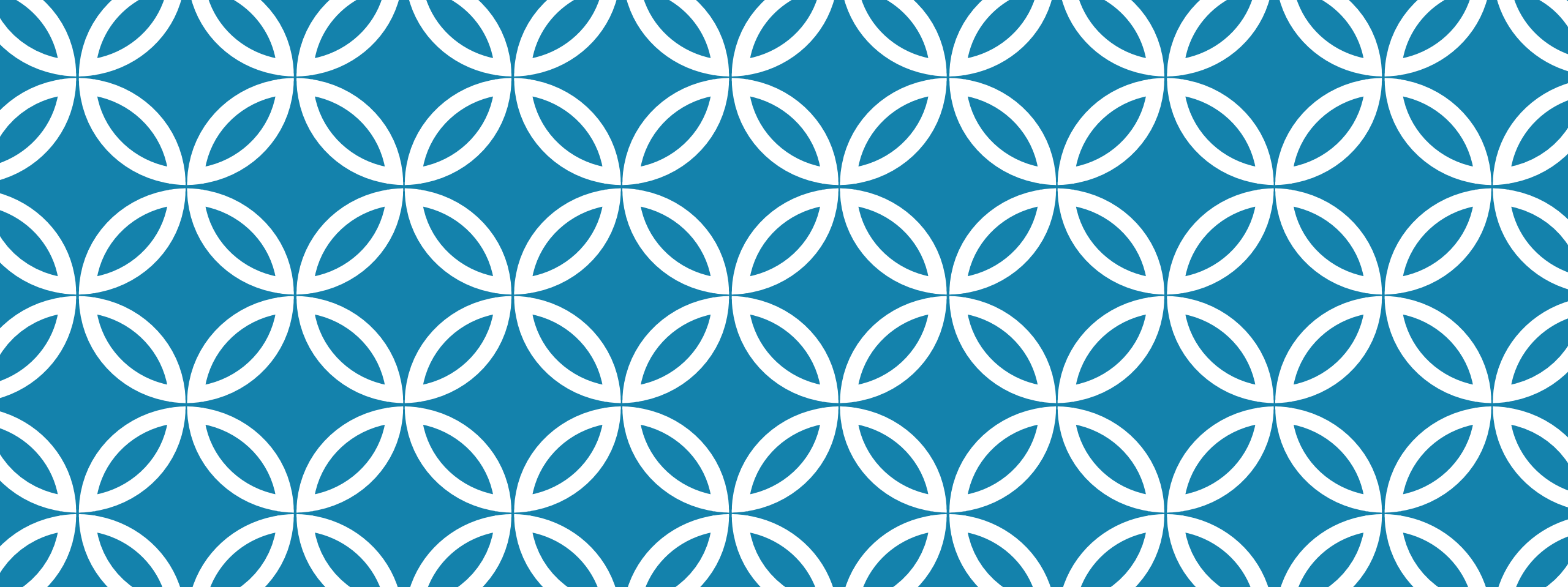




الگوریتم
رتبه

محسن هوشمند
دانشکده علم رایانه و تکنولوژی اطلاعات
دانشگاه تحصیلات تکمیلی علوم پایه زنجان

رتبه

حجم زیاد علیه درک مناسب

- سختی فهم انسان از استخراج اطلاعات در حجم زیاد داده‌های بدون ساختار
- تلخیص داده‌ها با استخراج کمیت‌های آماری
- از راه‌حل‌های مهم در مواجهه با چنین داده‌هایی

الگوریتم‌ها و ساختارهای داده محاسبه ویژگی‌های مبتنی بر رتبه داده‌ها

- با استفاده از حافظه کم و یک گذر از داده‌ها

چندک

پراکندگی در رتبه
چندکها

q - چندک

▪ $(0 \leq q \leq 1)$

▪ مقداری از دنباله مقادیر

▪ کسر q از مقادیر دنباله کمتر یا مساوی آن

▪ $1 - q$ باقی مانده بیشتر یا مساوی آن

در صورت دنباله دارای n مقدار

▪ مقدار q-چندک مقداری از دنباله با رتبه $q \cdot n$

صدکها یا درصدها

چندکهای تقسیم کننده دنباله مرتب شده به ۱۰۰ قسمت مساوی

95-امین صدک متناظر ۰,۹۵-چندک

چندکهای ۰ و ۱ به ترتیب حداقل و حداکثر مقادیر در دنباله

چندک ۰,۵ متناظر میانه مقادیر

چندک

ایان مونرو و مایکل پاترسون در سال ۱۳۵۹

- اثبات یافتن چندک خاص در دقیقه p گذر از داده‌ها نیازمند به حافظه $O(n^{\frac{1}{p}})$
- به معنای نبودن الگوریتم تک‌گذری تولید مقدار دقیق چندک در فضای زیرخطی
- تشویق به جستجوی الگوریتم‌های محاسبه چندک‌های تقریبی

