.

Definition 6.47 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Die Permutation $\bar{\pi}$ heißt Fortress, wenn sie eine ungerade Anzahl von Hurdles besitzt, die alle Super-Hurdles sind. $\bar{f}(\bar{\pi}) := 1$, wenn $\bar{\pi}$ ein Fortress ist, und $\bar{f}(\bar{\pi}) := 0$ sonst.

6.4.5 Eine untere Schranke für Min-SOR

Unser Ziel in diesem Abschnitt wird sein, die folgende Ungleichung zu zeigen:

$$d(\bar{\pi}) \geq \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \bar{h}(\bar{\pi}) + \bar{f}(\bar{\pi}).$$

Theorem 6.48 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $\bar{\rho}$ ein Reversion, dann gilt:*

$$\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi} \circ \bar{\rho}) - (\bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi})) \leq 1.$$

Beweis: Wir halten zunächst einmal ein paar grundlegende Fakten fest:

- Wenn $\bar{\rho}$ nur auf einer Komponente operiert, so kann sie auch nur an dieser etwas ändern und somit maximal eine Hurdle eliminieren.
- Wenn $\bar{\rho}$ auf zwei Reality-Edges verschiedener Komponenten operiert, dann sind folgende Fälle möglich:
 - a) Wenn beide Komponenten Hurdles sind, so können diese nach Definition nicht durch eine weitere Hurdle getrennt sein und nur diese beiden Hurdles können eliminiert werden.
 - b) Wenn beide Komponenten keine Hurdles sind, so können maximal 2 Hurdles dazwischen liegen, die von $\bar{\rho}$ verändert werden können (eine Reversion einer gesamten Hurdle stellt dabei keine Veränderung dar). Weitere schlechte Komponenten, die dazwischen liegen, würden ansonsten diese beiden Hurdles trennen und eine der beiden wäre keine Hurdle. Also können nur zwei Hurdles eliminiert werden.
 - c) Wenn genau eine Komponente ein Hurdle ist, so kann nach Definition einer Hurdle maximal eine weitere Hurdle zwischen den beiden Komponenten liegen. Also kann auch in diesem Fall $\bar{\rho}$ maximal zwei Hurdles verändern (eine Reversion einer gesamten Hurdle stellt dabei auch hier keine Veränderung dar). Also können ebenso nur zwei Hurdles eliminiert werden.

Aus diesen Überlegungen folgt, dass $\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi}) \geq -2$.

Weiter gilt die folgende Fallunterscheidung:

- 1) Gilt $\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi}) = 1$, dann muss $\bar{\rho}$ auf einem orientierten Kreis operieren, also kann keine Hurdle eliminiert werden, d.h. $\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi}) \geq 0$. Somit gilt:

$$\begin{aligned} \bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi} \circ \bar{\rho}) - (\bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi})) &= \bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi}) - (\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi})) \\ &\leq 1 - 0 = 1. \end{aligned}$$

- 2) Gilt $\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi}) = 0$, dann muss $\bar{\rho}$ auf genau einem Kreis operieren, da sich nur dann die Anzahl der Kreise nicht verändern kann. Also kann $\bar{\rho}$ maximal eine Hurdle entfernen oder hinzufügen. Also gilt insbesondere $\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi}) \geq -1$. Daraus folgt

$$\begin{aligned} \bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi} \circ \bar{\rho}) - (\bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi})) &= \bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi}) - (\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi})) \\ &\leq 0 - (-1) = 1. \end{aligned}$$

- 3) Gilt $\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi}) = -1$, dann muss nach der obigen Überlegung in jedem Falle $\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi}) \geq -2$ gelten, also ist

$$\begin{aligned} \bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi} \circ \bar{\rho}) - (\bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi})) &= \bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi}) - (\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi})) \\ &\leq -1 - (-2) = 1. \end{aligned}$$

Somit folgt die Behauptung. ■

Korollar 6.49 Sei $\bar{\pi} \in \bar{S}_n$, dann gilt

$$d(\bar{\pi}) \geq (n+1) - \bar{c}'(\bar{\pi}) + \bar{h}(\bar{\pi}) = \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \bar{h}(\bar{\pi}).$$

Beweis: Da nach dem vorhergehenden Lemma jede Reversion den folgenden Parameter $\bar{h}(\bar{\pi}) - \bar{c}'(\bar{\pi})$ nur um maximal eins erniedrigen kann und da

$$(n+1) - \bar{c}'(\text{id}) + \bar{h}(\text{id}) = (n+1) - (n+1) + 0 = 0$$

gilt, folgt

$$d(\bar{\pi}) \geq (n+1) - \bar{c}'(\bar{\pi}) + \bar{h}(\bar{\pi})$$

und somit die erste Behauptung.

Man überlegt sich wie schon früher, dass auch hier

$$(n+1) - \bar{c}'(\bar{\pi}) + \bar{h}(\bar{\pi}) = \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \bar{h}(\bar{\pi})$$

gilt und damit folgt die zweite Behauptung. ■

Das letzte Korollar lässt sich noch wie folgt verstärken.

Theorem 6.50 Sei $\bar{\pi} \in \bar{S}_n$, dann gilt

$$d(\bar{\pi}) \geq (n+1) - \bar{c}(\bar{\pi}) + \bar{h}(\bar{\pi}) + \bar{f}(\bar{\pi}).$$

Den Beweis, dass wir für eine Fortress eine weitere Reversion wirklich benötigen, lassen wir hier aus und verweisen den interessierten Leser auf die Originalliteratur.

6.4.6 Sortierung orientierter Komponenten

Im letzten Abschnitt haben wir eine untere Schranke für das Sortieren mit orientierten Permutationen kennen gelernt. Nun wollen wir zeigen, dass man diese untere Schranke auch in jedem Falle erreichen kann. Zuerst wollen wir uns mit dem Fall beschäftigen, dass die Permutation orientierte Komponenten enthält.

In Abbildung 6.39 ist das Reality-Desire-Diagramm einer Permutation angegeben, wobei die orientierten Kanten blau hervorgehoben sind.

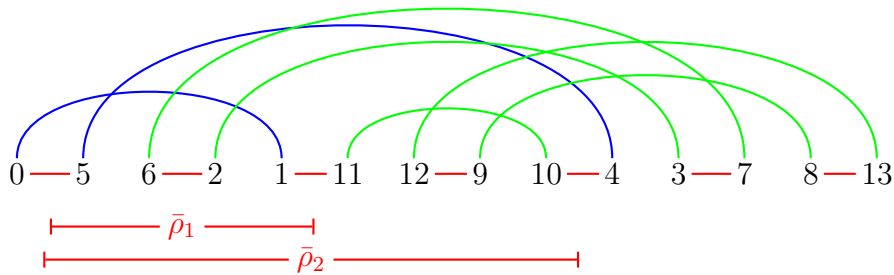


Abbildung 6.39: Beispiel: Orientierte Desire-Edges

Definition 6.51 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und $\bar{\rho}$ eine orientierte Reversion auf $\bar{\pi}$. Der Score von $\bar{\rho}$ ist definiert als die Anzahl orientierter Desire-Edges in $G(\bar{\pi} \circ \bar{\rho})$.

In den Abbildungen 6.40 und 6.41 sind die Ergebnisse der orientierten Reversionen, die zu orientierten Desire-Edges im Beispiel in Abbildung 6.39 gehören, und ihre Scores angegeben.

$$\bar{\pi} \circ \bar{\rho}_1$$

$$\text{Score}(\bar{\rho}_1) = 2$$

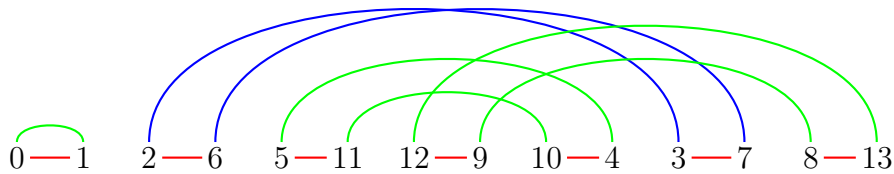
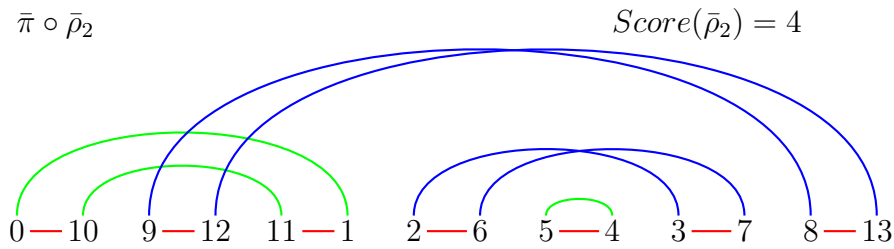


Abbildung 6.40: Beispiel: Score von $\bar{\pi} \circ \bar{\rho}_1$

Der zugehörige Greedy-Algorithmus geht nun wie folgt vor. Solange $G(\bar{\pi})$ eine orientierte Desire-Edge besitzt, wähle eine Reversion zu einer orientierten Desire-Edge mit maximalem Score. Ziel des Abschnittes wird der Beweis des folgenden Satzes sein.

Abbildung 6.41: Beispiel: Score von $\bar{\pi} \circ \bar{\rho}_2$

Theorem 6.52 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Wendet man mit dem oben angegebenen Algorithmus genau k Reversionen auf $\bar{\pi}$ an und erhält man somit $\bar{\pi}'$, dann gilt $d(\bar{\pi}) = d(\bar{\pi}') + k$.

Bevor wir zum Beweis kommen, zeigen wir erst noch, wie man algorithmisch effizient die ausgeführten Reversionen und ihre Änderungen auf dem Overlap-Graphen darstellen kann.

Lemma 6.53 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Eine Desire-Edge, die zu einem Knoten $v \in V(OV(\bar{\pi}))$ korrespondiert, ist genau dann orientiert, wenn v einen ungeraden Grad besitzt.

Für die Permutation in Abbildung 6.39 ist in Abbildung 6.42 der zugehörige Overlap-Graph angegeben. Die orientierten Desire-Edges sind durch Knoten mit einer doppelten Kreislinie dargestellt.

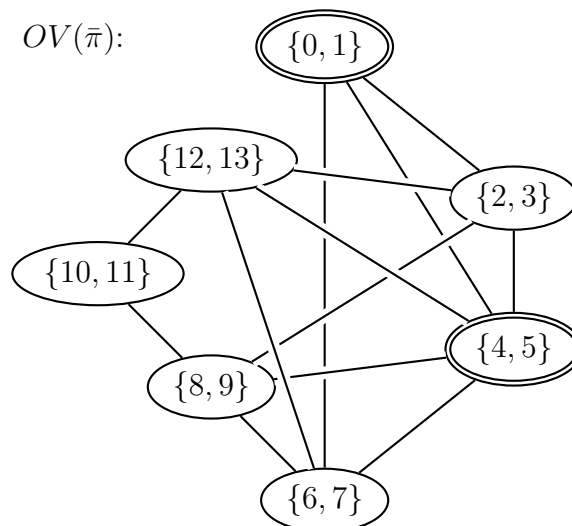


Abbildung 6.42: Beispiel: Overlap-Graph zur Permutation in Abbildung 6.39

Beweis: $\Rightarrow$: Sei $v \in V(OV(\bar{\pi}))$ eine orientierte Desire-Edge im Overlap-Graphen. Sei x die Anzahl der Desire-Edges, die vollständig innerhalb von der zu v gehörigen Desire-Edge verlaufen. Sei weiter y die Anzahl der Desire-Edges, die die zu v gehörige Desire-Edge schneiden. Dann gilt, dass die Anzahl Endpunkte innerhalb der betrachteten Desire-Edge ist $2x + y$. Die Anzahl der Endpunkte innerhalb der betrachteten Intervalls der Desire-Edge muss aber ungerade sein, da v ja orientiert ist. Also ist $2x + y$ ungerade und es muss dann auch y ungerade sein. Da y gerade der Grad des Knotens v im Overlap-Graphen ist, folgt die Behauptung. Dies ist in Abbildung 6.43 noch einmal illustriert.

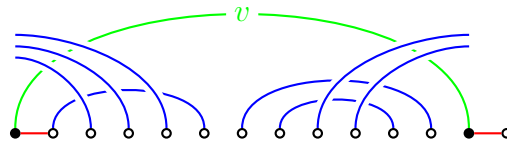


Abbildung 6.43: Skizze: Der Knoten v im Overlap-Graphen mit ungeradem Grad

$\Leftarrow$: Nach Voraussetzung gibt es eine ungerade Anzahl von Endpunkten von Kanten, die die zu v gehörige Desire-Edge schneiden. Die Anzahl Endpunkte von Kanten, die vollständig innerhalb der zu v gehörigen Desire-Edge verlaufen, ist offensichtlich gerade. Also muss die Anzahl der Endpunkte, die innerhalb der zu v gehörigen Desire-Edge liegen, ungerade sein. Diese Tatsache ist in Abbildung 6.43 illustriert.

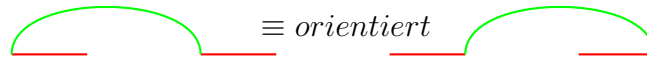


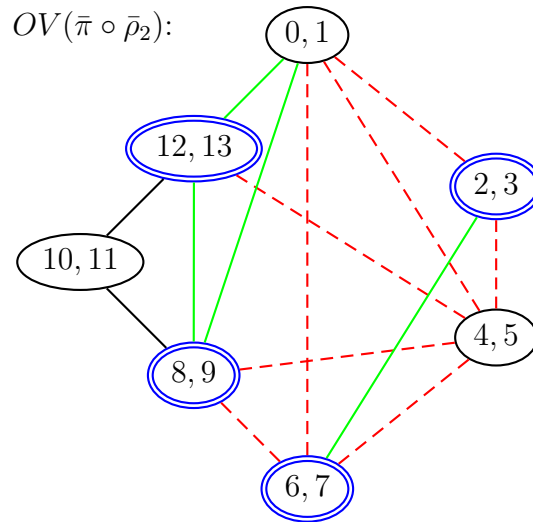
Abbildung 6.44: Skizze: Form der zugehörige Desire-Edge

Somit muss die zu v gehörige Desire-Edge eine der Formen besitzen, wie in Abbildung 6.44 abgebildet. Also ist diese Desire-Edge orientiert. ■

Lemma 6.54 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $\bar{\rho}$ eine orientierte Reversion, die zu einem orientierten Knoten $v \in V(OV(\bar{\pi}))$ korrespondiert. Dann entsteht $OV(\bar{\pi} \circ \bar{\rho})$ aus $OV(\bar{\pi})$ durch Komplementieren des Teilgraphen $\bar{N}(v)$, wobei

$$\begin{aligned} N(v) &:= \{w : \{v, w\} \in E\}, \\ \bar{N}(v) &:= N(v) \cup \{v\}. \end{aligned}$$

In Abbildung 6.45 ist ein Beispiel für die Anwendung der Reversion $\bar{\rho}_2$ auf die Permutation in Abbildung 6.42 und der zugehörige Overlap-Graphen angegeben. Die Kante, die verschwinden, sind rot gestrichelt gezeichnet, die neu entstandenen

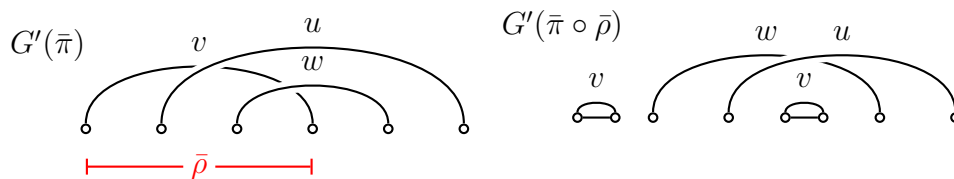
Abbildung 6.45: Beispiel: Der Overlap-Graph von $\bar{\pi} \circ \bar{\rho}_2$

Kanten sind grün gezeichnet. Die neu entstandenen orientierten Knoten sind blau gezeichnet.

Beweis: Zunächst einmal sei $v = (x, x+1)$ die betrachtete Desire-Edge, auf der $\bar{\rho}$ operiert. $\bar{\rho}$ sorgt dafür, dass x und $x+1$ in $\bar{\pi} \circ \bar{\rho}$ benachbart sind, d.h. im erweiterten Reality-Desire-Diagramm befinden sich x und $x+1$ in einem Kreis, der aus genau einer Reality- und einer Desire-Edge besteht. Somit ist der zugehörige Knoten v anschließend auch im Overlap-Graphen ein isolierter Knoten, da die zugehörige Desire-Edge keine anderen Desire-Edges mehr schneiden kann.

Als nächstes halten wir fest, dass nur Desire-Edges von $\bar{\rho}$ verändert werden können, die die zu v gehörige Desire-Edge schneiden oder vollständig innerhalb dieser Desire-Edge liegen.

Fall 1: Betrachten wir zuerst den Fall, dass zwei Knoten u und w beide v schneiden, sich aber selbst nicht schneiden. Die Auswirkung von $\bar{\rho}$ ist in Abbildung 6.46 illustriert, wobei die Lage von v von der Richtung der inzidenten Reality-Edges abhängt. Wie man leicht sieht, schneiden sich u und w nicht mehr mit v , aber durch die Verdrehung von u mit w schneiden sich jetzt u und w .

Abbildung 6.46: Skizze: Wirkung von $\bar{\rho}$, wenn sich u und w nicht schneiden

Fall 2: Betrachten wir jetzt den Fall, dass sich u und w sowohl mit v als auch mit sich selbst schneiden. Die Auswirkung von $\bar{\rho}$ ist in Abbildung 6.47 illustriert, wobei die Lage von v von der Richtung der inzidenten Reality-Edges abhängt. Wie man

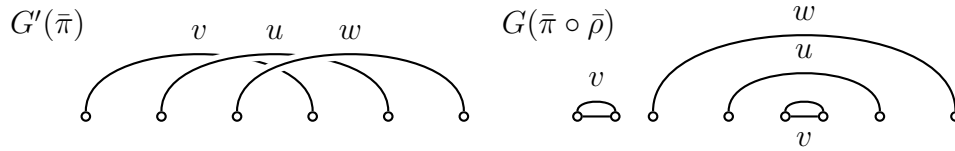


Abbildung 6.47: Skizze: Wirkung von $\bar{\rho}$, wenn sich u und w schneiden

leicht sieht, schneiden sich u und w nicht mehr mit v , aber durch die Entdrehung von u mit w schneiden sich jetzt auch u und w nicht mehr.

Fall 3: Jetzt betrachten wir den Fall, dass sich u mit v schneidet, aber v nicht mit w . Wenn sich u und w nicht schneiden, ändert sich bei der Reversion $\bar{\rho}$ daran nichts. Also nehmen wir an, dass sich u und w schneiden. Nach der Operation $\bar{\rho}$ schneiden sich u und w weiterhin.

Fall 4: Zum Schluss nehmen wir an, dass sich u und w nicht mit v schneiden. Wenn jetzt u und w außerhalb von v liegt, kann $\bar{\rho}$ an der Schnitt-Relation nichts ändern. Liegt u und w innerhalb von v , dann liegt die in Abbildung 6.48 skizzierte Situation vor.