چندک

در عمل، قابل تحمل بودن وجود خطا در محاسبه چندک

- اعمال چنین تخمین‌هایی روی داده‌های نوفه‌دار
- تقریب توزیع‌های داده ناشناخته

در نتیجه: در پی چندک ϵ -تقریبی q در بیشتر موارد

به معنای قرار گرفتن رتبه مقداری $[(q - \epsilon) \cdot n, (q + \epsilon) \cdot n]$

- n تعداد مقادیر و $0 < \epsilon < 1$ پارامتر خطا

- امکان واجد شرایط بودن بیش از یک مقدار

چندک

تخمین ویژگی‌های مختلف رتبه
▪ مانند خلاصه‌های چندک

نقشی مهم در روش‌های تشخیص نقاط پرت جریانی

مثال

- نظارت بر تراکنش‌های تجارت الکترونیک آنلاین جهت تشخیص تقلب در کارت اعتباری
- در پی موقعیت‌ها و مکان‌های پرداخت غیرمعمول
- مقادیری خارج از صدک نود و نهم توزیع مکان معمول برای مشتریان

مثال - تشخیص تقلب

تقلب مالی از مسائل اساسی پیش روی صنعت مالی

در سال ۱۳۹۴، تقلب در کارتهای اعتباری و نقدی جهانی منجر به زیان‌هایی بالغ بر ۲۱,۸۴ میلیارد دلار

تولید برنامه‌های جستجو و شناسایی علائم تقلب مالی

- غالباً استفاده از متغیرهای خاص متعددی

- بررسی «درجه دورافتادگی» یا درجه پرتی آنها در هر مشاهده

- مثال، استفاده از متغیرهایی مانند کل مبلغ هزینه شده در کارت اعتباری و مقدار هزینه شده در روز

برای هر مشاهده،

- امکان تقریب درجه دورافتادگی با چندک‌های برخی از توزیع‌های هزینه

- شناسایی و برچسب زدن مشاهدات به مشکوک به تقلب با نقاط پرت

- از طریق مقایسه با برخی چندک‌های بالا، مثل چندک ۰,۹۵

چندک

نظارت بر ترافیک وب از دیگر حوزه کاربردی عظیم خلاصه‌های رتبه
بررسی خلاصه‌ها متناظر با تشخیص زودتر مشکلات
بدون بررسی داده‌های واقعی

مثال - نظارت بر تارمانه

مدیریت میلیون‌ها کاربر در تارمانه‌های بزرگ به صورت روزانه

شهریور و مهر ۱۳۹۶

- ویکی‌پدیا: پردازش حدود ۵۰۰ میلیون بازدید در روز در تمام زبان‌های خود
- معادل تقریباً ۵,۷ هزار درخواست در ثانیه با استفاده از بیش از ۳۰۰ سرور در سراسر جهان

تأخیر در تارمانه یا به زبان دقیق‌تر تأخیر بین لحظه ایجاد محتوا و لحظه انتقال به بازدیدکننده

- از جمله گزینه‌های مهم عملکرد هر تارمانه
- شایع بودن چولگی در توزیع احتمالی میزان تأخیر
- اعمال نظارت معمولاً با ردیابی برخی چندک‌ها یا صدک‌های خاص بالا

رایج‌ترین سوالات

- تأخیر ۹۵ درصد از درخواست‌ها برای سرور وب منفرد چقدر است؟
- تأخیر ۹۹ درصد از درخواست‌ها برای کل وب‌سایت چقدر است؟
- تأخیر ۹۵ درصد از درخواست‌ها برای کل وب‌سایت در ۱۵ دقیقه گذشته چقدر بود؟

مثال - نظارت بر تارمانه

امکان پاسخ به چنین سوالاتی با محاسبه چندک

- دارای تفاوت از نظر فنی
- در نتیجه نیاز به استفاده از روش‌های مختلف

سوال اول

- محاسبه خلاصه برای هر جریان منفرد

سوال دوم

- نیازمند الگوریتم‌های توزیع شده جهت محاسبه آمار روی داده‌های جریان‌های متعدد

سوال سوم

- نیازمند به زیرمجموعه‌ای از داده‌های جریان تعریف شده در یک پنجره زمانی
- تغییر مداوم چنین زیرمجموعه‌هایی

پرسش یا پرس وجوی چندک

وظیفه یافتن q -چندک،

یا یافتن مقادیری از دنباله مرتب از n مقدار دارای رتبه $q.n$
▪ $q \in (0,1)$

پرس وجوی میانه

▪ موردی خاص از پرس وجوی چندک با $q = 0.5$

محاسبه چندک

متاخر نبودن مسئله محاسبه چندک

- دارای روش‌های مناسب در محاسبات کلاسیک

وجود چالش‌های جدیدی برای جریان‌های نامحدود

- رایج در برنامه‌های داده بزرگ

- با داشتن محدودیت در حافظه و امکان‌پذیری تک‌گذاری از داده‌ها

الگوریتم طرح کلی شمارش-کمینه

- امکان محاسبه مقادیر تقریبی چندک‌ها

- اما نیازمند حافظه بسیار بیشتر نسبت به الگوریتم‌های مورد بحث

پرس و جوی چندک معکوس

نوع دیگری از مسئله مربوط به چندک

یافتن رتبه مقدار داده شده در دنباله‌ای مرتب از n مقدار

در صورت یافتن آن و داشتن تعداد کل مقادیر n ، محاسبه چندک متناظر q :

$$q = \frac{1}{n} \cdot rank(x)$$

چندک بازه

در بسیاری از برنامه‌ها،

- اهمیت یافتن تعداد مقادیر از دنباله مرتب از n مقدار در بازه $[a, b]$
- غالباً شناخته شده به عنوان پرس‌وجوی بازه
- محاسبه چنین عددی
- تعیین رتبه‌های مرزهای محدوده و برگرداندن تفاوت آنها

چند روش

الگوریتم‌های نمونه‌برداری تصادفی

الگوریتم **q-digest** یا تلخیص-**q** بر مبنای درخت ساده

الگوریتم **t-digest**

▪ استفاده از خوشه‌بندی برای تخمین کارآمد آمار مبتنی بر رتبه در جریان‌های نامحدود

نمونه برداری تصادفی

نمونه برداری تصادفی

- انتخاب بدون جایگزینی زیرمجموعه‌ای تصادفی از داده‌ها
- مورد استفاده در بسیاری از الگوریتم‌های علم رایانش

در مسائل رتبه

- امکان استفاده از چنین روشی
- در تهیه گزارش چندک‌های محاسبه شده روی نمونه‌ها
- به عنوان تقریبی از چندک‌های کل جریان داده

نمونه برداری تصادفی

مزیت متمایز چنین مجموعه نمونه‌هایی

- کوچکتر بودن آنها
- امکان استفاده از الگوریتم‌های قطعی کلاسیک در پاسخ به چنین پرس‌وجوهای چندک رتبه روی مجموعه نمونه

چالش

- با توجه به چولگی داده‌ها
- نیاز به داشتن تضمین‌های قبلی در مورد خطای چنین تقریبی،
- روشی خاص در نمونه برداری تصادفی
- امکان وابستگی روش به داده‌ها

نمونه برداری تصادفی

مشکل دیگر نمونه برداری ساده

- بسیاری از طرح های نمونه برداری نیاز به اطلاع قبلی از اندازه مجموعه داده
- ناممکن یا مشکل آفرین در جریان های پیوسته معمول در برنامه های داده بزرگ

از راه حل های ممکن

- روش نمونه برداری مخزنی
- پیشنهاد جفری ویتز در سال ۱۳۶۴
- امکان تولید نمونه بدون اطلاع از تعداد
- استفاده از آن در مسئله چندک
- نیازمند به حافظه ای قابل توجه

نمونه برداری مخزن

الگوریتم ۱ - نمونه برداری مخزنی

ورودی: جریان داده D و میزان چندک q

ورودی: بافر خالی B با اندازه k

خروجی: بافر پر شده B

$B \leftarrow D[1:k]$

برای i از $k+1$ تا انتها

$r = rnd(1, i)$

اگر $r \leq k$

$B[r] = D[i]$

مرتب سازی B

$j = \lceil q \times k \rceil$

برگرداندن $B[j]$

نمونه برداری مخزنی

[3, 1, 7, 2, 9, 4, 8, 5, 6]