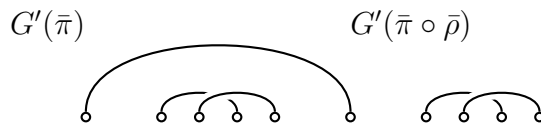
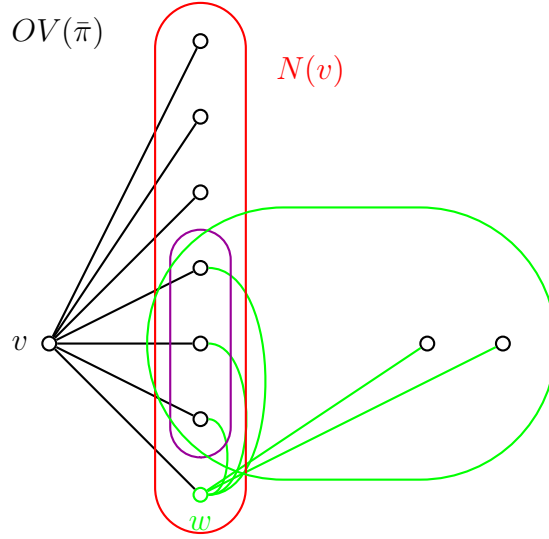


Abbildung 6.48: Skizze: v schneidet sich nicht mit u und w

Aus allen Fällen folgt, dass sich die Kanten im Teilgraphen $(\bar{N}(v), E \cap \bar{N}(v))$ komplementieren und die Behauptung ist gezeigt. ■

Lemma 6.55 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $\bar{\rho}$ eine Reversion mit maximalem Score, das zu einer orientierten Desire-Edge korrespondiert, die im Overlap-Graph durch den Knoten $v \in V(OV(\bar{\pi}))$ gegeben ist. Dann ändert jeder zu v adjazente Knoten in $OV(\bar{\pi})$ seine Orientierung in $OV(\bar{\pi} \circ \bar{\rho})$.

Beweis: Nach Voraussetzung ist $v \in V(OV(\bar{\pi}))$ ein orientierter Knoten und somit ist $\deg(v)$ ungerade. Wir betrachten jetzt die Nachbarschaft von v in $OV(\bar{\pi})$, wie in Abbildung 6.49 dargestellt.

Abbildung 6.49: Skizze: Die Nachbarschaft eines Knotens v

Da v orientiert ist, muss $|N(v)| = 2k + 1$ für ein $k \in \mathbb{N}$ gelten. Sei jetzt $w \in N(v)$ eine beliebiger zu v adjazenter Knoten. Weiter sei $j := |N(w) \cap N(v)|$. Im Folgenden bezeichne $\deg(v)$ bzw. $\deg'(v)$ den Grad des Knotens v im Overlap-Graphen $OV(\bar{\pi})$ bzw. $OV(\bar{\pi} \circ \bar{\rho})$. Es gilt

$$\deg'(w) = \deg(w) - 1 - j + (2k - j) = \deg(w) + 2(k - j) - 1.$$

Offensichtlich ist $2(k - j) - 1$ ungerade und somit muss sich die Parität von w ändern. ■

Notation 6.56 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Dann bezeichnet für ein festes $v \in V(OV(\bar{\pi}))$:

- $T(\bar{\pi})$ die Anzahl orientierter Knoten in $V(OV(\bar{\pi}))$;
- $O_v(\bar{\pi})$ die Anzahl orientierter Knoten in $V(OV(\bar{\pi})) \cap N(v)$;
- $U_v(\bar{\pi})$ die Anzahl unorientierter Knoten in $V(OV(\bar{\pi})) \cap N(v)$.

Mit Hilfe dieser Notation können wir das Ergebnis aus dem letzten Beweis auch wie folgt formulieren.

Lemma 6.57 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $\bar{\rho}$ eine orientiertes Reversion auf $\bar{\pi}$, die zu einer orientierten Desire-Edge korrespondiert, wobei v der zugehörige Knoten in $OV(\bar{\pi})$ ist. Für den Score von $\bar{\rho}$ gilt dann

$$\text{Score}_{\bar{\pi}}(\bar{\rho}) = T(\bar{\pi}) + U_v(\bar{\pi}) - O_v(\bar{\pi}) - 1.$$

Beweis: Der Term $T(\bar{\pi})$ beschreibt die Anzahl der orientierten Knoten in $OV(\bar{\pi})$. Der Term $U_v(\bar{\pi}) - O_v(\bar{\pi}) - 1$ beschreibt die relative Änderung, die durch die Reversion $\bar{\rho}$ ausgelöst wird. ■

Definition 6.58 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Eine orientierte Reversion $\bar{\rho}$ heißt sicher, wenn

$$\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi} \circ \bar{\rho}) > \bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi}).$$

Damit können wir den folgenden fundamentalen Satz beweisen.

Theorem 6.59 Eine zu einer orientierten Desire-Edge korrespondierenden Reversion mit maximalem Score ist sicher.

Beweis: Für den Beweis werden wir folgende stärkere Aussage zeigen: Es wird keine neue unorientierte Komponente erzeugt. Da keine neuen unorientierten Komponenten erzeugt werden, kann auch keine neue Hurdle erzeugt werden. Dies ist ausreichend, da eine orientierte Reversion immer die Anzahl der Kreise im erweiterten Reality-Desire $G'(\bar{\pi})$ erhöht.

Für einen Widerspruchsbeweis nehmen wir an, dass eine Reversion mit maximalen Score, die zur orientierten Desire-Edge $v \in V(OV(\bar{\pi}))$ gehört, eine neue unorientierte Komponente C erzeugt.

Dann muss einer der Knoten aus C in $OV(\bar{\pi})$ zu v adjazent gewesen sein, da sonst C schon in $OV(\bar{\pi})$ unorientiert gewesen wäre. Sei dieser Knoten $w \in V(OV(\bar{\pi}))$. Sei weiter $\text{Score}(v) = T + U_v - O_v - 1$ und $\text{Score}(w) = T + U_w - O_w - 1$.

Alle unorientierten Knoten aus $N(v)$ sind adjazent zu w in $OV(\bar{\pi})$. Sonst gäbe es einen orientierten Knoten in $OV(\bar{\pi} \circ \bar{\rho})$, der zu w in $OV(\bar{\pi})$ adjazent ist. Dann wäre C orientiert, was einen Widerspruch liefert. Somit gilt $U_w \geq U_v$.

Alle orientierten Knoten in $N(w)$ sind adjazent zu v in $OV(\bar{\pi})$. Sonst gäbe es einen orientierten Knoten $u \in N(w)$ in $OV(\bar{\pi})$, der zu v nicht adjazent ist. Da $u \notin \bar{N}(v)$ ist, wäre aber u auch in $OV(\bar{\pi} \circ \bar{\rho})$ orientiert und zu w adjazent. Dann wäre C orientiert, was einen Widerspruch liefert. Damit gilt $O_w \leq O_v$.

Somit gilt:

$$\text{Score}(v) = T + U_v - O_v - 1 \leq T + U_w - O_w - 1 = \text{Score}(w).$$

Wenn $\text{Score}(w) > \text{Score}(v)$, erhalten wir einen Widerspruch zur Voraussetzung. Also gilt $U_v = U_w$ und $O_v = O_w$. Dann gilt aber $\bar{N}(v) = \bar{N}(w)$.

Die Reversion $\bar{\rho}$ würde also aus w einen isolierten Knoten machen. Dies ist ein Widerspruch zur Tatsache, dass w in der unorientierten Komponente c enthalten ist. ■

Aus dem Beweis folgt unmittelbar auch das folgende Korollar.

Korollar 6.60 *Eine zu einer orientierten Desire-Edge korrespondierenden Reversion mit maximalem Score erzeugt keine neuen unorientierten Komponenten.*

Weiterhin gilt als Folge der Definition 6.58 zusammen mit Theorem 6.48

Korollar 6.61 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Eine orientierte Reversion $\bar{\rho}$ ist genau dann sicher, wenn*

$$\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{h}(\bar{\pi} \circ \bar{\rho}) = \bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi}) + 1.$$

Aus diesem Korollar folgt zusammen mit Theorem 6.59 sofort die Behauptung aus Theorem 6.52.

01.02.07

6.4.7 Eliminierung von Hurdles

In diesem Abschnitt wollen wir nun beschreiben wie wir mit Hurdles umgehen.

Definition 6.62 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Eine Reversion, die auf zwei Reality-Edges in zwei verschiedenen, nicht benachbarten Hurdles operiert, wird als Hurdle-Merging bezeichnet.*

Ein solches Hurdle-Merging ist in Abbildung 6.50 schematisch dargestellt. Hierbei ist links die Kreis-Darstellung des Reality-Desire-Diagrams angegeben, wobei Kreise dort Hurdles darstellen. Im rechten Teil ist die Intervall-Darstellung angegeben, oben vor und unten nach der Reversion.

Man kann sich leicht überlegen, dass eine Reversion, die auf zwei Kreisen der beiden Hurdles operiert, aus diesen beiden Kreisen einen macht. Weiterhin sorgt eine solche Reversion dafür, dass eine orientierte Desire-Edge entsteht. Somit nimmt nach einem Hurdle-Merging die Zahl der Hurdles um 2 ab, während die Anzahl der Kreise im erweiterten Reality-Desire-Diagramm um eins sinkt. Somit gilt aber insgesamt

$$(\bar{h}(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi} \circ \bar{\rho})) - (\bar{h}(\bar{\pi}) - \bar{c}'(\bar{\pi})) \leq -1.$$

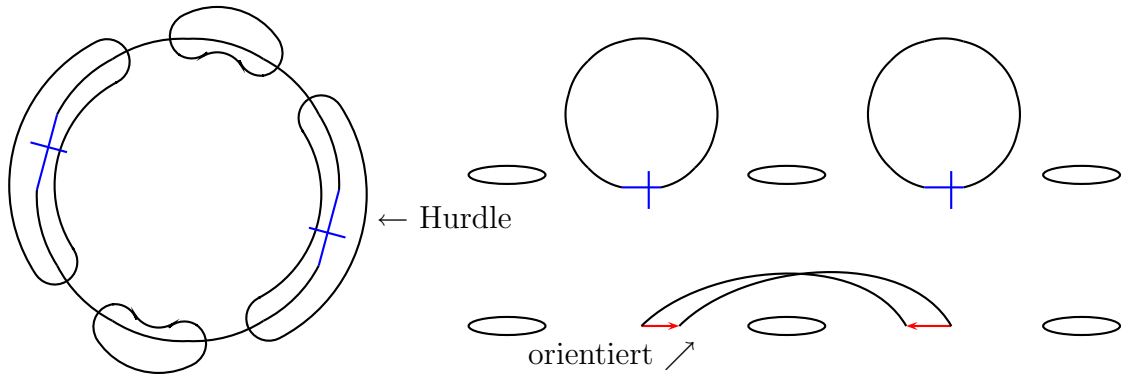


Abbildung 6.50: Skizze: Schematische Darstellung eines Hurdle-Merging

Weiterhin kann man sich überlegen, dass die Eigenschaft ein Fortress zu sein unverändert bleibt, da keine neuen Super-Hurdles generiert werden und maximal zwei Super-Hurdles eliminiert wurden (also die ungerade Anzahl an Hurdles, die alle Super-Hurdles sind, unverändert bleibt).

Hierbei werden nur nicht benachbarte Hurdles gemischt, damit beim Mischen von benachbarten Hurdles nicht eine neue Hurdle entstehen kann. Siehe hierzu auch Abbildung 6.51. Würde man hier die Hurdles A und B mischen, würde aus der

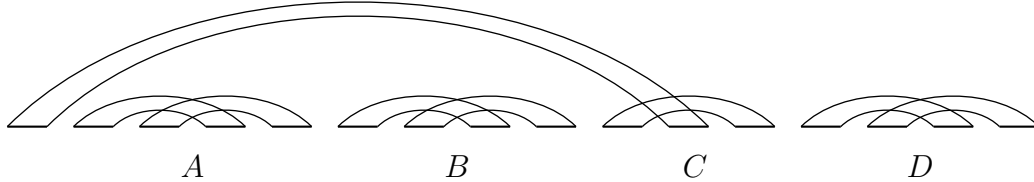


Abbildung 6.51: Skizze: Mögliche Probleme beim Hurdle-Merging

unorientierten Komponente eine neue Hurdle. Durch das Mischen nicht benachbarter Hurdles wird dies vermieden. Man kann also ein zielführendes Hurdle-Merging nur dann durchführen, wenn $\bar{h}(\bar{\pi}) \geq 4$ ist, da ansonsten jedes Paar von Hurdles benachbart wäre.

Definition 6.63 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Eine Reversion, die auf zwei Reality-Edges einer Hurdle operiert, die durch eine Desire-Edge verbunden sind, wird als Hurdle-Cutting bezeichnet.

Wir betrachten zunächst eine Hurdle, die ja eine unorientierte Komponente sein muss. Zuerst überlegt man sich leicht, dass diese eine Desire-Edge mit der Ausrichtung der inzidenten Reality-Edges, wie in Abbildung 6.52 angegeben, existieren muss.

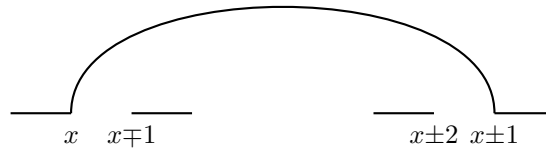


Abbildung 6.52: Skizze: Existenz einer schneidende Desire-Edge in einer Hurdle

Angenommen diese Desire Edge verbindet x und $x+1$ (bzw. $x-1$). Dann folgt man der Folge $(x, x-1, x-2, \dots, 0)$, die alternierend aus Desire-Edges und benachbarten Knoten, die zu Strips von orientierten Werten gehören, beschreiben. Dieser Pfad verläuft zunächst innerhalb der betrachteten Desire-Edge und muss irgendwann einmal aus dem Inneren dieser Desire-Edge ausbrechen. Dies kann nur über eine andere Desire-Edge geschehen, die dann die betrachtete schneiden muss. Also enthält eine Hurdle mindestens zwei sich schneidende Desire-Edges, wie in Abbildung 6.53 illustriert.

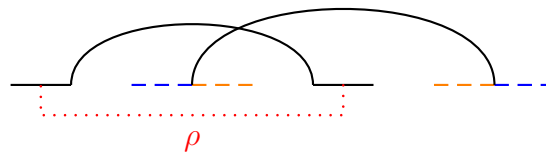


Abbildung 6.53: Skizze: Sich schneidende Desire-Edges in einer Hurdle

Wendet man jetzt die Reversion ρ (wie in Abbildung 6.53 angegeben) an, so sieht man leicht, dass die andere schneidende Desire-Edge anschließend orientiert sein muss, da die zur schneidenden Desire-Edge inzidenten blauen (bzw. orangefarbenen) Reality-Edges nun gegenläufig durchlaufen werden.

Also wird auch hier eine Hurdle eliminiert, da ein orientierter Kreis generiert wird. Die Anzahl der Zyklen steigt auch hier um 1. Ist beim Hurdle-Cutting die betroffene Hurdle einfach, dann entsteht durch Auflösung keine neue Hurdle (im Gegensatz zu Super-Hurdles). Damit bekommt man bei $\bar{h}(\pi) = 3$ Probleme, wenn eben alle Hurdles Super-Hurdles sind. Dann ist die betroffene Permutation aber ein Fortress und man darf sich eine zusätzliche Reversion erlauben, die durch Generierung einer Hurdle auch benötigt wird.

Somit haben wir für $\bar{h}(\pi) \geq 3$ (sofern es sich nicht um drei Super-Hurdles handelt) eine Strategie kennen gelernt, die die untere Anzahl Schranke für die benötigte Anzahl an Reversionen respektiert. Bei genau zwei Hurdles wenden wir ebenfalls wieder ein Hurdle-Merging an, da wir hierbei keine neue Hurdle erzeugen können. Bei einer Hurdle bleibt nichts anderes als ein Hurdle-Cutting übrig. Man überlegt sich hier leicht, dass diese dann eine einfache Hurdle sein muss und keine Super-Hurdle sein kann. Somit können wir auch hier keine neue Hurdle generieren.

6.4.8 Algorithmus für Min-SOR

Zusammenfassend erhalten wir also den in Abbildung 6.54 angegebenen Algorithmus zum Sortieren mit orientierten Reversionen.

SOR (permutation π)

```
{
  while ( $\pi \neq \text{id}$ )
  {
    if ( $OV(\bar{\pi})$  contains an oriented vertex)
      apply oriented reversal with maximal score;
    elseif ( $(\bar{h}(\bar{\pi}) \geq 4) \parallel (\bar{h}(\bar{\pi}) = 2)$ )
      Hurdle-Merging;
    elseif ( $(\bar{h}(\bar{\pi}) \in \{1, 3\}) \&\& (\text{there is a simple hurdle})$ )
      Hurdle-Cutting (on a simple hurdle);
    else /*  $\pi$  is a fortress */
      Hurdle-Merging;
  }
}
```

Abbildung 6.54: Algorithmus: Pseudo-Code für Min-SOR

Halten wir zum Abschluss das erzielte Ergebnis fest, wobei wir den Beweis der Laufzeit dem Leser überlassen.

Theorem 6.64 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Zum Sortieren von $\bar{\pi}$ sind genau $d(\bar{\pi}) = n + 1 - \bar{c}'(\bar{\pi}) - \bar{h}(\bar{\pi}) + \bar{f}(\bar{\pi})$ orientierte Reversionen nötig und ausreichend. Eine optimale sortierende Folge von orientierten Reversionen kann in polynomieller Zeit gefunden werden, d.h. $\text{Min-SOR} \in \mathcal{P}$.*

Zum Abschluss noch ein paar Bemerkungen zu Verbesserungen des angegebenen Algorithmus, insbesondere zu den Laufzeiten:

- Die Laufzeit des ersten vorgeschlagenen Algorithmus von Hannenhalli und Pevzner beträgt $O(n^4)$.
- Eine Verbesserung von Berman und Hannenhalli benötigt $O(n^2\alpha(n))$, wobei $\alpha(n)$ die inverse der Ackermann-Funktion ist (für die Praxis quasi eine Konstante).
- Kaplan, Shamir und Tarjan haben eine weitere Laufzeitverbesserung entwickelt, die einen Zeitbedarf von $O((d(\bar{\pi}) + \alpha(n)) \cdot n) = O(n^2)$ besitzt.

- Bader, Moret und Yan konnten die Laufzeit auf $O(d(\bar{\pi}) \cdot n) = O(n^2)$ senken.
- Eine vereinfachte Darstellung zur Berechnung der Reversal-Distanz (ohne die Bestimmung der sortierenden Reversionen) wurde von Bergeron, Mixtacki und Stoye mit Laufzeit $O(n)$ entwickelt.

Man beachte, dass man in der Regel die Eingabe-Sequenz durch Zusammenfassen von Strips (Bereichen ohne Breakpoints) so modifiziert, dass $d(\bar{\pi}) \geq \bar{b}(\bar{\pi})/2 = n/2$ gilt, also $d(\bar{\pi}) = \Theta(n)$.

06.02.07

6.5 Sorting by Transpositions (*)

In diesem Abschnitt wollen wir nun die Distanz zwischen zwei Genomen bestimmen, wenn nur Transpositionen als elementare Operationen zugelassen sind.

6.5.1 Elementare Definitionen

Zuerst halten wir in diesem Abschnitt noch ein paar grundlegende Definitionen und Beobachtungen fest. Wir bemerken als erstes, dass wir jetzt wieder unorientierte Permutationen betrachten, da wir mit einer Transposition nicht die Orientierung eines Teilstückes ändern können.

SORTING BY TRANSPOSITIONS (MIN-SBT)

Eingabe: Eine Permutation $\pi \in S_n$.

Gesucht: Eine kürzeste Folge von Transpositionen $\tau_1, \dots, \tau_k$ mit $\pi \circ \tau_1 \circ \dots \circ \tau_k = \text{id}$.