K=4

B=[3, 1, 7, 2]

نمونه برداری مخزنی

[3, 1, 7, 2, 9, 4, 8, 5, 6]

$K=4$

$B=[3, 1, 7, 2]$

$i=5$

▪ مقدار ۹

▪ $r=3$

▪ $r \leq k$

▪ $B=[3, 1, 9, 2]$

نمونه برداری مخزنی

[3, 1, 7, 2, 9, 4, 8, 5, 6]

$K=4$

$B=[3, 1, 9, 2]$

$i=6$

▪ مقدار 4

▪ $r=5$

▪ $r > k$

▪ $B=[3, 1, 9, 2]$

نمونه برداری مخزنی

[3, 1, 7, 2, 9, 4, 8, 5, 6]

$K=4$

$B=[3, 1, 9, 2]$

$i=7$

▪ مقدار 8

▪ $r=2$

▪ $r \leq k$

▪ $B=[3, 8, 9, 2]$

نمونه برداری مخزنی

[3, 1, 7, 2, 9, 4, 8, 5, 6]

$K=4$

$B=[3, 8, 9, 2]$

$i=8$

▪ مقدار 5

▪ $r=6$

▪ $r > k$

▪ $B=[3, 8, 9, 2]$

نمونه برداری مخزنی

[3, 1, 7, 2, 9, 4, 8, 5, 6]

$K=4$

$B=[3, 8, 9, 2]$

$i=9$

▪ مقدار 6

▪ $r=1$

▪ $r \leq k$

▪ $B=[6, 8, 9, 2]$

مرتب سازی

▪ $B=[2, 6, 8, 9]$

▪ میانه برابر ۷ در اصلی برابر ۵

نیاز به حافظه $O\left(\frac{1}{\epsilon^2}\right)$

نمونه برداری تصادفی مرل MRL

الگوریتم نمونه برداری تصادفی

▪ MRL

▪ پیشنهاد گورمیت سینگ مانکو، سریده‌هار راجاگوپالان و بروس لیندسی

▪ در سال ۱۳۷۸

▪ جهت حل مسئله نمونه برداری صحیح و تخمین چندک

شامل

▪ روش نمونه برداری غیریکنواخت

▪ الگوریتم یافتن چندک قطعی

نمونه برداری تصادفی مرل

پشتیبانی از پردازش جریان داده پیوسته با نیازهای فضای کم

- پیشنهاد اصلاح غیریکنواخت از نمونه برداری مخزنی
- گنجاندن مقادیری که زودتر در دنباله ظاهر می شوند با احتمال بیشتر نسبت به سایرین
- بهبود کارایی فضایی
- دقیق تر از نمونه برداری مخزنی اصلی

نمونه برداری تصادفی مرل

معایب اصلی

- تعیین پارامترهای پیکربندی آن با حل مسئله بهینه سازی
- پیچیدگی آن و استفاده از رویه های پیچیده

معرفی نسخه ای ساده از الگوریتم مرل

- پیشنهادی گه لو، لو وانگ، که یی، و گراهام کورمودی
- در سال ۱۳۹۲
- در مقالات اصلی با عنوان Random

الگوریتم ت

پردازش داده‌های جریان داده در قطعات با اندازه متغیر

نمونه‌برداری از آنها

تولید نهایی نمونه‌برداری غیریکنواخت

الگوریتم ت

استفاده از داده ساختار یا آرایه ای شخصی شده برای ذخیره نمونه های مقادیر

حاصل از b واحد داده $B_1, B_2, \dots, B_b$

▪ معروف به بافر

هر یک امکان ذخیره تا حداکثر k مقدار

مرتبط شدن با سطح L یا سطح پر شدن

الگوریتم ت

پارامتر سطح L

▪ نمایانگر احتمال انتخاب مقادیر

وابسته به

▪ تعداد مقادیر n تاکنون پردازش شده

▪ حداکثر ارتفاع مجاز h درخت نمایش دهنده دنباله عملیات انجام شده الگوریتم

$$L = L(n, h) = \max\left(0, \left\lceil \log \frac{n}{k^{2^{h-1}}} \right\rceil\right)$$

$$L(0, h) = 0$$

الگوریتم ت

برای پر کردن بافر خالی B_i^L
▪ $i \in 0 \dots b$
▪ در سطح L

انتخاب k مقدار تصادفی از $k \cdot 2^L$ مقدار ورودی متوالی،
هر مقدار به ازای هر بلوک از 2^L بلوک

ذخیره مقادیر در B_i^L

در پایان رویه
▪ بافر کمتر از k مقدار
▪ به دلیل نبود مقدار به اندازه کافی در دنباله ورودی
▪ در صورت وجود حداقل یک مقدار در بافر
▪ برچسب گذاری به عنوان پر

الگوریتم ت

احتمال انتخاب مقداری خاص از جریان داده ورودی و ذخیره در بافر
▪ وابسته به سطح L
▪ اندازه قطعه 2^L کنترل کننده استخراج مقادیر از آن

پیاده سازی عملی نمونه برداری غیریکنواخت مورد استفاده الگوریتم

الگوریتم ۱- پر کردن بافرهای خالی

ورودی: جریان داده D

ورودی: بافر خالی B^L با اندازه k در سطح L

خروجی: بافر پر شده B^L و برچسب آن

برای i از ۰ تا $k - ۱$ انجام بده

$S \leftarrow next(\mathfrak{L}, D)$ // خواندن $\mathfrak{L}$ مقدار بعدی از D

اگر $S = \emptyset$ آنگاه

break

$x \leftarrow$ نمونه $(s \in S)$ // انتخاب تصادفی مقدار از S

$B^L \leftarrow B^L \cup \{x\}$

خالی $\leftarrow$ برچسب

اگر $۰ < (B^L)$ تعداد آنگاه

پر $\leftarrow$ برچسب

برگرداندن B^L ، برچسب

الگوریتم ت

امکان ادغام دو بافر از یک سطح L جهت بازیابی فضای بافر

- منجر به بافری جدید با اندازه مشابه در سطح $L + 1$

جهت ادغام دو بافر

- مرتب‌سازی دنباله مقادیر هر دو بافر

- انتخاب تصادفی نیمی از مقادیر

- به عنوان مثال، با انتخاب همه مقادیر در موقعیت‌های زوج یا فرد

علامت‌گذاری بافرهای ادغام شده به عنوان خالی

- علامت‌گذاری بافر خروجی به عنوان پر

الگوریتم ۲- ادغام دو بافر غیر خالی

ورودی: بافرهای غیر خالی B_i^L ، B_j^L با اندازه k در سطح L

خروجی: بافر پر شده B^{L+1} در سطح L و برجسب آن

$S \leftarrow \text{sort}(B_i^L \cup B_j^L)$

$\text{free}(B_i^L)$

$\text{free}(B_j^L)$

$B^{L+1} \leftarrow \text{نمونه}(S, k)$ // انتخاب تصادفی k مقدار از بافرهای پیوسته

برگرداندن B^{L+1} ، پر

الگوریتم ت

عملیات ادغام

- مرتب‌سازی بافرها نیاز به $O(k \log k)$
- آماده کردن جمعیت بعدی بافر در زمان $O(k)$

فرآیند ساخت داده‌ساختار

- شامل سری مراحل جمعیت بافر و عملیات ادغام

آغاز

- برچسب‌گذاری هر بافر به عنوان خالی

پردازش جریان ورودی با فعال‌سازی سطح L

- در ابتدا برابر با صفر به دلیل عدم وجود مقادیر پردازش شده

در صورت خالی بودن بافر خالی B

- پر کردن با استفاده از الگوریتم ۱ با خواندن $k \cdot 2^L$ مقدار از جریان

پر شدن همه بافرها

- یافتن حداقل دو بافر با پایین‌ترین سطح
- ادغام دو بافر

مثال - ساخت بافرهای نمونه

مجموعه داده از ۲۵ عدد صحیح

$\{0,0,3,4,1,6,0,5,2,0,3,3,2,3,0,2,5,0,3,1,0,3,1,6,1\}$

استفاده از ارتفاع $h = 3$

بافر $b = 4$

B_1, B_2, B_3, B_4 ▪

▪ هر یک $k=4$ مقدار

محاسبه سطح فعال با استفاده از معادله

$$L = L(n) = \max(0, \lfloor \log(n) - 4 \rfloor)$$

در ابتدا،

تعداد مقادیر پردازش شده $n = 0$

پس: شروع از $L = 0$

▪ خواندن اولین $N_1 = 4$ مقدار از جریان ورودی $\{0,0,3,4\}$

- پر کردن یکی از بافرهای خالی، مثلاً B_1

- ظرفیت بافر برابر ۴ در نتیجه عدم نیاز به کنار گذاشتن مقادیر تصادفی و ذخیره همه آنها