Im Folgenden bezeichnen wir mit d_T die *Transpositions-Distanz*. Wie im Falle orientierter Permutationen definieren wir noch den Begriff einer zugehörigen unorientierten Permutation, auch wenn dies zunächst sehr seltsam klingen mag.

Definition 6.65 Sei $\pi \in S_n$ eine Permutation, dann ist $\pi' \in S_{2n}$ die zugehörige unorientierte Permutation, wobei $\pi' = (\pi'_1, \dots, \pi'_{2n})$ mit $(\pi'_{2i-1}, \pi'_{2i}) = (2\pi_i - 1, 2\pi_i)$.

Auch hier werden wir im Folgenden die zugehörige unorientierte Permutation immer als erweiterte zugehörige unorientierte Permutation $\pi' = (0, \pi'_1, \dots, \pi'_{2n}, 2n+1)$ inter-

pretieren. Damit können wir analog zu den orientierten Permutation das Reality-Desire-Diagramm definieren.

Definition 6.66 Sei $\pi \in S_n$ eine erweiterte Permutation und sei π' die zugehörige unorientierte Permutation. Dann ist das (erweiterte) Reality-Desire-Diagramm (kurz RDD) $G(\pi)$ bzw. $G'(\pi)$ der (erweiterte) Breakpoint-Graph $G(\pi')$ bzw. $G'(\pi')$ mit der Einschränkung, dass es für jedes $i \in [1 : n]$ keine Reality- bzw. Desire-Edges zwischen $(2i - 1)$ und $(2i)$ gibt.

In Abbildung 6.55 ist ein Beispiel eines (erweiterten) Reality-Desire-Diagrams für die erweiterte Permutation $(0, 7, 4, 1, 3, 2, 6, 5, 8)$ angegeben.

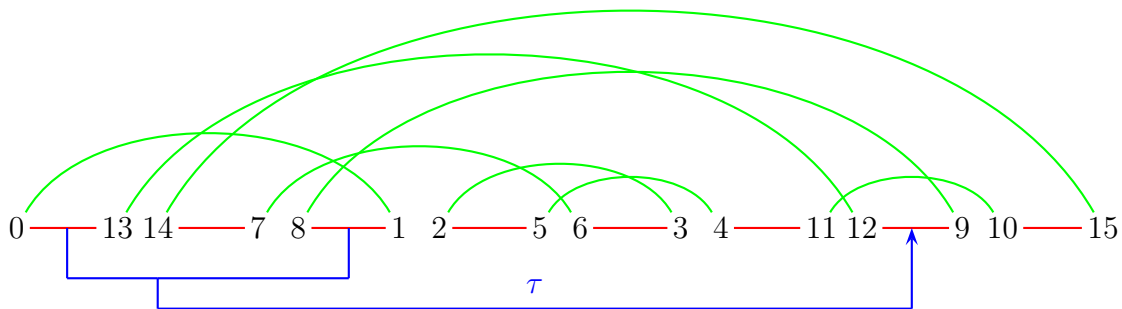


Abbildung 6.55: Beispiel: Reality-Desire-Diagramm für $(0, 7, 4, 1, 3, 2, 6, 5, 8)$

In Abbildung 6.56 ist ein Beispiel angegeben, wie sich die Transposition $(1, 2, 6)$ auf die erweiterte Permutation $(0, 7, 4, 1, 3, 2, 6, 5, 8)$ im (erweiterten) Reality-Desire-Diagramm für die zugehörige unorientierte Permutation auswirkt. Weiterhin ist dort in blau die Zyklenzerlegung in alternierende Kreise angegeben.

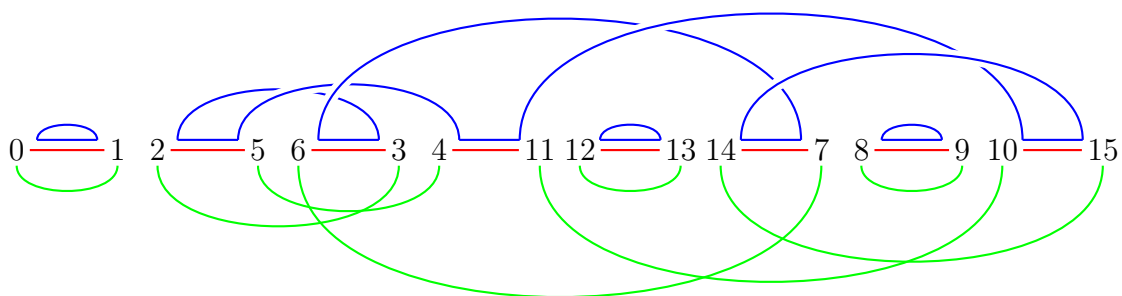


Abbildung 6.56: Beispiel: Erweitertes Reality-Desire-Diagramm für die Permutation $(0, 7, 4, 1, 3, 2, 6, 5, 8)$ nach der Transposition $\tau = (1, 2, 6)$

6.5.2 Untere Schranken für die Transpositions-Distanz

Zunächst wollen wir aus dem Reality-Desire-Diagramm eine untere Schranke für die mindestens benötigte Anzahl von Transpositionen herleiten. Da jede Transposition auf maximal drei Breakpoints operieren kann und somit auch maximal drei Breakpoints eliminieren kann, gilt die folgende Beobachtung.

Beobachtung 6.67 Für jede Permutation $\pi \in S_n$ gilt $d_T(\pi) \geq b(\pi)/3$.

Mit Hilfe des folgenden Lemmas können wir eine bessere untere Schranke herleiten.

Lemma 6.68 Sei $\pi \in S_n$ eine erweiterte Permutation und sei τ eine Transposition auf S_n . Dann gilt $c'(\pi \circ \tau) - c'(\pi) \in \{-2, 0, 2\}$

Der Beweis erfolgt durch eine einfache Fallunterscheidung unter Berücksichtigung der Gestalt von Kreisen im erweiterten Reality-Desire-Diagramm, wie sie im folgenden Abschnitt erläutert werden, und sei dem Leser als Übungsaufgabe überlassen.

Theorem 6.69 Für jede Permutation $\pi \in S_n$ gilt $d_T(\pi) \geq \frac{(n+1)-c'(\pi)}{2} = \frac{b(\pi)-c(\pi)}{2}$.

Diese Schranke lässt sich noch verbessern, wenn man nur Kreise ungerader Länge betrachtet (d.h. mit einer ungeraden Anzahl von Reality- bzw. Desire-Edges).

Notation 6.70 Sei $\pi \in S_n$ eine erweiterte Permutation. Dann bezeichnet $\hat{c}(\pi)$ die Anzahl der Kreise in $G'(\pi)$, die eine ungerade Anzahl von Reality-Edges (oder Desire-Edges) besitzen.

Korollar 6.71 Sei $\pi \in S_n$ eine erweiterte Permutation und sei τ eine Transposition auf S_n . Dann gilt $\hat{c}(\pi \circ \tau) - \hat{c}(\pi) \in \{-2, 0, 2\}$

Der folgende Satz folgt aus der Tatsache, dass die Identität im erweiterten Reality-Desire-Diagramm aus genau $n + 1$ Kreise der Länge 1 (wenn man nur Reality- oder Desire-Edges bei der Länge mitzählt) besteht.

Theorem 6.72 Für jede Permutation $\pi \in S_n$ gilt $d_T(\pi) \geq \frac{(n+1)-\hat{c}(\pi)}{2} = \frac{b(\pi)-\hat{c}(\pi)}{2}$.

6.5.3 Orientierte und Unorientierte Desire-Edges

Da das (erweiterte) Reality-Desire-Diagramm nur aus unorientierten Permutationen erstellt wird (im Gegensatz zu Min-SOR), treten die Paare $(2i-1, 2i)$ für den Wert $+i$ in der zugehörigen unorientierten Permutation nur in dieser Reihenfolge auf. Somit sind im Reality-Desire-Diagramm auch nur Desire-Edges bestimmter Gestalt möglich, wie sie in Abbildung 6.57 dargestellt sind (die rot durchgestrichenen können nicht auftreten).

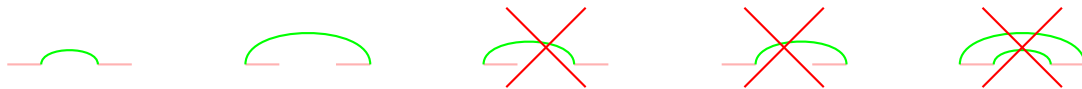


Abbildung 6.57: Skizze: Vorkommende Desire-Edges

Der rechteste Fall in Abbildung 6.57 illustriert, dass zwischen zwei Reality-Edges, die mit zwei Desire-Edges verbunden sind, mindestens eine weitere Reality-Edge vorkommen muss. Andernfalls hätte die entsprechende Teilsequenz der zugehörigen unorientierten Permutation die Form: $(2i-1, 2i, 2j-1, 2j, 2i+1, 2i+2)$. Da wir aber im Reality-Desire-Diagramm zwischen $2j-1$ und $2j$ keine Desire-Edge haben, kann dies nicht eintreten. Damit haben wir folgende Beobachtung gezeigt.

Beobachtung 6.73 Sei $\pi \in S_n$ eine erweiterte Permutation und sei $G'(\pi)$ das zu π gehörige erweiterte Reality-Desire-Diagramm. Seien (π_i, π_{i+1}) und (π_j, π_{j+1}) mit $i+1 < j \in [0 : 2n+1]$ zwei Reality-Edges, deren äußere oder innere Endpunkte mit einer Desire-Edge verbunden sind, dann existiert eine weitere Reality-Edge (π_k, π_{k+1}) mit $k \in [i+1 : j-2]$.

Beim Sortieren mit Transpositionen ist die Definition von orientierten Kreisen etwas anders als beim Sortieren mit Reversionen, da alle Reality-Edges automatisch immer in derselben Richtung durchlaufen werden.

Definition 6.74 Sei $\pi \in S_n$ eine erweiterte Permutation und $G'(\pi)$ das zugehörige erweiterte Reality-Desire-Diagramm. Ein Kreis C in $G'(\pi)$ heißt orientiert, wenn alle bis auf eine Desire-Edge (die zu $\pi_{\max(C)}$ inzident ist) nicht in derselben Richtung durchlaufen werden. Andernfalls heißt der Kreis unorientiert.

Wir merken hier noch an, dass ein Kreis aus einer Reality- und einer Desire-Edge nach dieser Definition unorientiert ist.

In Abbildung 6.58 ist links ein unorientierter Kreis und in der Mitte sowie rechts jeweils ein orientierter Kreis schematisch dargestellt. Man beachte, dass alle unori-

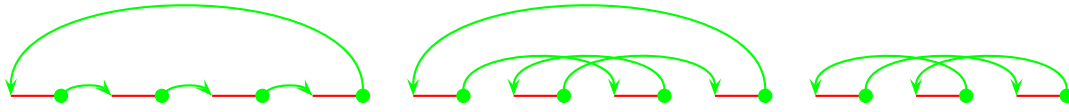


Abbildung 6.58: Skizze: Orientierte und unorientierte Kreise im RDD

entierten Kreise eine Gestalt wie der linkeste Kreis in Abbildung 6.58 haben müssen. Wenn nämlich die Desire-Edge mit dem rechtesten Ende von rechts nach links durchlaufen wird, müssen aufgrund der Unorientiertheit alle andere Desire-Edges von links nach rechts durchlaufen werden.

6.5.4 Eine 2-Approximation

In diesem Abschnitt wollen wir einen Approximationsalgorithmus zur Lösung für Min-SBT entwickeln. Zunächst geben wir die Definition von x -Moves an.

Definition 6.75 Sei $\pi \in S_n$ eine erweiterte Permutation. Eine Transposition τ auf π heißt 2-Move, 0-Move bzw. -2 -Move, wenn $c(\pi \circ \tau) - c(\pi)$ gleich 2, 0 bzw. -2 ist.

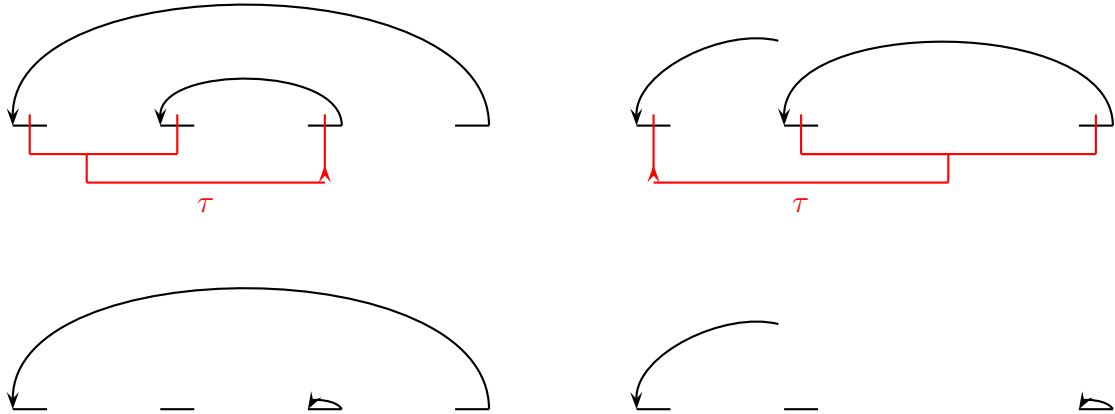
Vermutlich sollte man beim Sortieren mit Transpositionen nach Möglichkeit nur 2-Moves verwenden. Wir merken aber bereits jetzt an, dass es noch ungeklärt ist, ob eine optimale Lösung ohne -2 -Moves auskommen kann.

Das folgende Lemma wird zentral für unseren Approximationsalgorithmus sein.

Lemma 6.76 Sei $\pi \in S_n$ eine erweiterte Permutation. Besitzt $G'(\pi)$ einen orientierten Kreis, dann existiert ein 2-Move τ .

Beweis: Nach Voraussetzung gibt es also einen orientierten Kreis. Neben der am weitesten rechts beginnenden Desire-Edge muss eine weitere von rechts nach links durchlaufen werden. Dies ist in Abbildung 6.59 oben dargestellt, wobei wir dort unterscheiden, ob die zweite Desire-Edge vollständig innerhalb der ersten Desire-Edge (links) verläuft oder nicht (rechts). Da alle Reality-Edges von links nach rechts durchlaufen werden, ist die Form dieser zweiten Desire-Edge auch klar.

Verläuft die zweite Desire-Edge vollständig innerhalb der ersten, wie links in Abbildung 6.59 dargestellt, dann wenden wir die dort angegebene Transposition an. Andernfalls liegt die Situation vor, wie im rechten Teil der Abbildung 6.59 skizziert. Dann wenden wir die dort angegebene Transposition an. Man beachte, dass beide Transpositionen jeweils auf einem Kreis operieren, d.h. die Anzahl der Kreise in keinem Fall sinken kann. In jedem Fall wird aber nach Ausführung der jeweiligen

Abbildung 6.59: Skizze: Orientierter Kreis in $G'(\pi)$

Transposition die Anzahl der Kreise in $G'(\pi \circ \tau)$ um mindestens eins erhöht, da wir ja eine Adjazenz erzeugen. Nach Lemma 6.68 muss dann aber $G'(\pi \circ \tau)$ genau zwei Kreise mehr besitzen. Also ist τ ein 2-Move. ■

Lemma 6.77 Sei $\pi \in S_n$ eine erweiterte Permutation. Besitzt $G'(\pi)$ keinen orientierten Kreis, dann existiert kein 2-Move τ auf π .

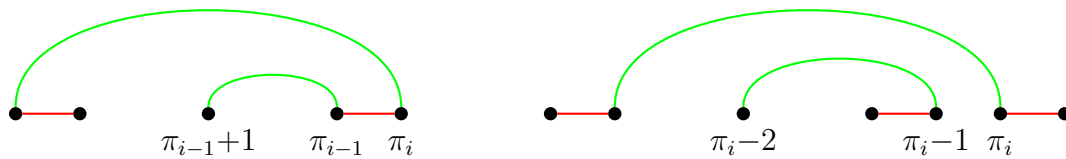
Den Beweis dieses Lemmas überlassen wir dem Leser als Übungsaufgabe.

Lemma 6.78 Sei $id \neq \pi \in S_n$ eine erweiterte Permutation. Dann besitzt π einen 2-Move oder es existieren zwei Transposition τ und τ' , wobei τ ein 0-Move für π und τ' ein 2-Move für $\pi \circ \tau$ ist.

Beweis: Enthält $G'(\pi)$ einen orientierten Kreis, dann folgt die Behauptung aus dem Lemma 6.76. Also seien alle Kreise in $G'(\pi)$ unorientiert. Wir erinnern noch einmal an die Beobachtung, dass unorientierte Kreise die Gestalt der linken Illustration in Abbildung 6.58 besitzen müssen.

Wir betrachten einen unorientierten Kreis C in $G'(\pi)$, der mindestens zwei Desire-Edges enthält (sonst wäre die Permutation bereits sortiert). Weiterhin muss es eine Desire-Edge zwischen $\max(C)$ und $\min(C)$ geben, da sonst C orientiert wäre. Weiter gilt $(\min(C), \max(C)) \neq (0, n+1)$.

Sei $k = \max(C)$ und sei also $(\pi_k - 1, \pi_k)$ die am weitesten rechts endende Desire-Edge in C (die dann auch in $\min(C)$ beginnen muss, da C unorientiert ist). Wir setzen zunächst $i = k$. Dann betrachten wir die von rechts her nächste darin beginnende bzw. endende Desire-Edge. Ist (π_{i-1}, π_i) eine Reality-Edge, dann hat diese die Form



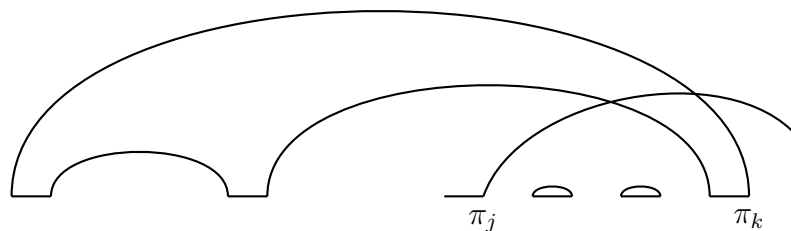
Abbildungung 6.60: Skizze: Mögliche nächstinnere Desire-Edges

$(\pi_j - 1, \pi_j) = (\pi_{i-1} - 1, \pi_{i-1})$ und sonst $(\pi_j + 1, \pi_j) = (\pi_i - 1, \pi_i - 2)$. Diese beiden Fälle sind in Abbildung 6.60 illustriert.

Schneidet $(\pi_j - 1, \pi_j)$ bzw. $(\pi_j + 1, \pi_j)$ keine der bislang betrachteten Desire-Edges, so setzen wir $i = j$ und führen dieses Verfahren fort. Andernfalls untersuchen wir die zu π_j inzidente Desire-Edge genauer. Irgendwann müssen wir eine Desire-Edge finden, die aus dem betrachteten Intervall ausbricht (sonst würde die betrachtete unorientierte Permutation nicht aus einer Permutation mit ausschließlich positiven Vorzeichen stammen).

Für die so gefundenen ausbrechende Desire-Edge unterscheiden wir zwei Fälle, je nachdem, wie die zu π_j inzidente Desire-Edge aussieht.

Fall 1: Wir nehmen zuerst an, dass die zu π_j inzidente Kante zu einer Position größer als k führt. Man beachte, dass die betrachte Desire-Edge nicht an den Positionen $i + 1$ mit k enden kann und dass diese Desire-Edge zu einem anderen Kreis gehören muss. Dies ist in Abbildung 6.61 illustriert.