B ₁ ⁰				B ₂				B ₃				B ₄			
0	0	3	4												

مثال - ساخت بافرهای نمونه

دوم،

- تعریف دوباره سطح فعال
- تعداد مقادیر پردازش شده تاکنون $n = N_1 = 4$
- $L = L(4) = \max(0, 2 - 4) = 0$
- سطح فعال برابر سطح صفر
- خواندن $N_2 = 4$ مقدار بعدی $\{1, 6, 0, 5\}$
- و پر کردن بافر B_2 به همان ترتیب

B_1^0				B_2^0				B_3				B_4			
0	0	3	4	1	6	0	5								

مثال - ساخت بافرهای نمونه

پردازش $n = N_1 + N_2 = 8$ مقدار تاکنون

سطح فعلی برابر

$$L = L(8) = \max(0, 3 - 4) = 0$$

خواندن $N_3 = 4$ مقدار بعدی $\{2, 0, 3, 3\}$ نمایه‌سازی آنها به بافر B_3

B_1^0				B_2^0				B_3^0				B_4			
0	0	3	4	1	6	0	5	2	0	3	3				

مثال - ساخت بافرهای نمونه

پردازش $n = N_1 + N_2 + N_3 = 12$ مقدار تاکنون

سطح فعال هنوز برابر صفر

خواندن $N_4 = 4$ مقدار بعدی $\{2,3,0,2\}$ و پر کردن آخرین بافر خالی B_2

B_1^0				B_2^0				B_3^0				B_4^0			
0	0	3	4	1	6	0	5	2	0	3	3	2	3	0	2

مثال - ساخت بافرهای نمونه

پیش از این همه بافرها در این مرحله

نیاز به عملیات ادغام

• پایین‌ترین لایه دارای حداقل دو بافر: سطح

انتخاب تصادفی دو بافر از آن سطح

▪ مثلاً، B_2^0 و B_3^0

▪ ادغام و سپس مرتب‌سازی

$$\{1,6,0,5\} \cup \{2,0,3,3\} = \{1,6,0,5,2,0,3,3\} \rightarrow \{0,0,1,2,3,3,5,6\}$$

مثال - ساخت بافرهای نمونه

گام بعدی

- آزاد کردن بافرهای B_2^0 و B_3^0
- پر کردن بافر B_3 در سطح ۱ با ۵۰ درصد از مقادیر قبلی آنها
- جهت سادگی انتخاب مقادیر فرد

B_1^0				B_2				B_3^1				B_4^0			
0	0	3	4					0	1	3	5	2	3	0	2

مثال - ساخت بافرهای نمونه

تاکنون پردازش $n = N_1 + N_2 + N_3 + N_4 = 16$ مقدار

لایه فعال باقی ماندن لایه صفر و پر کردن B_2 با $N_5 = 4$ مقدار بعدی از جریان داده
▪ $\{5,0,3,1\}$

B_1^0				B_2^0				B_3^1				B_4^0			
0	0	3	4	5	0	3	1	0	1	3	5	2	3	0	2

مثال - ساخت بافرهای نمونه

عدم وجود بافر خالی

▪ نیاز به ادغامی دیگر

سطح • شامل سه بافر پر

انتخاب تصادفی دو مورد از آنها

▪ مثلاً، B_1^0 و B_4^0

ادغام مقادیر آنها و در ادامه مرتب‌سازی

$$\{0,0,3,4\} \cup \{2,3,0,2\} = \{0,0,3,4,2,3,0,2\} \rightarrow \{0,0,0,2,2,3,3,4\}$$

مثال - ساخت بافرهای نمونه

برچسب‌زنی بافرهای B_1^0 و B_4^0 به عنوان خالی

پر کردن بافر B_4 در سطح ۱ با ۵۰ درصد از مقادیر آنها با گرفتن مقادیر در موقعیت‌های زوج

B_1				B_2^0				B_3^1				B_4^1			
				5	0	3	1	0	1	3	5	0	2	3	4

مثال - ساخت بافرهای نمونه

گام بعد

$$n = N_1 + N_2 + N_3 + N_4 + N_5 = 20 \text{ تاکنون پردازش}$$

$$L = L(20) = 4.32 - 4 = 1 \text{ سطح فعال برابر}$$

$$N_6 = 4 \cdot 2^1 = 8 \text{ خواندن مقدار بعدی از جریان داده}$$

▪ باقی نماندن مقادیر کافی در جریان داده

▪ خواندن $\{0,3,1,6,1\}$

▪ پرکردن B_1 با نمونه برداری یک مقدار از هر گروه دو مقداری

B_1^1				B_2^0				B_3^1				B_4^1			
3	1	1		5	0	3	1	0	1	3	5	0	2	3	4

پرسش چندک معکوس

امکان پاسخ به پرسش چندک معکوس با استفاده از بافر نمونه

تخمین رتبه مقدار داده شده x به صورت مجموع وزنی لایه از شمارش مقادیر کوچکتر از x برای هر بافر غیر خالی

$$rank(x) = \sum_{i=1}^b 2^{L(B_i)} \cdot |\{e < x : e \in B_i^L(B_i)\}|$$

مثال - پرسش چندک معکوس با نمونه برداری تصادفی

جریان داده از مثال قبلی

پرس و جوی چندک معکوس با تخمین رتبه مقدار ۴

داده ساختار به بافر نمونه

B_1^1				B_2^0				B_3^1				B_4^1			
3	1	1		5	0	3	1	0	1	3	5	0	2	3	4

مثال - پرسش چندک معکوس با نمونه برداری تصادفی

B_1^1				B_2^0				B_3^1				B_4^1			
3	1	1		5	0	3	1	0	1	3	5	0	2	3	4

با استفاده از

$$rank(x) = \sum_{i=1}^b 2^{L(B_i)} \cdot |\{e < x : e \in B_i^L(B_i)\}|$$

محاسبه

$$rank(4) = 2^1 \cdot 3 + 2^0 \cdot 3 + 2^1 \cdot 3 + 2^1 \cdot 3 = 21$$

تخمین رتبه مقدار ۴ رتبه $rank(4) = 21$

پرس وجوی چندک و یافتن Q -چندک

جهت حل پرسش چندکی و یافتن q -چندک
▪ از داده ساختار بافر نمونه

در پی مقداری با بدست آوردن رتبه تخمینی آن حاصل از معادله
▪ نزدیکترین به $q \cdot n$

امکان محاسبه با تعدادی پرس وجوی چندک معکوس برای هر یک از مقادیر در داده ساختار

روش دیگر

- استفاده از جستجوی دودویی برای تسریع فرآیند
- توقف به محض یافتن مقدار به اندازه کافی نزدیک

مثال - پرس و جوی چندک با نمونه برداری تصادفی

جریان داده مثال قبلی

محاسبه چندک ۰,۶۵

تعداد کل اعداد در داده ساختار $n = 25$

▪ بنابراین مقدار مرزی حدود $q \cdot n = 0.65 \cdot 25 = 16.25$

B_1^1				B_2^0				B_3^1				B_4^1			
3	1	1		5	0	3	1	0	1	3	5	0	2	3	4

مثال - پرس و جوی چندک با نمونه برداری تصادفی

وجود مقادیر $\{0,1,2,3,4,5\}$ در بافر نمونه

- آغاز با مقدار

- تخمین رتبه آن

- برابر با صفر $\text{rank}(0) = 0$

سپس، بررسی مقدار ۱

- تخمین رتبه آن $\text{rank}(1) = 5$

رتبه مقدار ۲ برابر $\text{rank}(2) = 12$

رتبه مقدار ۳ برابر $\text{rank}(3) = 14$

از مثال قبل رتبه ۴ برابر $\text{rank}(4) = 21$

رتبه مقدار ۵ برابر $\text{rank}(5) = 23$

مثال - پرس و جوی چندک با نمونه برداری تصادفی

نزدیک ترین مقدار به مقدار مرزی ۱۶,۲۵ مقدار ۳ با $\text{rank}(3) = 14$

گزارش مقدار ۳ به عنوان تقریبی از چندک ۰,۶۵

امکان تسریع فرایند با استفاده از جستجوی دودویی روی مجموعه مرتب شده اعداد،
▪ با در نظر گرفتن اینکه رتبه تابعی یکنواخت از آرگومان خود

ویژگی‌ها

برای محاسبه ϵ -تقریبی q - چندک،

نیاز به مقدار ثابتی از حافظه

▪ متناسب با $b \cdot k$

▪ فقط وابسته به ϵ

با خطای تقریب داده شده، امکان داشتن درخت محاسباتی با ارتفاع $h = \log \frac{1}{\epsilon}$