Abbildungung 6.61: Skizze: Fall 1: Eine Kante bricht nach rechts aus

Diese Situation ist in Abbildung 6.62 etwas genauer dargestellt. Dort ist (f, i) die ausbrechende Kante und die rot-gepunkteten Linien entsprechend alternierenden Pfaden im Reality-Desire-Diagramm.

Wie man der Abbildung 6.62 entnimmt, gibt es eine Transposition τ , die ein 0-Move ist: Da diese Transposition auf zwei Kreisen operiert und mindestens zwei Kreise übrig lässt (unter anderem erzeugen wir ja die Adjazenz (e, j)), muss es sich um einen 0-Move oder 2-Move handeln (nach Lemma 6.77 ist letzteres aber nicht

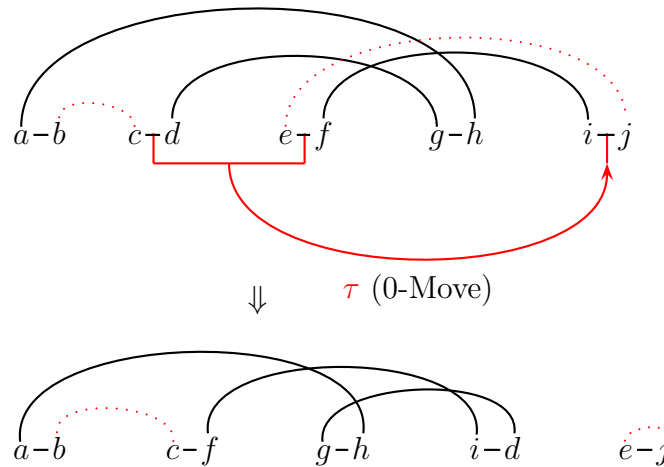


Abbildung 6.62: Skizze: Fall 1: 0-Move der einen orientierten Kreis generiert

möglich). Da nach der Ausführung von τ ein orientierter Kreis generiert wurde, gibt es nach Lemma 6.76 jetzt einen 2-Move.

Fall 2: Jetzt nehmen wir an, dass diese Kante nach links ausbricht und eine der vorher betrachteten Desire-Edges schneidet. Dies ist in Abbildung 6.63 illustriert. Hierbei ist jetzt jedoch nicht klar, ob der andere Endpunkt noch im aufspannenden Intervall des zugehörigen Kreises liegt oder nicht. Auch hier muss die schneidende Kante wieder zu einem andere Kreis gehören

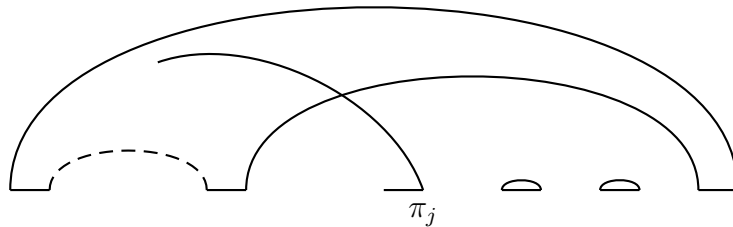


Abbildung 6.63: Skizze: Fall 2: Desire-Edge bricht nach links aus

Wir betrachten diese Situation jetzt in Abbildung 6.64 genauer. Hier ist die Desire-Edge (e, h) eine der zuletzt betrachteten Kanten, die von der aktuellen Kante (c, g) geschnitten wird. Man beachte, dass c auch links von a liegen kann.

Wie man der Abbildung 6.64 entnimmt, gibt es eine Transposition τ , die ein 0-Move ist: Da diese Transposition auf zwei Kreisen operiert und mindestens zwei Kreise übrig lässt (unter anderem erzeugen wir ja die Adjazenz (h, e)), muss es sich um einen 0-Move oder 2-Move handeln (nach Lemma 6.77 ist letzteres aber nicht möglich). Da nach der Ausführung von τ ein orientierter Kreis generiert wurde, gibt es nach Lemma 6.76 jetzt einen 2-Move. ■

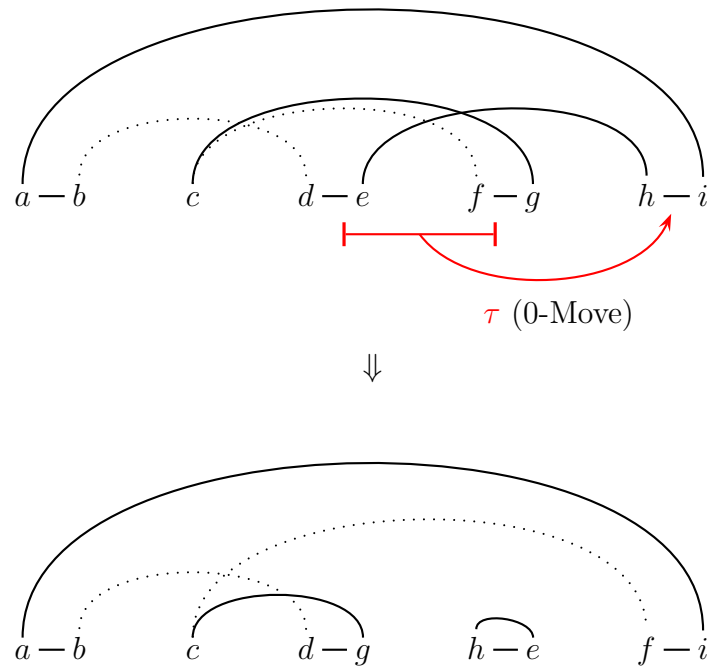


Abbildung 6.64: Skizze: Fall 2: 0-Move der einen orientierten Kreis generiert

Aus diesem Lemma folgt sofort, dass nach maximal 2 Transpositionen ein 2-Move erfolgt und dass ansonsten nur 0-Moves erfolgen können. Es sind also maximal $(n + 1) - c'(\pi)$ Transpositionen nötig, um π zu sortieren. Damit haben wir den folgenden Satz bewiesen.

Theorem 6.79 *Es existiert eine polynomielle 2-Approximation für Min-SBT.*

Zum Abschluss dieses Abschnitts wollen wir noch auf die folgenden Verbesserungen bei der Approximationsgüte hinweisen.

- Man kann zeigen, dass nach einem 0-Move mindestens zwei 2-Moves möglich sind, sofern $G'(\pi)$ lange orientierte Kreise mit mindestens 3 Reality-Edges enthält. Damit ist eine Approximation mit Güte 1,75 machbar.
- Eine Approximation mit Güte von 1,5 ist in Zeit $O(n^2)$ zu erstellen. Für Details verweisen wir auf die Originalarbeit von Hartman und Shamir.
- Es ist bislang noch ungeklärt, ob eine optimale Lösung -2 -Moves enthalten kann oder nicht.
- Bislang ist auch noch unbekannt, ob Min-SBT in $\mathcal{P}$ ist oder ob es $\mathcal{NP}$ -hart ist.

6.6 Sorting by Transpositions and Reversals (*)

In diesem Abschnitt wollen wir jetzt die Distanz von zwei Permutationen ermitteln, wenn sowohl Reversionen als auch Transpositionen erlaubt sind.

6.6.1 Problemstellung

In diesem Abschnitt erlauben wir neben den normalen Operationen Transposition und Reversionen teilweise auch noch so genannte *Transversals*. Hierbei handelt es sich um eine Transposition, wobei anschließend einer der beiden vertauschten Teile noch umgedreht wird. Weiter erlauben wir teilweise noch *doppelte Reversionen*, d.h. es werden gleichzeitig zwei benachbarte Teilstücke umgedreht ohne sie zu vertauschen. Auch diese Transversals bzw. doppelten Reversionen können wir in Form von Permutation angeben:

$$\begin{aligned} \text{Transversal}_1(i, j, k) &= \begin{pmatrix} 1 \cdots i-1 & i \cdots j & j+1 \cdots k & k+1 \cdots n \\ 1 \cdots i-1 & j+1 \cdots k & i \cdots j & k+1 \cdots n \end{pmatrix}, \\ \text{Transversal}_2(i, j, k) &= \begin{pmatrix} 1 \cdots i-1 & i \cdots j & j+1 \cdots k & k+1 \cdots n \\ 1 \cdots i-1 & k \cdots j+1 & i \cdots j & k+1 \cdots n \end{pmatrix}, \\ \text{DReversal}(i, j, k) &= \begin{pmatrix} 1 \cdots i-1 & i \cdots j & j+1 \cdots k & k+1 \cdots n \\ 1 \cdots i-1 & j \cdots i & k \cdots j+1 & k+1 \cdots n \end{pmatrix}. \end{aligned}$$

Dabei werden wir im Folgenden bei einer Reversion innerhalb eines Transversals oder einer doppelten Reversion auch immer die Vorzeichen ändern, da wir hier nur orientierte Permutationen untersuchen werden. In Abbildung 6.65 ist ein Beispiel für eine sortierende Folge von Operationen angegeben, die auch ein Transversal und eine doppelte Reversion verwendet.

$$\begin{array}{rcl} +0 & \boxed{+2 \quad -3 \quad +5 \quad +4 \quad -1} & +6 \\ & & R \\ +0 & +1 \quad \boxed{-4 \quad -5 \quad +3} \quad -2 & +6 \\ & & T \\ +0 & +1 \quad \boxed{+3 \quad -4 \quad -5} \quad \boxed{-2} & +6 \\ & & TR \\ +0 & +1 \quad +2 \quad +3 \quad \boxed{-4} \quad \boxed{-5} & +6 \\ & & DR \\ +0 & +1 \quad +2 \quad +3 \quad +4 \quad +5 & +6 \end{array}$$

Abbildung 6.65: Beispiel: Sortierung mit Transversals und doppelter Rotation

Somit können wir nun die zu untersuchenden Probleme genau definieren.

SORTING BY REVERSALS AND TRANSPOSITIONS (MIN-SBRT)

Eingabe: Eine Permutation $\pi \in S_n$.

Gesucht: Eine kürzeste Folge von Reversionen und Transpositionen $\rho_1, \dots, \rho_k$ mit $\pi \circ \rho_1 \circ \dots \circ \rho_k = \text{id}$.

SORTING BY TRANSVERSALS (MIN-SBTR)

Eingabe: Eine Permutation $\pi \in S_n$.

Gesucht: Eine kürzeste Folge von Reversionen, Transpositionen, Transversals und doppelten Reversionen $\rho_1, \dots, \rho_k$ mit $\pi \circ \rho_1 \circ \dots \circ \rho_k = \text{id}$.

Notation 6.80 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. d_{TR} bezeichnet die minimale Anzahl von Reversionen und Transpositionen, die nötig sind um $\bar{\pi}$ zu sortieren. d'_{TR} bezeichnet die Transversal-Distanz, d.h. die minimale Anzahl von Reversionen, Transpositionen, Transversals und doppelten Reversionen, die nötig sind um $\bar{\pi}$ zu sortieren

Weiterhin kürzen wir die Operationen Reversionen, Transpositionen, Transversals bzw. doppelte Reversionen oft mit T , R , TR bzw. DR ab. Man überlege sich ferner, warum man bei Transversals nicht die Reversion beider Teile nach einer Transposition zulässt. Wir wiederholen auch hier noch einmal kurz zur Erinnerung die Definition eines Reality-Desire-Diagrams.

Definition 6.81 Sei $\bar{\pi} \in \bar{S}_n$ eine erweiterte orientierte Permutation. Dann ist das (erweiterte) Reality-Desire-Diagramm (kurz RDD) $G(\bar{\pi})$ bzw. $G'(\bar{\pi})$ der (erweiterte) Breakpoint-Graph $G(\pi)$ bzw. $G'(\pi)$, wobei π die zu $\bar{\pi}$ gehörige erweiterte unorientierte Permutation ist, mit der Einschränkung, dass es für jedes $i \in [1 : n]$ keine Reality- bzw. Desire-Edges zwischen $(2i - 1)$ und $(2i)$ gibt.

6.6.2 Untere Schranken für die Transversal-Distanz

Nun wollen wir eine untere Schranke für das orientierte Sortieren mit Transversals herleiten.

Lemma 6.82 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $\bar{\rho}$ eine Operation vom Typ R , T , TR_1 , TR_2 oder DR , dann gilt:

$$|\bar{b}(\bar{\pi} \circ \bar{\rho}) - \bar{c}(\bar{\pi} \circ \bar{\rho}) - (\bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}))| \leq 2 \quad \text{bzw.} \quad |\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi})| \leq 2.$$

Beweis: Auch hier gilt immer noch $\bar{c}'(\bar{\pi}) = \bar{c}(\pi) + (n+1) - \bar{b}(\bar{\pi})$. Wir müssen also nur $|\bar{c}'(\bar{\pi} \circ \bar{\rho}) - \bar{c}'(\bar{\pi})| \leq 2$ zeigen. Es sind in jeder Operation jedoch maximal drei Kreise involviert, d.h. im schlimmsten Fall werden aus einem Kreis drei Kreise bzw. aus drei Kreisen ein Kreis. Damit folgt im Wesentlichen auch schon die Behauptung. ■

Theorem 6.83 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Dann gilt:

$$d'_{TR}(\bar{\pi}) \geq \frac{\bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi})}{2} = \frac{n+1 - \bar{c}'(\bar{\pi})}{2}.$$

Der Beweis erfolgt analog wie beim Sortieren mit Reversionen und sei dem Leser zur Übung überlassen.

Korollar 6.84 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Dann gilt:

$$d_{TR}(\bar{\pi}) \geq \frac{\bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi})}{2} = \frac{n+1 - \bar{c}'(\bar{\pi})}{2}.$$

Notation 6.85 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Dann bezeichnet $\hat{c}(\bar{\pi})$ die Anzahl der Kreise in $G'(\bar{\pi})$, die eine ungerade Anzahl von Reality-Edges (oder Desire-Edges) besitzen.

Im Folgenden sagen wir, dass ein Kreis gerade bzw. ungerade Länge hat, wenn er eine gerade bzw. ungerade Anzahl an Reality-Edges enthält.

Lemma 6.86 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $\bar{\rho}$ eine Operation vom Typ R, T, TR_1, TR_2, DR , dann gilt $|\hat{c}(\bar{\pi} \circ \bar{\rho}) - \hat{c}(\bar{\pi})| \leq 2$.

Beweis: Die Behauptung folgt im Wesentlichen aus folgender Beobachtung in Kombination mit dem Beweis von Lemma 6.82. Aus einem Kreis gerader Länge können maximal zwei Kreise ungerader Länge entstehen, da alle Kanten, die vorher in einem Kreis waren, auch nachher wieder in einem Kreis enthalten sein müssen. Andererseits kann aus drei Kreisen ungerader Länge kein Kreis gerader Länge entstehen. Da ein Transversal auf maximal drei Kreisen operiert, kann sich die Anzahl ungerader Kreise nur um höchstens zwei ändern. ■

Theorem 6.87 Für jede orientierte Permutation $\bar{\pi} \in \bar{S}_n$ gilt: $d'_{TR}(\bar{\pi}) \geq \frac{n+1-\hat{c}(\bar{\pi})}{2}$.

Korollar 6.88 Für jede orientierte Permutation $\bar{\pi} \in \bar{S}_n$ gilt: $d_{TR}(\bar{\pi}) \geq \frac{n+1-\hat{c}(\bar{\pi})}{2}$.

6.6.3 Eine 2-Approximation

In diesem Abschnitt wollen wir jetzt einen polynomiellen Approximationsalgorithmus mit einer Güte von 2 für Min-SBRT entwickeln.

Lemma 6.89 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Sei weiter C ein Kreis in $G'(\bar{\pi})$ mit einem Paar von sich schneidender Desire-Edges. Dann existiert eine Operation τ vom Typ T oder R , die die Länge von C um eins bzw. zwei verkürzt, eine bzw. zwei neue Adjazenzen erzeugt und die Länge der übrigen Kreise in $G'(\bar{\pi})$ unverändert lässt.*

Beweis: Wir werden dazu zeigen, dass wir jeweils aus einem Kreis mindestens eine neue Adjazenz erzeugen können. Die restlichen Kanten des Kreises müssen dann weiterhin einen Kreis bilden. Wir betrachten hierzu die folgende Fallunterscheidung. Für die linke der beiden sich schneidenden Desire-Edges betrachten wir die vier möglichen Kombinationen, wie die hierzu inzidenten Reality-Edges liegen können.

Fall 1: Wir nehmen zuerst an, dass die beiden inzidenten Reality-Edges nach außen führen. Siehe hierzu auch Abbildung 6.66.

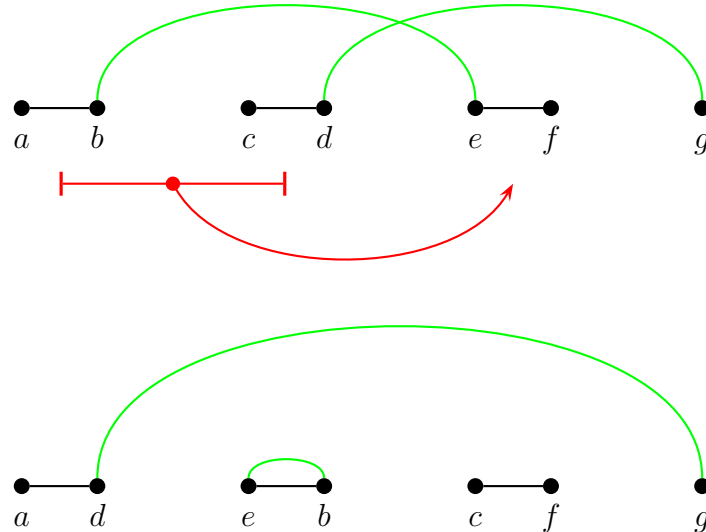


Abbildung 6.66: Skizze: Fall 1: Reality-Edges führen nach außen

Wie man in Abbildung 6.66 sieht, gibt es eine Transposition, die mindestens eine neue Adjazenz erzeugt. Hierbei ist die Ausrichtung der Reality-Edge (c, d) unerheblich (d.h., ob sie von d nach links oder rechts führt), die Transposition muss nur zwischen c und d operieren.

Fall 2: Wir nehmen jetzt an, dass die beiden inzidenten Reality-Edges nach links führen. Siehe hierzu auch Abbildung 6.67.

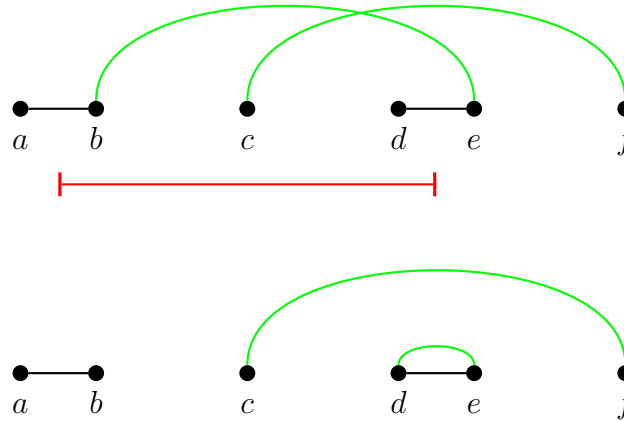


Abbildung 6.67: Skizze: Fall 2: Reality-Edges führen nach links

Wie man in Abbildung 6.67 sieht, gibt es eine Reversion, die mindestens eine neue Adjazenz erzeugt.

Fall 3: Wir nehmen jetzt an, dass die beiden inzidenten Reality-Edges nach rechts führen. Siehe hierzu auch Abbildung 6.68.

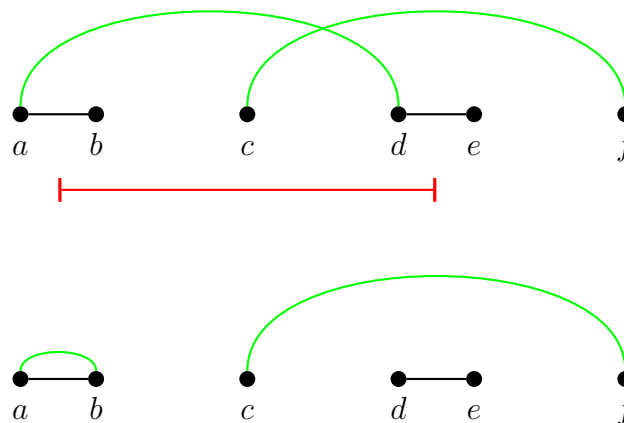


Abbildung 6.68: Skizze: Fall 3: Reality-Edges führen nach rechts

Wie man in Abbildung 6.68 sieht, gibt es eine Reversion, die mindestens eine neue Adjazenz erzeugt.

Fall 4: Wir nehmen zum Schluss an, dass die beiden inzidenten Reality-Edges nach innen führen. Siehe hierzu auch Abbildung 6.69.

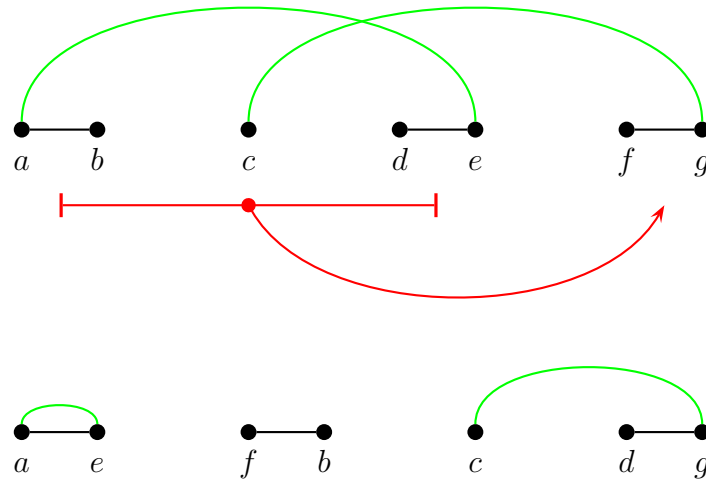


Abbildung 6.69: Skizze: Fall 4: Reality-Edges führen nach innen

Wie man in Abbildung 6.69 sieht, gibt es eine Transposition, die eine Adjazenz erzeugt. Hierbei ist die Ausrichtung der Reality-Edge (f, g) unerheblich (d.h., ob sie von g nach links oder rechts führt), die Transposition muss nur zwischen f und g operieren. ■

Lemma 6.90 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Sei weiter C ein Kreis in $G'(\bar{\pi})$ ohne sich schneidende Desire-Edges und sei $(x, x+1)$ eine Desire-Edge in C . Dann existiert eine Reality-Edge (y, y') eines anderen Kreises C' in $G'(\bar{\pi})$ zwischen $(x, x+1)$, die nicht in einer Adjazenz enthalten ist.

Beweis: Ein Kreis in $G'(\bar{\pi})$ mit sich nicht schneidenden Kanten muss eine wie in Abbildung 6.70 angegebene Gestalt besitzen.

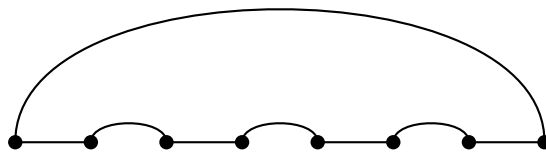


Abbildung 6.70: Skizze: Kreis ohne sich schneidende Desire-Edges

Für einen Widerspruchsbeweis nehmen an, dass in einer Lücke $(x, x+1)$ nur Reality-Edges liegen, die zu Adjazenzen gehören. Da die betrachtete Desire-Edge x mit $x+1$ verbindet, muss nach Konstruktion des erweiterten Reality-Desire-Graphen neben x der Wert $x-1$ stehen (und x muss gerade sein). Da $x-1$ zu einer Adjazenz gehört, ist der folgende Wert $x-2$. Darauf folgt aufgrund der Konstruktion des erweiterten

Reality-Desire-Diagrams der Wert $x-3$. Da nun auch $x-3$ zu einer Adjazenz gehört, ist der folgenden Wert $x-4$ usw. Dies ist in Abbildung 6.71 illustriert.

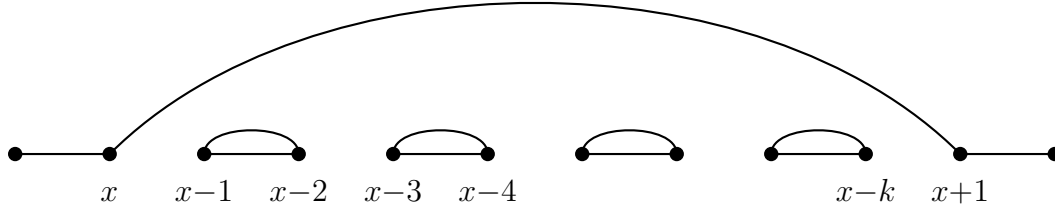


Abbildung 6.71: Skizze: Nur Adjazenzen in einer Desire-Lücke

Somit muss am Ende dieser Desire-Edge $x-k$ neben $x+1$ stehen und diese müssen aus einem vorzeichenbehafteten Element einer orientierten Permutation vermöge der Vorschrift $\{2i-1, 2i\}$ entstanden sein. Dies ist offensichtlich der gewünschte Widerspruch. ■

Lemma 6.91 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Gibt es in $G'(\bar{\pi})$ keine Kreise mit sich schneidenden Desire-Edges, dann existieren überlappende Kreise in $G'(\bar{\pi})$.

Beweis: Sei π die zu $\bar{\pi}$ gehörige unorientierte Permutation und sei e der linkeste Breakpoint in π im Kreis C . Betrachte die darauf folgende Reality-Edge e' , die nicht zu einer Adjazenz gehört. Schneidet eine der zugehörigen inzidenten Desire-Edges den Kreis C , ist die Behauptung gezeigt, da diese zu einem anderen Kreis gehören müssen. Dies ist im oberen Teil der Abbildung 6.72 illustriert.

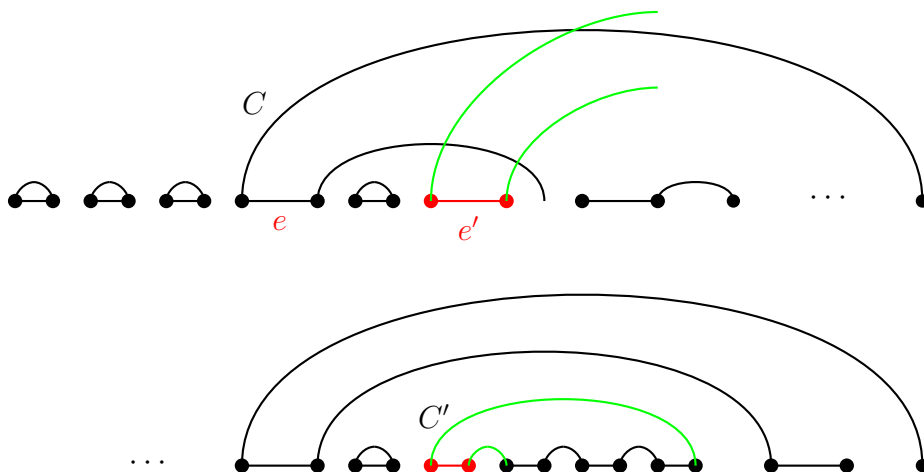


Abbildung 6.72: Skizze: Zum Beweis von Lemma 6.91

Schneiden die Desire-Edges des Kreises C' nicht den Kreis C , so betrachten wir jetzt den Kreis C' statt C . Dies ist im unteren Teil von Abbildung 6.72 dargestellt. Wir führen dieselbe Argumentation mit e' und C' fort. Entweder wir finden sich überlappende Kreise oder wir würden nach Lemma 6.90 unendlich viele Kreise inspizieren müssen. Dies ist offensichtlich nicht möglich, also finden wir ein Paar überlappender Kreise. ■

Lemma 6.92 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Seien C und C' zwei sich überlappende Kreise in $G'(\bar{\pi})$, die keine sich selbst schneidenden Desire-Edges enthalten. Dann existieren zwei Transpositionen τ und τ' , die die Länge von C und C' jeweils um eins reduzieren und zwei neue Adjazenzen einführen und darüber hinaus die Längen der übrigen Kreise unverändert lassen.*

Beweis: Zuerst einmal überlegt man sich, dass es bei zwei sich überlappenden Kreisen, zwei sich schneidende Desire-Edges geben muss, deren inzidente Reality-Edges nach außen weisen. Dazu benötigen wir wiederum die Voraussetzung, dass sich die Kreise nicht selbst schneiden und damit wieder eine feste Gestalt haben müssen, wie bereits in Abbildung 6.70 illustriert. Dies ist in Abbildung 6.73 im oberen Teil dargestellt, wobei die beiden sich schneidenden Desire-Edges grün gezeichnet sind.

Wendet man jetzt die beiden Transpositionen an, wie es in Abbildung 6.73 illustriert ist, so sieht man leicht, dass beide Kreise am Ende um jeweils eine Reality-Edge kürzer sind und dass zwei neue Adjazenzen erzeugt wurden. Die Längen der anderen Kreise blieben offensichtlich unverändert. ■

Theorem 6.93 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. Es existiert eine Folge von k Reversionen und Transpositionen, die $\bar{\pi}$ in Zeit $O(n^2)$ sortieren, wobei $k \leq 2 \cdot d_{TR}(\bar{\pi})$.*

Beweis: Wie aus den vorhergehenden Lemmata folgt, finden wir immer ein Folge Operationen, die im Schnitt die Länge eines Kreises von $G'(\bar{\pi})$ um eins verkürzt. Seien $\ell_1, \dots, \ell_{\bar{c}'(\bar{\pi})}$ die Längen der $\bar{c}'(\bar{\pi})$ Kreise in $G'(\bar{\pi})$ (wobei hier mit Länge wieder die Anzahl von Reality-Edges gemeint ist). Da es im erweiterten Reality-Desire-Diagramm zu $\bar{\pi}$ genau $n + 1$ Reality-Edges gibt, ist die Anzahl ausgeführter Operationen durch

$$\sum_{i=1}^{\bar{c}'(\bar{\pi})} (\ell_i - 1) \leq (n + 1) - \bar{c}'(\bar{\pi})$$

nach oben beschränkt. Mit dem Korollar 6.84 folgt die Behauptung über die Approximationsgüte.

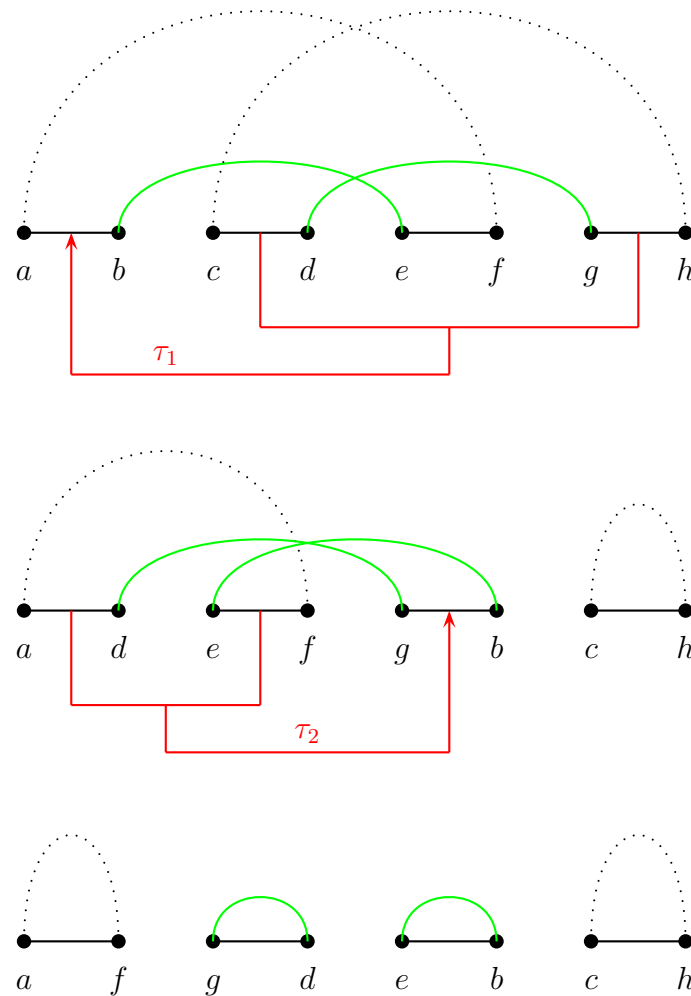


Abbildung 6.73: Skizze: Zwei Transpositionen, die die Längen von zwei Kreisen in $G'(\bar{\pi})$ um jeweils eins verkürzen

Dass dieser Algorithmus in quadratischer Zeit implementiert werden kann, bleibt dem Leser als Übungsaufgabe überlassen. ■

Ein Verbesserung des vorhergehenden Ergebnisses halten wir im folgenden Satz ohne Beweis fest.

Theorem 6.94 *Jede Permutation $\bar{\pi} \in \bar{S}_n$ lässt sich in Zeit $O(n^2)$ mit Operationen vom Typ R , T , TR und DR mit einer Approximationsgüte von 1,75 sortieren.*

Die entsprechenden Approximationsgüten gelten trivialerweise auch noch, wenn man Transversals oder doppelte Reversionen zulässt.

Selbst im orientierten Fall ist bislang nicht bekannt, ob die Fragstellung $\mathcal{NP}$ -hart ist. Für den unorientierten Fall sind bislang noch keine (außer den trivialen) polynomiellen Approximationen bekannt.

6.7 Sorting by weighted Transversals (*)

In diesem Abschnitt wollen wir nun untersuchen, ob man eine optimale Lösung für das Sortieren mit Reversionen und Transpositionen effizient konstruieren kann, wenn man beide Operationen unterschiedlich gewichtet. Bei einer Reversion muss die DNA nur an zwei Stellen geschnitten werden, bei einer Transposition an drei Stellen. Wir werden nun versuchen die wahrscheinlicheren Reversionen billiger als die aufwendigeren Transpositionen zu machen.

6.7.1 Gewichtete Transversal-Distanz

Um die neue gewichtete Transversal-Distanz zu definieren benötigen wir erst noch eine Notation.

Notation 6.95 Sei s eine Folge von Reversionen und Transpositionen. Mit $rev(s)$ bezeichnet man die Anzahl von Reversionen in s und mit $trp(s)$ die Anzahl von Transpositionen in s .

Nun zur Definition der gewichteten Transversal-Distanz.

Definition 6.96 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation. $\mathcal{R}(\bar{\pi})$ bezeichnet alle Folgen von Reversionen, die $\bar{\pi}$ sortieren und $\mathcal{TR}(\bar{\pi})$ alle Folgen von Reversionen und Transpositionen, die $\bar{\pi}$ sortieren.

$$\begin{aligned}\bar{d}(\bar{\pi}) = \bar{d}_R(\bar{\pi}) &:= \min \{ rev(s) : s \in \mathcal{R}(\bar{\pi}) \} \\ \bar{d}_{TR}(\bar{\pi}), &:= \min \{ rev(s) + trp(s) : s \in \mathcal{TR}(\bar{\pi}) \}, \\ \tilde{d}(\bar{\pi}) = \tilde{d}_{TR}(\bar{\pi}) &:= \min \{ rev(s) + 2trp(s) : s \in \mathcal{TR}(\bar{\pi}) \}\end{aligned}$$

$\bar{d} = \bar{d}_R$ heißt Reversions-Distanz, $\bar{d}_{TR}$ heißt einfache Transversal-Distanz und $\tilde{d}$ heißt gewichtete Transversal-Distanz.

Wir erinnern uns, dass $\bar{d}_R(\bar{\pi}) = \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \bar{h}(\bar{\pi}) + \bar{f}(\bar{\pi})$ gilt.

Zunächst einmal halten wir fest, dass es für die gewichtete Transversal-Distanz unerheblich ist, ob wir Transversals bzw. doppelte Reversionen zulassen oder nicht, da wir diese jeweils durch zwei Reversionen mit gleichen Kosten ersetzen können (siehe Abbildung 6.74 für die Ersetzung eines Transversals).

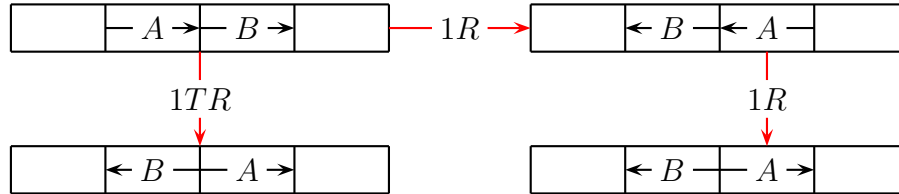


Abbildung 6.74: Skizze: Ersetzung von Transversals

Damit können wir nun auch die Problemstellung angeben.

SORTING BY WEIGHTED TRANSVERSALS (MIN-SWT)

Eingabe: Eine orientierte Permutation $\bar{\pi} \in \bar{S}_n$.

Gesucht: Eine Folge $s = (\bar{\rho}_1, \dots, \bar{\rho}_k)$ bestehend aus Reversionen und Transpositionen mit $\bar{\pi} \circ \bar{\rho}_1 \circ \dots \circ \bar{\rho}_k = \text{id}$, so dass $\tilde{d}(\bar{\pi}) = \text{rev}(s) + 2\text{trp}(s)$.

6.7.2 Starke Orientierung und starke Hurdles

Für das Sortieren mit gewichteten Transversals benötigen wir noch die Begriffe der starken Orientierung und starken Hurdles.

Definition 6.97 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Für eine Zusammenhangskomponente C bezeichnet $\tilde{d}(C)$ die Anzahl der Operationen, die benötigt werden, um alle Elemente aus C zu sortieren (interpretiert als eine neue erweiterte orientierte Permutation).

- C heißt stark unorientiert, wenn $\tilde{d}(C) > \bar{b}(C) - \bar{c}(C)$.
- Trennt eine stark unorientierte Komponente C keine anderen stark unorientierten Komponenten, dann heißt C eine starke Hurdle.
- Eine starke Hurdle heißt starke Super-Hurdle, wenn deren Löschung eine andere starke unorientierte Komponente in eine starke Hurdle umwandelt.
- Besteht eine Permutation aus einer ungeraden Anzahl von starken Hurdles, die alle starke Super-Hurdles sind, dann ist $\bar{\pi}$ ein starkes Fortress.

Darauf aufbauend ergänzen wir nun noch die entsprechenden Notationen für starke Hurdles und starke Fortresses.

Notation 6.98 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. $\tilde{h}(\bar{\pi})$ bezeichnet die Anzahl starker Hurdles in $\bar{\pi}$. Es gilt $\tilde{f}(\bar{\pi}) = 1$, wenn $\bar{\pi}$ ein starkes Fortress ist, und $\tilde{f}(\bar{\pi}) = 0$ sonst.

Die Begriffe orientiert, unorientiert, Hurdle und Fortress beziehen sich im Folgenden immer auf das Sortieren mit Reversionen. Im Zusammenhang mit gewichteten Transversals reden wir immer von stark orientiert, stark unorientiert, starken Hurdles und starken Fortresses.

Lemma 6.99 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Jede stark unorientierte Komponente in $G'(\bar{\pi})$ ist auch unorientiert.

Beweis: Für eine orientierte Komponente C in $G'(\bar{\pi})$ gilt: $\bar{d}_R(C) = \bar{b}(C) - \bar{c}(C)$. Angenommen, eine stark unorientierte Komponente C sei orientiert. Dann gilt

$$\bar{d}_R(C) \geq \tilde{d}_{TR}(C) > \bar{b}(C) - \bar{c}(C).$$

Damit erhalten wir den gewünschten Widerspruch. ■

Es kann jedoch unorientierte Komponenten geben, die stark orientierte Komponenten sind. Man bedenke insbesondere an unorientierte Komponenten, bei denen man mit einer Transposition im erweiterten Reality-Desire-Diagramm die Anzahl der Kreise um zwei erhöhen kann.

6.7.3 Eine untere Schranke für die Transversal-Distanz

In diesem Abschnitt wollen wir eine zur Reversal-Distanz analoge Formel für die gewichtete Transversal-Distanz angeben.

Theorem 6.100 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Es gilt:

$$\begin{aligned} \tilde{d}(\bar{\pi}) &= \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi}), \\ \tilde{d}(\bar{\pi}) &= (n+1) - \bar{c}'(\bar{\pi}) + \tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi}). \end{aligned}$$

Wir geben hier nur die Idee des Beweises an und beginnen mit einer groben Interpretation der Formel. Wie im Fall orientierter Reversionen soll also gelten

$$\tilde{d}(\bar{\pi}) = (n+1) - \bar{c}'(\bar{\pi}) + \tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi}) = \bar{b}(\bar{\pi}) - \bar{c}'(\bar{\pi}) + \tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi}).$$

Beachte, dass eine Reversion den Wert $(n+1) - \bar{c}'(\bar{\pi})$ im besten Falle um eins und eine Transposition um zwei erniedrigen kann. Somit ist nach Definition der gewichteten Transversal-Distanz $(n+1) - \bar{c}'(\bar{\pi})$ schon eine triviale untere Schranke. Der Term $\tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi})$ gibt hierbei auch die Anzahl zusätzlicher Operationen an, um aus starken Hurdles stark orientierte Komponenten zu machen.

Beweis: Zunächst einmal halten wir fest, dass wir jede orientierte Komponente C wie beim Sortieren mit Reversionen mit nur $\bar{b}(C) - \bar{c}(C)$ Reversionen sortieren können. Weiter können unorientierte Komponenten, die nicht stark unorientiert sind, nach Definition mit Hilfe von Transpositionen mit $(\bar{b}(C) - \bar{c}(C))$ Operationen sortiert werden. Wir müssen uns also nur noch um stark unorientierte Komponenten kümmern, die nach Lemma 6.99 auch unorientiert sind.

Wir bemerken, dass eine Reversion angewendet auf $\bar{\pi}$ den Wert $\bar{c}'(\bar{\pi})$ um maximal eins verändern kann. Eine Transposition angewendet auf $\bar{\pi}$ kann $\bar{c}'(\bar{\pi})$ um die Werte -2 , 0 oder $+2$ verändern. Man beachte, dass wir nur noch unorientierte Komponenten (und somit auch Permutationen) betrachten, d.h. wir können annehmen, dass alle Vorzeichen positiv sind. Somit folgt diese Tatsache aus Lemma 6.68.

Es bleibt also die Frage, inwieweit Transpositionen beim Aufbrechen starker Hurdles helfen können. Wir unterscheiden jetzt drei Fälle, je nachdem, wie eine Transposition auf Kreisen einer Komponente operiert.

Fall 1: Eine Transposition bricht einen Kreis in 3 Kreise auf. War der Kreis Teil einer starken Hurdle C , dann muss sich mindestens einer der neuen Kreise in einer starken Hurdle befinden, da sonst nach Definition die ganze starke Hurdle mit $\bar{b}(C) - \bar{c}(C)$ Operationen aufgelöst werden kann. Somit kann sich der Parameter $\tilde{h}(\bar{\pi})$ nicht verringern.

Fall 2: Eine Transposition auf 2 Kreisen liefert 2 neue Kreise. Um die gewichtete Transversal-Distanz um 3 erniedrigen zu können, müssten 3 Hurdles eliminiert werden. Dies ist unmöglich, da nur 2 Kreise verändert werden.

Fall 3: Wir betrachten jetzt eine Transpositionen, die auf 3 Kreisen operiert. Dabei entsteht ein einziger Kreis ergibt. Da $\bar{c}'(\bar{\pi})$ um 2 wächst, müssten 5 Hurdles eliminiert werden, was offensichtlich nicht möglich ist.

Eine Transposition kann also den Parameter $\tilde{h}(\bar{\pi})$ nicht noch zusätzlich zum Parameter $(n+1) - \bar{c}'(\bar{\pi})$ erniedrigen. Wir müssen uns nur noch überlegen, ob eine Transposition beim Eliminieren eines Fortresses auch noch den Parameter $\tilde{h}(\bar{\pi})$ erniedrigen

kann. Nach unserer Analyse muss aber eine Transposition, die $(n+1) - \tilde{c}'(\bar{\pi}) + \tilde{h}(\bar{\pi})$ um zwei erniedrigt, auch $\tilde{h}(\bar{\pi})$ um eine gerade Zahl verändern und dabei werden keine außer den entfernten starken Hurdles verändert. Also kann keine Transposition den Parameter $(n+1) - \tilde{c}'(\bar{\pi}) + \tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi})$ um drei erniedrigen. Somit gilt die untere Schranke

$$\tilde{d}(\bar{\pi}) \geq (n+1) - \tilde{c}'(\bar{\pi}) + \tilde{h}(\bar{\pi}) + \tilde{f}(\bar{\pi}),$$

wie in der Behauptung gefordert. ■

Man beachte, dass hier im Gegensatz zum Sortieren mit Reversionen die Begriffe starke Hurdle und starkes Fortress über die Eigenschaft definiert wurde, ob man zusätzliche Operationen über den Parameter $\bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi})$ hinaus zur Eliminierung der entsprechenden Komponenten benötigt. Es ist also momentan algorithmisch überhaupt nicht klar, wie man entscheidet, ob eine Komponente stark orientiert bzw. stark unorientiert ist oder ob es eine starke Hurdle ist.

6.7.4 Eine Approximation für Min-SWT

Nun versuchen wir einen Algorithmus zu entwickeln, der mit möglichst wenig Operationen auskommt, d.h. der unteren Schranke aus dem vorherigen Abschnitt sehr nahe kommt. Zuerst geben wir eine algorithmisch verwendbare Bedingung an, wann eine Komponente stark unorientiert ist.

Lemma 6.101 *Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Enthält eine unorientierte Komponente von $G'(\bar{\pi})$ einen Kreis gerader Länge, dann ist sie stark unorientiert.*

Auch hier verstehen wir bei der Länge eines Kreises wieder die Anzahl seiner Reality-Edges.

Beweis: Da die Komponente unorientiert ist, kann eine Reversion die Anzahl der Kreise im erweiterten Reality-Desire-Diagramm nicht erhöhen. Weiter wissen wir, dass eine Transposition die unorientierten Kreise unorientiert lässt und die Anzahl der unorientierten Kreise nur um -2 , 0 , oder $+2$ ändern kann. Somit ist eine Anwendung einer Transposition, die die Anzahl der Kreise um zwei erhöht, die einzige sinnvolle Wahl. Damit bleibt die Parität der Anzahl der Kreise unverändert, da jede Anwendung einer Transposition auf einen unorientierten Kreis gerader Länge einen Kreis gerader Länge sowie zwei Kreisen mit entweder gerader oder ungerader Länge erzeugt. Ein Kreis gerader Länge kann also nie mit Transpositionen endgültig eliminiert werden.

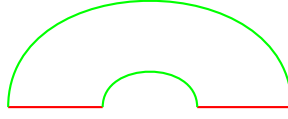


Abbildung 6.75: Skizze: ein unorientierter Kreis der Länge 2

Es bleibt also ein Kreis der Länge 2 übrig. Dieser muss eine Form, wie in Abbildung 6.75 illustriert, besitzen. Es existiert auch hier keine Transposition mehr, die $\bar{c}'(C)$ erhöhen kann. Also ist C eine stark unorientierte Komponente. ■

Es bleibt die Frage, wie man mit unorientierten Komponenten umgeht, die nur aus Kreisen ungerader Länge bestehen. Hier ist a priori nicht klar, ob sie stark orientiert oder stark unorientiert sind. Man kann leicht überprüfen, dass es nur eine unorientierte Komponente mit genau drei Reality-Edges gibt. Da sich diese jedoch mit einer Transposition auflösen lässt, ist sie stark orientiert. Diese Komponente ist in Abbildung 6.76 ganz links dargestellt.

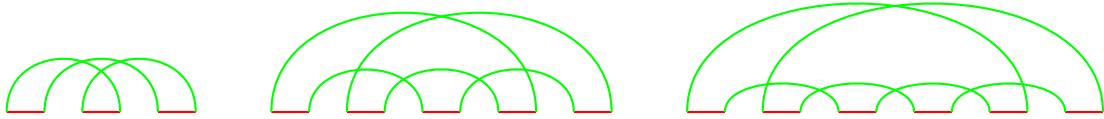


Abbildung 6.76: Skizze: Kleine unorientierte Komponenten

In Abbildung 6.76 sind in der Mitte und rechts die einzigen stark unorientierten Komponenten mit höchstens 6 Reality-Edges dargestellt, die nur aus Kreisen ungerader Länge bestehen. Mit dieser vollständigen Aufzählung lassen sich also kleine Komponenten leicht daraufhin testen, ob sie stark unorientiert sind oder nicht.

Definition 6.102 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Das Gewicht einer Komponente C von $G'(\bar{\pi})$ ist durch $\gamma(C) := \bar{b}(C) - \bar{c}(C)$ definiert.

In Abbildung 6.76 sind in der Mitte und rechts alle stark unorientierten Komponenten mit einem Gewicht von höchstens 4 dargestellt.

Lemma 6.103 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. Das Gewicht aller stark unorientierten Zusammenhangskomponenten in $G'(\bar{\pi})$, die nur aus Kreisen ungerader Länge bestehen, ist gerade.

Den Beweis diese Lemmas überlassen wir dem Leser zur Übung. Für das Folgende ist nur wichtig, dass es keine stark orientierten Komponenten mit Gewicht 5 gibt, die nur aus Kreisen ungerader Länge bestehen.

Der Algorithmus geht also wie folgt vor. Orientierte Komponenten werden mit Hilfe von Reversionen wie im Falle des Sortierens mit Reversionen behandelt. Unorientierte Komponenten, die einen Kreis gerader Länge enthalten, werden als stark unorientierte Komponente klassifiziert. Unorientierte Komponenten, die nur aus ungeraden Kreisen bestehen, werden gegen die bekannten stark unorientierten Komponenten kleinen Gewichtes (im einfachsten Falle kleiner als 6) verglichen und als stark orientiert bzw. stark unorientiert klassifiziert. Die stark orientierten werden dann mit Hilfe von Transpositionen und Reversionen aufgelöst (deren prinzipielle Lösung ebenfalls vorab gespeichert wurde). Alle übrigen unorientierten Komponenten, die nur aus ungeraden Kreisen bestehen, werden aus Unwissenheit der Einfachheit halber als stark unorientiert klassifiziert. Stark unorientierte Komponenten werden mit Hilfe von Reversionen wie im Falle von Sortieren mit Reversionen aufgelöst. Hierbei werden durch die möglichen Klassifizierungsfehler eventuell überflüssige Reversionen ausgeführt.

Man beachte, dass der Algorithmus hauptsächlich Reversionen verwendet und nur bei stark orientierten Komponenten, die unorientiert sind und ein kleines Gewicht besitzen, auf Transpositionen zurückgreift.

6.7.5 Approximationsgüte

Im letzten Teil kommen wir zur Bestimmung der Approximationsgüte des im eben vorgestellten polynomiellen Approximationsalgorithmus. Hierfür benötigen wir erst noch ein paar Notationen.

Notation 6.104 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. $\tilde{h}_u(\bar{\pi})$, $\bar{c}_u(\bar{\pi})$ bzw. $\bar{b}_u(\bar{\pi})$ bezeichnet die Anzahl der unorientierten Komponenten, Kreise bzw. Breakpoints in unorientierten Komponenten, die nur aus Kreisen ungerader Länge bestehen und mindestens ein Gewicht von 6 besitzen. $\bar{c}_o(\bar{\pi})$ bzw. $\bar{b}_o(\bar{\pi})$ bezeichnet die Anzahl der Kreise bzw. Breakpoints in den übrigen Komponenten.

Notation 6.105 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweiterte Reality-Desire-Diagramm. $\tilde{h}_{\min}(\bar{\pi})$ bzw. $\tilde{h}_{\max}(\bar{\pi})$ bezeichnet die Anzahl der starken Hurdles mit Gewicht mindestens 6, wenn keine bzw. alle unorientierten Komponenten, die nur aus Kreisen ungerader Länge bestehen, starke Hurdle sind.

Beobachtung 6.106 Sei $\bar{\pi} \in \bar{S}_n$ eine orientierte Permutation und sei $G'(\bar{\pi})$ das zugehörige erweitert Reality-Desire-Diagramm, dann gilt

$$\tilde{h}_{\min}(\bar{\pi}) \leq \tilde{h}(\bar{\pi}) \leq \tilde{h}_{\max}(\bar{\pi})$$

sowie

$$\begin{aligned} \tilde{h}_{\max}(\bar{\pi}) &= \tilde{h}_{\min}(\bar{\pi}) + \tilde{h}_u(\bar{\pi}), \\ \tilde{h}_u(\bar{\pi}) &\leq \bar{c}_u(\bar{\pi}), \\ \tilde{h}_u(\bar{\pi}) &\leq \frac{\bar{b}_u(\bar{\pi}) - \bar{c}_u(\bar{\pi})}{6}. \end{aligned}$$

Die Abschätzung für $\tilde{h}_u(\bar{\pi})$ folgt aus der Tatsache, dass hier nur unorientierte Komponenten mit Gewicht mindestens 6 gezählt werden. Im Regelfall ist diese Abschätzung nicht sehr genau, aber für unsere Zwecke ausreichend.

Die konkrete Zahl 6 in den vorhergehenden Notationen und Beobachtung stammt daher, dass wir für unorientierte Komponenten mit einem Gewicht kleiner 6, die nur aus ungeraden Kreisen bestehen, wissen, welche davon stark orientiert sind und welche nicht. Wenn wir das auch für größere Gewichte wissen, kann diese Konstante entsprechend erhöht werden.

Sei für die folgenden Betrachtungen $\bar{\pi} \in \bar{S}_n$ kein starkes Fortress. Dann gilt für die gewichtete Anzahl ausgeführter Operationen $A(\bar{\pi})$ nach dem im vorigen Abschnitt angegebenen Algorithmus:

$$\begin{aligned} A(\bar{\pi}) &= \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \tilde{h}(\bar{\pi}) \\ &\leq \bar{b}_o(\bar{\pi}) + \bar{b}_u(\bar{\pi}) - \bar{c}_o(\bar{\pi}) - \bar{c}_u(\bar{\pi}) + \tilde{h}_{\max}(\bar{\pi}) \\ &= \bar{b}_o(\bar{\pi}) + \bar{b}_u(\bar{\pi}) - \bar{c}_o(\bar{\pi}) - \bar{c}_u(\bar{\pi}) + \tilde{h}_{\min}(\bar{\pi}) + \tilde{h}_u(\bar{\pi}) \\ &\leq \bar{b}_o(\bar{\pi}) - \bar{c}_o(\bar{\pi}) + \bar{b}_u(\bar{\pi}) - \bar{c}_u(\bar{\pi}) + \tilde{h}_{\min}(\bar{\pi}) + \frac{\bar{b}_u(\bar{\pi}) - \bar{c}_u(\bar{\pi})}{6}. \end{aligned}$$

Weiter gilt offensichtlich mit Theorem 6.100:

$$\begin{aligned} \tilde{d}(\bar{\pi}) &\geq \bar{b}(\bar{\pi}) - \bar{c}(\bar{\pi}) + \tilde{h}(\bar{\pi}) \\ &\geq \bar{b}_o(\bar{\pi}) - \bar{c}_o(\bar{\pi}) + \bar{b}_u(\bar{\pi}) - \bar{c}_u(\bar{\pi}) + \tilde{h}_{\min}(\bar{\pi}). \end{aligned}$$

Somit gilt für die Approximationsgüte:

$$\frac{A(\bar{\pi})}{\tilde{d}(\bar{\pi})} \leq \frac{\bar{b}_o(\bar{\pi}) - \bar{c}_o(\bar{\pi}) + \frac{7}{6}(\bar{b}_u(\bar{\pi}) - \bar{c}_u(\bar{\pi})) + \tilde{h}_{\min}(\bar{\pi})}{\bar{b}_o(\bar{\pi}) - \bar{c}_o(\bar{\pi}) + \bar{b}_u(\bar{\pi}) - \bar{c}_u(\bar{\pi}) + \tilde{h}_{\min}(\bar{\pi})} \leq \frac{7}{6}.$$

Ein analoger Beweis für diese Approximationsrate ist auch für Permutationen möglich, die starke Fortresses sind. Halten wir das Ergebnis im folgenden Satz fest.

Theorem 6.107 *Es gibt es eine $7/6$ -Approximation, die in polynomieller Zeit eine orientierte Permutation mittels gewichteter Reversionen und Transpositionen sortiert.*

Man beachte, dass die Approximationsgüte im Wesentlichen aus der nicht sehr genauen Abschätzung für $\hat{h}_u(\bar{\pi})$ aus Beobachtung 6.106 stammt. Meist ist die gefundene Lösung also viel genauer als die Approximationsgüte vorhersagt.

Wenn für alle unorientierten Komponenten mit einem Gewicht kleiner als k , die nur aus Kreisen ungerader Länge bestehen, vorab nach der starken Orientiertheit (samt einer Lösung) klassifiziert wird, dann kann man auch den folgenden Satz zeigen, wobei dort $\varepsilon = 1/k$ ist.

Theorem 6.108 *Für jedes $\varepsilon > 0$ gibt es eine $(1 + \varepsilon)$ -Approximation, die in polynomieller Zeit eine orientierte Permutation mittels gewichteter Reversionen und Transpositionen sortiert.*

Wie schon bei Angabe des Algorithmus angegeben, verwendet dieser hauptsächlich Reversionen und nur in einigen Ausnahmefällen werden Transpositionen eingesetzt. Das legt die Vermutung nahe, dass die Gewichtung von Reversion zu Transposition für die Reversion zu günstig ausfällt. Insbesondere basiert das ganze Verfahren auch auf dem Trick, dass die Kosten einer Operation gerade genau groß wie die Anzahl der hinzugefügten Kreise im günstigsten Falle ist. Vermutlich wäre eine andere Gewichtung, wie etwa 2 pro Reversion und 3 pro Transposition günstiger, da sie den echten Aufwand, die Anzahl Stellen, an denen die DNA aufgeschnitten wird, besser repräsentiert.

6.8 Weitere Modelle (*)

Zum Abschluss dieses Kapitels wollen wir noch kurz auf einige Erweiterungen der Modelle und auf bereits dort erzielte Ergebnisse eingehen.

6.8.1 Gewichtung der Operationen durch ihre Länge

Unter der Annahme, dass kurze Operationen (wie Reversionen oder Transpositionen) wahrscheinlicher sind als lange, kann man das Modell dahingehend erweitern, die Kosten pro Operation von der Länge abhängig zu machen.

Eine erste, mehr mathematische Methode, die Länge ins Spiel zu bringen, ist, nur Reversionen fester Länge k zuzulassen. Für $k = 3$ gibt es dann bereits Permutationen, die mit solchen Reversionen überhaupt nicht mehr sortiert werden können. Erste Ergebnisse hierzu findet man in der Arbeit von Chen und Skiena.

Heath und Vergara haben eine ähnliche Variante für Transpositionen untersucht, wobei die Anzahl der transponierten Elemente beschränkt ist. Ist diese Beschränkung nicht konstant, sondern von der Länge der Permutation abhängig, hat dieses Problem bereits dieselbe (unbekannte) Komplexität wie das normale Sortieren mit Transpositionen. Für Transposition, die maximal drei Elemente transponieren, wird hingegen die Existenz eines polynomiellen Algorithmus vermutet.

Ein erster Ansatz für allgemeinere Gewichtsfunktionen, die wirklich von der Länge der Operationen abhängen, stammt von Pinter und Skiena. Für das Sortieren unorientierter Permutationen mittels Reversionen konnten für sehr spezielle Gewichtsfunktionen erste Approximationsalgorithmen für die Reversions-Distanz ermittelt werden. Bender et al. haben eine Klasse von Gewichtsfunktion der Gestalt $f_\alpha(\ell) = \ell^\alpha$ mit $\alpha > 0$ für unorientierte Reversionen genauer untersucht und untere bzw. obere Schranken sowie auch einige Approximationen sowohl für den Durchmesser (Maximum der minimalen Distanz über alle Paare von Permutationen) als auch für den minimalen Abstand zweier gegebener Permutationen ermittelt.

Auch die Erweiterung auf so genannte Block-Interchange Operationen wurden inzwischen genauer untersucht. Hierbei werden im Gegensatz zu Transpositionen zwei nicht benachbarte Blöcke vertauscht. Für den Fall unorientierter Permutationen konnte Christie einen quadratischen Algorithmus angeben.

6.8.2 Duplikationen

In Eukaryonten tauchen häufig Gene nicht nur einmal, sondern mehrfach im Genom auf. Von daher sind auch Varianten von Interesse, die solche Duplikationen berücksichtigen. Leider lassen sich dann die Genome nicht mehr als Permutationen modellieren. Da für unorientierte Permutationen Min-SBR bereits ein Spezialfall ist, bleibt dieses Problem $\mathcal{NP}$ -hart. Für eine feste Anzahl von Genen (unabhängig von der Länge des Genoms) konnten Christie und Irving zeigen, dass das Problem auch hier $\mathcal{NP}$ -hart bleibt. Sogar für nur zwei Gene (mit entsprechender Vielfachheit) bleibt die Bestimmung der Reversions-Distanz (im orientierten Fall) $\mathcal{NP}$ -hart. Für die Transpositions-Distanz ist dies noch offen. Für orientierte Permutationen gibt es bislang noch keine Ergebnisse.

6.8.3 Multi-chromosomale Genome

Genome höherer Lebewesen bestehen ja nicht nur aus einem, sondern in der Regel aus mehreren Chromosomen. Im Falle orientierter Permutationen (aufgeteilt auf mehrere Chromosomen) konnte Hannenhalli einen ersten polynomiellen Algorithmus zur Bestimmung der so genannten Translokations-Distanz angeben. Diesen Algorithmus konnten Hannenhalli und Pevzner noch um die Berücksichtigung von Reversionen erweitern. Dieser kann die exakte Distanz bis auf 1 ermitteln, was für praktische Zwecke jedoch völlig ausreichend ist.

Berücksichtigt man bei mehreren Chromosomen nur die Operationen Fusion, Fission sowie Transposition und belastet die Transpositionen mit doppelten Kosten im Vergleich zu Fusions und Fissions, so lässt sich die Distanz leicht berechnen, wie Meidanis und Dias zeigen konnten.

Unter der Prämisse, dass die Duplikation nur auf eine Duplikation aller Chromosomen eines urzeitlichen Genoms zurückzuführen sind, konnten El-Mabrouk und Sankoff eine effiziente Lösung vorstellen.

Betrachtet man bei multi-chromosomalen Genomen nur die Zuordnung der Gene auf die Chromosomen und vergisst deren Ordnung innerhalb der Chromosome, so kommt man zu den so genannten *Syntenry-Problemen*. Dabei werden in der Regel nur die Operationen Translokation, Fission und Fusion betrachtet. DasGupta et al. konnten zeigen, dass auch dieses Problem bereits $\mathcal{NP}$ -hart ist und haben eine 2-Approximation entwickelt. Kleinberg und Liben-Nowell konnten zeigen, dass die syntenische Distanz von zwei Genomen mit n Chromosomen höchstens $2n - 4$ ist.

6.8.4 Multiple Genome Rearrangements

Man kann auch versuchen, aus den Rearrangements der Genome mehrerer Spezies direkt einen evolutionären Baum aufzubauen, anstatt nur die daraus ermittelten Distanzen für Paare zur Rekonstruktion von evolutionären Bäumen zu verwenden. Dies führt zu dem so genannten Median Problem, das wir hier nur in seiner einfachsten Form darstellen wollen.

GENERIC MEDIAN PROBLEM

Eingabe: Drei Permutationen $\pi_1, \pi_2, \pi_3 \in S_n$ und eine Distanz $d : S_n \times S_n \rightarrow \mathbb{N}_0$.

Gesucht: Eine Permutation $\sigma \in S_n$, so dass $\sum_{i=1}^3 d(\pi_i, \sigma)$ minimal ist.

Man kann sich leicht überlegen, wie dieses Problem zu erweitern ist, wenn man für k Permutationen einen Baum mit k Blättern (die die Genome respektive ihre

Spezies darstellen) und minimaler Distanz sucht. Hierbei besteht als erste, einfachste Möglichkeit den Baum als einen Stern auszuwählen oder den Grad eines jedes inneren Knoten auf drei festzusetzen.

Als Abstandsfunktion kann man beispielsweise die Reversions-Distanz verwenden, dann erhält man das so genannte Reversal Median Problem.

REVERSAL MEDIAN PROBLEM

Eingabe: Drei Permutationen $\bar{\pi}_1, \bar{\pi}_2, \bar{\pi}_3 \in \bar{S}^n$.

Gesucht: Eine Permutation $\bar{\sigma} \in \bar{S}^n$, so dass $\sum_{i=1}^3 \bar{d}_R(\bar{\pi}_i, \bar{\sigma})$ minimal ist.

Dieses Problem ist bereits dann $\mathcal{NP}$ -hart, wenn man orientierte Permutationen betrachtet. Für k statt drei Sequenzen ist das Problem sogar $\mathcal{APX}$ -vollständig, wie Caprara zeigen konnte.

Eine andere Wahl für die Distanzfunktion ist die so genannte *Breakpoint-Distanz* d_b . Diese liefert die Anzahl derjenigen Gene, die in beiden Genomen auftreten und in den beiden Genomen unterschiedliche Nachfolger besitzen. Manchmal wird diese Distanz auch noch durch die Anzahl der in beiden Genomen gemeinsamen Gene normalisiert.

BREAKPOINT MEDIAN PROBLEM

Eingabe: Drei Permutationen $\pi_1, \pi_2, \pi_3 \in S_n$.

Gesucht: Eine Permutation $\sigma \in S_n$, so dass $\sum_{i=1}^3 d_b(\pi_i, \sigma)$ minimal ist.

Auch hier konnten Bryant bzw. Pe'er und Shamir zeigen, dass dieses Problem bereits $\mathcal{NP}$ -hart ist.

Eine weitere Variante des Median Problems ist die Steiner-Baum Variante. Hier betrachtet man die Menge der Permutationen als Graphen (den so genannten Reversionsgraphen) und verbindet zwei Permutationen, wenn es eine Reversion gibt, die die beiden ineinander überführt. Gesucht ist dann ein minimaler Teilbaum dieses Graphen, der alle gegebenen Permutationen aufspannt.

STEINER REVERSAL MEDIAN PROBLEM

Eingabe: Eine Menge von Permutationen $\pi_1, \dots, \pi_k \in S_n$.

Gesucht: Ein für $\pi_1, \dots, \pi_k \in S_n$ minimal aufspannender Teilbaum des Reversionsgraphen von S_n .

Auch dieses Problem konnte Caprara als $\mathcal{APX}$ -vollständig nachweisen. Man kann sich auch hier leicht Varianten des Modells vorstellen, beispielsweise für andere Permutation, mit Gewichten auf den Operationen oder man kann den vollständigen Graphen auf S_n betrachten, wobei dann die Kantengewichte durch die entsprechenden Distanzen gegeben sind.

A.1 Lehrbücher zur Vorlesung

- S. Aluru (Ed.): *Handbook of Computational Molecular Biology*; Chapman and Hall/CRC, 2006.
- H.-J. Böckenhauer, D. Bongartz: *Algorithmische Grundlagen der Bioinformatik: Modelle, Methoden und Komplexität*; Teubner, 2003.
- D. Gusfield: *Algorithms on Strings, Trees, and Sequences — Computer Science and Computational Biology*; Cambridge University Press, 1997.
- M. Lothaire: *Applied Combinatorics on Words*, Encyclopedia of Mathematics and Its Applications, Cambridge University Press, 2005.
- P.A. Pevzner: *Computational Molecular Biology — An Algorithmic Approach*; MIT Press, 2000.
- J.C. Setubal, J. Meidanis: *Introduction to Computational Molecular Biology*; PWS Publishing Company, 1997.

A.2 Skripten anderer Universitäten

- B. Berger, M. Sing: *Introduction to Computational Molecular Biology*, Massachusetts Institute of Technology,
theory.lcs.mit.edu/~mona/18.417-home.html;
- R. Karp, L. Ruzzo: *Algorithms in Molecular Biology*; CSE 590BI, University of Washington, Winter 1998.
www.cs.washington.edu/education/courses/590bi/98wi/
- S. Kurtz: *Lecture Notes for Foundations of Sequence Analysis*, Universität Bielefeld, 2001.
www.techfak.uni-bielefeld.de/techfak/persons/kurtz/PS/GSA.ps.gz
- S. Rahman: *Foundations of Sequence Analysis*, Universität Bielefeld, 2006.
gi.cebitec.uni-bielefeld.de/teaching/2006winter/sequenzanalyse/SkriptSequenzanalyse.pdf

R. Shamir: *Algorithms in Molecular Biology*, Tel Aviv University,
www.math.tau.ac.il/~rshamir/algmb.html.

M. Tompa: *Computational Biology*, CSE 527, University of Washington, Winter
2000. www.cs.washington.edu/education/courses/527/00wi/

A.3 Originalarbeiten

A.3.1 Optimal Scoring Subsequences

P. Fariselli, M. Finelli, D. Marchignoli, P.L. Martelli, I. Rossi, R. Casadio:
MaxSubSeq: An Algorithm for Segment-Length Optimization. The Case
Study of the Transmembrane Spanning Segments, *Bioinformatics*, Vol. 19,
500–505, 2003.

K.-M. Chung, H.-I. Lu: An Optimal Algorithm for the Maximum-Density Segment
Problem, *SIAM Journal on Computing*, Vol. 34, No. 2, 373–387, 2004.

M. Csűrös: Maximum-Scoring Segment Sets, *IEEE/ACM Transactions on
Computational Biology and Bioinformatics*, Vol. 1, No. 4, 139–150, 2004.

M.H. Goldwasser, M.-Y. Kao, H.-I. Lu: Linear-Time Algorithms for Computing
Maximum-Density Sequence Segments with Bioinformatics Applications, in
*Proceedings of the 2nd Workshop on Algorithms in Bioinformatics
(WABI'02)*, Lecture Notes in Computer Science, Vol. 2452, , 2002; also
Journal of Computer and System Sciences, Vol.70, No. 2, 128–144, 2005.

S.K. Kim: Linear-Time Algorithm for Finding a Maximum-Density Segment of a
Sequence, *Information Processing Letter*, Vol. 86, 339–342, 2003.

Y.-L. Lin, T. Jiang, K.-M. Chao: Efficient Algorithms for Locating the
Length-Constrained Heaviest Segments with Applications to Biomolecular
Sequence Analysis, *Journal of Computer and System Sciences*, Vol. 65,
570–586, 2002; also in *Proceedings of the 27th International Symposium on
Mathematical Foundations of Computer Science, MFCS'02*, Lecture Notes in
Computer Science, Vol. 2420, 459–470, 2002.

W.L. Ruzzo, M. Tompa: A Linear Time Algorithm for Finding All Maximal
Scoring Subsequences, *Proceedings of the 7th International Conference on
Intelligent Systems for Molecular Biology, ISMB'99*, 234–241, 1999.

A.3.2 Suffix-Trees

- R. Giegerich, S. Kurtz, J. Stoye: Efficient Implementation of Lazy Suffix Trees, *Software — Practice and Experience*, Vol. 33, 1035–1049, 2003.
- M. Maaß: *Suffix Trees and Their Applications*, Ausarbeitung von der Ferienakademie, Kurs 2, Bäume: Algorithmik und Kombinatorik, 1999.
www14.in.tum.de/konferenzen/Ferienakademie99/
- E.M. McCreight: A Space-Economical Suffix Tree Construction Algorithm; *Journal of the ACM*, Vol. 23, 262–272, 1976.
- E. Ukkonen: On-Line Construction of Suffix Tress, *Algorithmica*, Vol. 14, 149–260, 1995.

A.3.3 Repeats

- A.S. Fraenkel, J. Simpson: How Many Squares Can a String Contain?, *Journal of Combinatorial Theory, Series A*, Vol. 82, 112–120, 1998.
- D. Gusfield, J. Stoye: Linear Time Algorithm for Finding and Representing All Tandem Repeats in a String, *Journal of Computer and System Sciences*, Vol. 69, 525–546, 2004; see also *Technical Report CSE-98-4*, Computer Science Department, UC Davis, 1998.
- R. Kolpakov, G. Kucherov: Finding Maximal Repetitions in a Word in Linear Time, *Proceedings of the 40th Annual IEEE Symposium on Foundations of Computer Science (FOCS'99)*, 596–604, 1999.
- R. Kolpakov, G. Kucherov: On Maximal Repetitions in Words, *Proceedings of the 12th International Symposium on Fundamentals of Computation Theory (FCT'99)*, Lecture Notes in Computer Science, Vol. 1684, 374–385, 1999.
- R. Kolpakov, G. Kucherov: Finding Approximate Repetitions Under Hamming Distance, *Theoretical Computer Science*, Vol. 303, 135–156, 2003.
- G.M. Landau, J.P. Schmidt: An Algorithm for Approximate Tandem Repeats, *Proceedings of the 4th Symposium on Combinatorial Pattern Matching (CPM'93)*, Lecture Notes in Computer Science, Vol. 684, 120–133, 1993.
- G.M. Landau, J.P. Schmidt, D. Sokol: An Algorithm for Approximate Tandem Repeats, *Journal of Computational Biology*, Vol. 8, No. 1, 1–18, 2001.
- M.G. Main, R.J. Lorentz: An $O(n \log n)$ Algorithm for Finding All Repetitions in a String, *Journal of Algorithms*, Vol. 5, No. 3, 422–432, 1984.

- J. Stoye, D. Gusfield: Simple and Flexible Detection of Contiguous Repeats Using a Suffix Tree, *Theoretical Computer Science*, Vol. 270, 843–856, January 2002.

A.3.4 Lowest Common Ancestors and Range Minimum Queries

- S. Alstrup, C. Gavoille, H. Kaplan, T. Rauhe: Nearest Common Ancestors: A Survey and a New Distributed Algorithm, *Proceedings of the 14th Annual ACM Symposium on Parallel Algorithms and Architectures (SPAA'02)*, 258–264, ACM Press, 2002; to appear in *Theory of Computing Systems*.
- M.A. Bender, M. Farach-Colton: The LCA Problem Revisited, *Proceedings of the 4th Latin American Symposium on Theoretical Informatics (LATIN'00)*, Lecture Notes in Computer Science, Vol. 1776, 88–94, 2000.
- O. Berkman, U. Vishkin: Recursive Star-Tree Parallel Data Structure, *SIAM Journal on Computing*, Vol. 22, 221–242, 1993.
- J. Fischer, V. Heun: Theoretical and Practical Improvements on the RMQ-Problem, with Applications to LCA and LCE. *Proceedings of the 17th Annual Symposium on Combinatorial Pattern Matching (CPM'06)*, Lecture Notes in Computer Science, Vol. 4009, 36–48, Springer, 2006.
- J. Fischer, V. Heun: A New Succinct Representation of RMQ-Information and Improvements in the Enhanced Suffix Array, *Proceedings of the International Symposium on Combinatorics, Algorithms, Probabilistic and Experimental Methodologies (ESCAPE'07)*, to appear in Lecture Notes in Computer Science, Springer, 2007.
- B. Schieber, U. Vishkin: On Finding Lowest Common Ancestors: Simplification and Parallelization, *SIAM Journal on Computing*, Vol. 17, 1253–1262, 1988.

A.3.5 Construction of Suffix-Arrays

- S. Burkhardt, J. Kärkkäinen: Fast Lightweight Suffix Array Construction and Checking, *Proceedings of the 14th Symposium on Combinatorial Pattern Matching (CPM'03)*, Lecture Notes in Computer Science, Vol. 2676, 55–69, 2003.
- D.K. Kim, J.S. Sim, H. Park, K. Park: Linear-Time Construction of Suffix Arrays, *Proceedings of the 14th Symposium on Combinatorial Pattern Matching (CPM'03)*, Lecture Notes in Computer Science, Vol. 2676, 186–199, 2003.

- P. Ko, A. Aluru: Space Efficient Linear Time Construction of Suffix Arrays, *Proceedings of the 14th Symposium on Combinatorial Pattern Matching (CPM'03)*, Lecture Notes in Computer Science, Vol. 2676, 200–210, 2003.
- J. Kärkkäinen, P. Sanders: Simple Linear Work Suffix Array Construction, *Proceedings of the 30th International Colloquium on Automata, Languages and Programming (ICALP'03)*, Lecture Notes in Computer Science, Vol. 2719, 943–955, 2003.
- U. Manber, G. Myers: Suffix Arrays: A New Method for On-Line String Searches, *SIAM Journal on Computing*, Vol. 22, 935–948, 1993.

A.3.6 Applications of Suffix-Arrays

- M.I. Abouelhoda, S. Kurtz, E. Ohlebusch: The Enhanced Suffix Array and Its Applications to Genome Analysis, *Proceedings of the 2nd Workshop on Algorithms in Bioinformatics (WABI'02)*, Lecture Notes in Computer Science, Vol. 2452, 449–463, 2002.
- M.I. Abouelhoda, E. Ohlebusch, S. Kurtz: Optimal Exact String Matching Based on Suffix Arrays, *Proceedings of the 9th International Symposium on String Processing and Information Retrieval (SPIRE'02)*, Lecture Notes in Computer Science, Vol. 2476, 31–43, 2002.
- M.I. Abouelhoda, S. Kurtz, E. Ohlebusch: Replacing Suffix Trees with Enhanced Suffix Arrays, *Journal of Discrete Algorithms*, Vol. 2, 53–86, 2004.
- J. Fischer, V. Heun: A New Succinct Representation of RMQ-Information and Improvements in the Enhanced Suffix Array, *Proceedings of the International Symposium on Combinatorics, Algorithms, Probabilistic and Experimental Methodologies (ESCAPE'07)*, to appear in Lecture Notes in Computer Science, Springer, 2007.
- T. Kasai, G. Lee, H. Arimura, S. Arikawa, K. Park: Linear-Time Longest-Common-Prefix Computation in Suffix Arrays and Its Applications, *Proceedings of the 12th Symposium on Combinatorial Pattern Matching (CPM'01)*, Lecture Notes in Computer Science, Vol. 2089, 181–192, 2001.
- G. Manzini: Two Space Saving Tricks for Linear Time LCP Array Computation, *Proceedings of the 9th Scandinavian Workshop on Algorithm Theory (SWAT'04)*, Lecture Notes in Computer Science, Vol. 3111, 372–383, 2004.
- K. Sadakane: Compressed Suffix Trees with Full Functionality, Theory of Computing Systems, in press, 2007.

A.3.7 Sorting by Reversals

- V. Bafna, P.A. Pevzner: Genome Rearrangements and Sorting by Reversal, *SIAM Journal on Computing*, Vol. 25, 272–289, 1996.
- P. Berman, S. Hannenhalli, M. Karpinski: A 1.375-Approximation Algorithm for Sorting by Reversals, *Proceedings of the 10th Annual European Symposium on Algorithms (ESA'02)*, Lecture Notes in Computer Science, Vol. 2461, 200–210, 2002.
- P. Berman, M. Karpinski: On Some Tighter Inapproximability Results, *Proceedings of the 26th International Colloquium on Automata, Languages and Programming (ICALP'99)*, Lecture Notes in Computer Science, Vol. 1644, 200–209, 1999.
- D. Christie: A $3/2$ -Approximation Algorithms for Sorting by Reversals, *Proceedings of the 9th ACM Symposium on Discrete Algorithms (SODA'98)*, 244–252, 1998.

A.3.8 Sorting by Oriented Reversals

- D.A. Bader, N.M.E. Moret, M. Yan: A Linear-Time Algorithm for Computing Inversion Distance Between Signed Permutations With an Experimental Study, *Journal of Computational Biology*, Vol. 8, No. 5, 483–491, 2001.
- A. Bergeron: A Very Elementary Presentation of the Hannenhalli-Pevzner Theory, *Proceedings of the 12th Symposium on Combinatorial Pattern Matching (CPM'01)*, 106–117, 2001.
- A. Bergeron, J. Mixtacki, J. Stoye: Reversal Distance without Hurdles and Fortresses, *Proceedings of the 15th Symposium on Combinatorial Pattern Matching (CPM'04)*, Lecture Notes in Computer Science, Vol. 3109, 388–399, 2004.
- P. Berman, S. Hannenhalli: Faster Sorting by Reversals, *Proceedings of the 7th Symposium on Combinatorial Pattern Matching (CPM'96)*, 168–185, 1996.
- S. Hannenhalli, P. Pevzner: Transforming Cabbage into Turnip: Polynomial Algorithm for Sorting Signed Permutations by Reversals, *Journal of the ACM*, Vol. 46, 1–27, 1999; also in *Proceedings of the 27th Annual ACM Symposium on Computing (STOC'95)*, 178–189, 1995.
- S. Hannenhalli, P. Pevzner: To Cut ... or not to Cut (Applications of Comparative Physical Maps in Molecular Evolution), *Proceedings of the 7th ACM Symposium on Discrete Algorithms (SODA'96)*, 304–313, 1996.

- H. Kaplan, R. Shamir, R. Tarjan: Faster and Simpler Algorithm for Sorting Signed Permutations by Reversals, *SIAM Journal on Computing*, Vol. 29, 880–892, 1999; also in *Proceedings of the 8th ACM-SIAM Symposium on Discrete Algorithms (SODA'97)*, 344–351, 1997.
- A.C. Siepel: An Algorithm to Enumerate all Sorting Reversals, *Journal Computational Biology*, Vol. 10, 575–597, 2003; also in *Proceedings of the 6th Annual International Conference on Computational Biology (RECOMB'02)*, 281–290, 2002.

A.3.9 Sorting by Transpositions

- V. Bafna, P.A. Pevzner: Sorting by Transpositions, *SIAM Journal on Discrete Mathematics*, Vol. 11, No. 2, 224–240, 1998.
- I. Elias, T. Hartman: A 1.375-Approximation Algorithm for Sorting by Transpositions, *Proceedings of the Fifth Workshop on Algorithms in Bioinformatics (WABI'05)*, Lecture Notes in Computer Science, Vol. 3692, 204–215, 2005.
- T. Hartman, R. Shamir: A Simpler 1.5-Approximation Algorithms for Sorting by Transpositions, *Proceedings of the 14th Symposium on Combinatorial Pattern Matching (CPM'03)*, Lecture Notes in Computer Science, Vol. 2676, 1560–169, 2003.

A.3.10 Sorting by Transversals

- M. Bader, E. Ohlebusch: Sorting by Weighted Reversals, Transpositions, and Inverted Transpositions, *Proceedings of the 10th Annual International Conference on Research in Computational Molecular Biology (RECOMB'06)*, Lecture Notes in Computer Science, Vol. 3909, 563–577, Springer, 2006.
- N. Eriksen: $(1 + \varepsilon)$ -Approximation of Sorting by Reversals and Transpositions, *Theoretical Computer Science*, Vol. 289, 517–529, 2002.
- Q.-P. Gu, S. Peng, I.H. Sudborough: A 2-Approximation for Genome Rearrangements by Reversals and Transpositions, *Theoretical Computer Science*, Vol. 210, 327–339, 1999.
- G-H. Lin, G. Xue: Signed Genome Rearrangements by Reversals and Transpositions: Models and Approximations, *Theoretical Computer Science*, Vol. 259, 513–531, 2001

A.3.11 Erweiterungen zu Genome Rearrangements

- M.A. Bender, D. Ge, S. He, H. Hu, R.Y. Pinter, S. Skiena, F. Swidan: Improved Bounds on Sorting with Length-Weighted Reversals, *Proceedings of the 15th Annual ACM-SIAM Symposium on Discrete Algorithms, SODA'04*, 919–928, 2004.
- A. Bergeron, J. Mixtacki, J. Stoye: On Sorting by Translocations, *Proceedings of the 9th Annual International Conference on Computational Biology (RECOMB'05)*, 2005, to appear.
- D. Bryant: A Lower Bound for the Breakpoint Phylogeny Problem, *Proceedings of the 11th Workshop on Combinatorial Pattern Matching, CPM'00*, 235–247, 2000.
- A. Caprara: Formulations and Hardness of Multiple Sorting by Reversals, *Proceedings of the 3rd Annual International Conference on Computational Biology (RECOMB'99)*, 84–93, 1999.
- T. Chen, S.S. Skiena: Sorting with Fixed-Length Reversals, *Discrete Applied Mathematics*, Vol. 71, 269–295, 1996.
- D.A. Christie, R.W. Irving: Sorting Strings by Reversals and by Transpositions, *SIAM Journal on Discrete Mathematics*, Vol. 14, 193–206, 2001.
- D.A. Christie: Sorting Permutations by Block-Interchanges, *Information Processing Letters*, Vol. 60, 165–169, 1996.
- B. DasGupta, T. Jiang, S. Kannan, M. Li, E. Sweedyk: On the Complexity and Approximation of Syntenic Distance, *Discrete Applied Mathematics*, Vol. 88, 59–82, 1998.
- N. El-Mabrouk, D. Sankoff: The Reconstruction of Doubled Genomes, *SIAM Journal on Computing*, Vol. 32, 754–792, 2003.
- S. Hannehalli: Polynomial-Time Algorithm for Computing Translocation Distance Between Genomes, *Discrete Applied Mathematics*, Vol. 71, 137–151, 1996.
- S. Hannenhalli, P.A. Pevzner: Transforming Men into Mice (Polynomial Algorithm for Genomic Distance), *Proceedings of the 36th Annual IEEE Symposium on Foundations of Computer Science (FOCS'95)*, 581–592, 1995.
- L.S. Heath, J.P. Vergara: Sorting by Bounded Block-Moves, *Discrete Applied Mathematics*, Vol. 88, 181–206, 1998.

-
- J. Kleinberg, D. Liben-Nowell: The Syntenic Diameter of the Space of n -Chromosome Genomes, in *Comparative Genomics*, David Sankoff and Joseph H. Nadeau (Eds.), Kluwer Academic Press, 2000.
- J. Meidanis, Z. Dias: Genome Rearrangements Distance by Fusion, Fission, and Transposition is Easy, *Proceedings of the 8th Symposium on String Processing and Information Retrieval, SPIRE'01*, 13–15, 2001
- R.Y. Pinter, S. Skiena: Genomic Sorting with Length-Weighted Reversals, *Genome Informatics*, Vol. 13, 103–111, 2002.

Symbole

–2-Move, 252
0-Move, 252
2-Move, 252

A

α -MSS, 16
Adjazenz, 219, 225
 orientierte, 225
aktives Suffix, 59
Alignment-Distanz, 69
All Maximal Scoring Subsequences, 8
alternierende Zyklenzerlegung, 220
AMSS, 8
 $\mathcal{APX}$, 210
atomic suffix tree, 47
ausreichend, 113

B

BAMSS, 24
Baum
 Tiefe, 121
 vollständiger, 133
binärer Baum, 133
Blattliste, 82
BMSS, 27
Bounded All Maximum Scoring
 Subsequences, 24
Bounded Maximal Scoring
 Subsequence, 27
Breakpoint, 211, 225
 orientierter, 225
Breakpoint Median Problem, 277
Breakpoint-Distanz, 277
Breakpoint-Graph, 218
 erweiterter, 218

C

Child-Tabelle, 184
compact suffix tree, 47

D

Darstellung eines Wortes, 44
Deletion, 205
Desire-Edges, 218
Divide-and-Conquer, 5
doppelte Reversion, 257
Duplikation, 206
dynamische Programmierung, 4, 123

E

EDIT-Distanz, 69
einfache Hurdle, 234
einfache Transversal-Distanz, 266
Elter-Intervall, 182, 196
Enhanced-Suffix-Array, 184, 191
enthalten von ℓ -Intervallen, 182, 196
erfolgloser Suffix-Link-Walk, 113
erfolgreicher Suffix-Link-Walk, 113
erweiterte orientierte Permutation,
 225
erweiterte Permutation, 210
erweiterter Breakpoint-Graph, 218
erweitertes Reality-Desire-Diagram,
 226, 249, 258
Euler-Tour, 120
exaktes Paar, 69
exaktes Repeat, 69

F

fallend rechtsschiefe Partition, 32
fallender Strip, 211
Feld
 inkrementelles, 125
 normalisiertes, 125

Fission, 206

Folge

links-negative, 27

rechtsschiefe, 32

Fortress, 234

starkes, 267

Fusion, 206

G

GC-reiche Regionen, 3

Geburtstagsparadoxon, 51

Generisches Median Problem, 276

Gewicht, 271

gewichtete Transversal-Distanz, 266

Gruppe

orientierte symmetrische, 224

symmetrische, 208

gute Komponente, 233

H

Hamming-Distanz, 69

Hashfunktion, 50

Hashing, 50

Höhe

eines Knotens, 137

Hurdle, 233

einfache, 234

starke, 267

Hurdle-Cutting, 245

Hurdle-Merging, 244

I

Incremental Range Minimum Query,
124

inkrementelles Feld, 125

Inorder-Nummerierung, 133, 134

Insertion, 205

Inversion, 206

iterierter rechtsschiefer Zeiger, 41

J

j -ter Block, 100

K

k -difference Repeat, 69

k -mismatch Repeat, 69

k -mismatch Tandem-Repeat, 93

kanonische Lokation, 64

Kante

orientierte, 228

unorientierte, 228

Kartesischer Baum, 127

Kette von Suffix-Link-Walks, 113

Kind-Intervall, 182, 196

Knoten

Höhe, 137

orientierter, 230

Tiefe, 121

unorientierter, 230

Kollision, 51

kompakter Σ^+ -Baum, 44

kompaktifizierter Trie, 44

komplementäres Zeichen, 78

Komponente

gute, 233

orientierte, 230

schlechte, 233

unorientierte, 230

konservierte Regionen, 3

Kopf des Runs, 140

Kreis

orientierter, 228, 251

unorientierter, 228, 251

L

ℓ -Index, 181, 195

ℓ -Intervall, 180, 194

Länge eines Tandem-Repeat-Paares,
80

Länge eines Tandem-Repeats, 80

längster gemeinsamer Präfix, 175

lca, 91

lce, 90

lcp, 175

LCP-Intervall vom Typ ℓ , 180, 194

LCP-Intervall-Baum, 182, 196
 least common ancestor, 119
 Lempel-Ziv-Zerlegung, 100
 linkeste Überdeckung, 97
 links-negative Folge, 27
 links-negativer Zeiger, 28
 linksdivers, 74
 Linksrotation, 80
 linksverzweigend, 80
 linksverzweigendes
 Tandem-Repeat-Paar, 80
 Linkszeichen
 einer Position, 74
 eines Blattes, 74
 Lokation, 63
 kanonische, 64
 offene, 64
 longest common backward extension,
 90
 longest common extension, 90, 91
 longest common forward extension, 90
 longest common prefix, 175
 lowest common ancestor, 91, 119
 Lowest Common Ancestor Query, 119

M

MASS, 31
 Matching, 219
 Maximal Average Scoring
 Subsequence, 31
 maximal bewertete Teilfolge, 9
 Maximal Scoring Subsequence, 1
 maximales Paar, 72
 maximales Repeat, 72
 Median Problem
 Breakpoint, 277
 generisches, 276
 Reversal, 277
 Steiner, 277
 Menge aller maximalen Paare, 72
 Menge aller maximalen Repeats, 72
 Mensch-Maus-Genom, 3

Min-SBR, 209
 Min-SBRT, 258
 Min-SBT, 248
 Min-SBTR, 258
 Min-SOR, 225
 Min-SWT, 267
 minimal links-negative Partition, 27
 MSS, 1
 MSS(a), 10

N

nearest common ancestor, 119
 nested suffix, 57
 niedrigster gemeinsamer Vorfahre, 91,
 119
 normalisiertes Feld, 125

O

offene Lokation, 64
 offene Referenz, 64
 orientierte Adjazenz, 225
 Orientierte Kante, 228
 orientierte Komponente, 230
 orientierte Permutation, 224
 erweiterte, 225
 orientierte Reversion, 225
 orientierte symmetrische Gruppe, 224
 orientierter Breakpoint, 225
 orientierter Knoten, 230
 orientierter Kreis, 228, 251
 Overlap-Graph, 229

P

Paar
 exaktes, 69
 maximales, 72
 Tandem-Repeat-, 80
 Pancake Flipping, 223
 Partition
 fallend rechtsschiefe, 32
 minimal links-negative, 27
 PCP-Theorem, 210
 Permutation, 208

- als Abbildung, 209
- als Umordnung, 209
- erweiterte, 210
- erweiterte orientierte, 225
- orientierte, 224
- zugehörige unorientierte, 226, 248
- Position vom Typ L, 164
- Position vom Typ S, 164
- Präfix
 - längster gemeinsamer, 175
- primitiv, 81
- primitives Tandem-Repeat, 81
- probabilistisch verifizierbare Beweise, 210

Q

- Query
 - Lowest Common Ancestor, 119
 - Range Minimum, 120

R

- Range Minimum Query, 120, 193, 202
- Reality-Desire-Diagram, 218, 226, 249, 258
 - erweitertes, 226, 249, 258
- Reality-Edges, 218
- Rearrangement Distance, 208
- Rechtsrotation, 80
- rechtsschiefe Folge, 32
- rechtsschiefer Zeiger, 36
 - iterierter, 41
- rechtsverzweigend, 80
- rechtsverzweigendes
 - Tandem-Repeat-Paar, 80
- rechtsverzweigendes Teilwort, 57
- Referenz, 47
 - offene, 64
- Repeat, 69
 - exaktes, 69
 - k -difference, 69
 - k -mismatch, 69
 - maximales, 72
 - revers-komplementäres, 78

- Tandem, 80
- revers-komplementäres Repeat, 78
- revers-komplementäres Wort, 78
- Reversal Median Problem, 277
- Reversion, 206, 225
 - doppelte, 257
 - orientierte, 225
 - sichere, 243
- Reversions-Distanz, 266
- Reversionsgraph, 277
- rightbranching, 57
- Run, 139
 - Kopf, 140
- Run von Tandem-Repeats, 97

S

- S-Distanz, 165
- schlechte Komponente, 233
- Score, 237
- sichere Reversion, 243
- Σ -Baum, 43
- Σ^+ -Baum, 43
 - kompakter, 44
- Singletons, 182, 196
- Sorting by Oriented Reversals, 225
- Sorting by Prefix Reversals, 223
- Sorting by Reversals, 209
- Sorting by Reversals and Transpositions, 258
- Sorting by Transpositions, 248
- Sorting by Transversals, 258
- Sorting by weighted Transversals, 267
- stark unorientiert, 267
- starke Hurdle, 267
- starke Super-Hurdle, 267
- starkes Fortress, 267
- steigender Strip, 211
- Steiner Reversal Median Problem, 277
- Strip, 211
 - fallend, 211
 - steigend, 211
- Suffix

- aktives, 59
- nested, 57
- verschachteltes, 57
- Suffix-Array, 151
 - Enhanced-, 191
- Suffix-Baum, 46
- Suffix-Link, 56
- Suffix-Link an Position i , 192, 201
- Suffix-Link-Intervall, 193, 202
- Suffix-Link-Walk, 113
 - erfolgloser, 113
 - erfolgreicher, 113
 - Kette, 113
- Suffix-Tree, 46
- Suffix-Trie, 45
- Super-Hurdle, 234
 - starke, 267
- symmetrische Gruppe, 208
 - orientierte, 224
- Synteny, 276

T

- Tandem-Arrays, 81
- Tandem-Repeat, 80
 - k -mismatch, 93
 - Länge, 80
 - primitives, 81
- Tandem-Repeat-Paar, 80
 - gleicher Typ, 96
 - Länge, 80
 - linksverzweigend, 80
 - rechtsverzweigend, 80
 - Zentrum, 102
- Teilwort, 45
 - rechtsverzweigendes, 57
- Teilwort vom Typ L, 164
- Teilwort vom Typ S, 164
- Tiefe
 - eines Baums, 121
 - eines Knotens, 121
- Translokation, 206
- Translokations-Distanz, 276

- Transmembranproteine, 2
- Transposition, 206
- Transpositions-Distanz, 248
- Transversal-Distanz, 258
 - einfache, 266
 - gewichtete, 266
- Transversals, 257
- Trie, 43
 - kompaktifizierter, 44
- Tripel, 170
- Typ
 - gleicher, 96
- Typ L, 158
- Typ S, 158

U

- überdecken, 97
- Überdeckung, 97
 - linkeste, 97
- ungapped local alignment, 3
- unorientiert
 - stark, 267
- unorientierte Kante, 228
- unorientierte Komponente, 230
- unorientierter Knoten, 230
- unorientierter Kreis, 228, 251

V

- verschachteltes Suffix, 57
- Vokabular, 96
- vollständiger Baum, 133
- Vorfahre
 - niedrigster gemeinsamer, 119

W

- Wort
 - revers-komplementäres, 78
- Worttiefe, 44, 82
- WOTD-Algorithmus, 53

X

- x -Move, 252

Z

Zeichen

komplementäres, 78

Links-, 74

Zeiger

iterierter rechtsschiefer, 41

links-negativer, 28

rechtsschiefer, 36

Zentrum eines

Tandem-Repeats-Paares, 102

zugehörige unorientierte Permutation,

226, 248

Zyklenzerlegung

alternierende, 220