تعداد بهینه بافرها

$$b = \log \frac{1}{\epsilon} + 1,$$

اندازه هر بافر

$$k = \frac{1}{\epsilon} \sqrt{\log \frac{1}{\epsilon}}$$

ویژگی‌ها

کراندار بودن تصادفی بودن الگوریتم ت تقریب‌های چندک با احتمال خطای ثابت با $\frac{1}{2}\epsilon$
ناشی از نمونه‌برداری تصادفی و مراحل ادغام تصادفی

الگوریتم ت

تعداد کل عملیات ادغام در کل جریان داده $O\left(\frac{n}{k}\right)$
▪ هر بهروزرسانی حدود $O(1)$

مرتب سازی هر بهروزرسانی $O(\log k)$

زمان **سرشکنی** $O(\log k)$

Q-DIGEST

خلاصه چندک

▪ یا q – digest

الگوریتم خلاصه جریان مبتنی بر درخت

نیشیت شریواستاوا، چیرانجیب بوراگوهین، دیویاکانت آگروال و سابهاش سوری

در سال ۱۳۸۳

پیشنهاد جهت استفاده در زمینه نظارت بر داده‌های توزیع شده از حسگرها

تلخیص چ Q-DIGEST

تبدیل یا حل مسئله محاسبه چندک

- به عنوان مسئله هیستوگرام

- تلخیص داده‌ها به تعدادی سطل

نگهداری مجموعه‌ای از چنین سطل‌هایی در داده‌ساختاری درختی

ادغام سطل‌های کوچک و تقسیم سطل‌های بزرگ

الگوریتمی «قطعی» با «اتلاف»

تلخیص چ

اجرای الگوریتم با مقادیر صحیح در بازه‌ای از پیش معین

تقسیم‌بندی دودویی بازه صحیح $[0, N - 1]$

▪ به مثابه درخت دودویی کامل

▪ مقدار ریشه آن شامل کل بازه $[0, N - 1]$

▪ فرزندان چپ و راست آن دارای بازه $\left[0, \left\lfloor \frac{N-1}{2} \right\rfloor\right]$ و $\left[\left\lfloor \frac{N-1}{2} \right\rfloor + 1, N - 1\right]$

گره‌های برگ نمایش دهنده مقادیر صحیح منفرد

عمق درخت $\log N$

تلخیص چ

هر گره v از درخت دارای

- بازه متناظر $[v_{کم}, v_{بیش}]$
- شمارنده متناظر شمارنده v
- نمایش تعداد اعضای سطل (شامل تکراری ها)

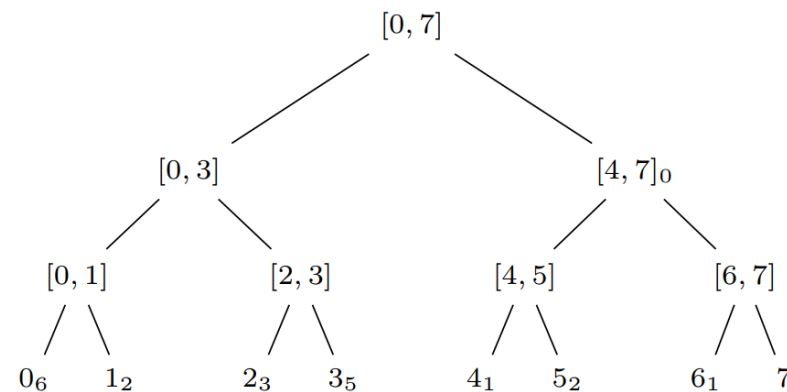
مثال - تقسیم‌بندی دودویی در تلخیص چ

مجموعه داده‌ای از $n = 20$ عدد صحیح از محدوده $[0, 7]$:
 $\{0, 0, 3, 4, 1, 6, 0, 5, 2, 0, 3, 3, 2, 3, 0, 2, 5, 0, 3, 1\}$

ایجاد درخت دودویی با تقسیم‌بندی دودویی بازه

سطل‌بندی داده‌های ورودی

▪ هر برگ شامل مقدار و بسامد



تلخیص چ

گره‌های داخلی داده شامل بسامدهای مقادیر ذخیره شده متناظر

در بدترین حالت ذخیره‌سازی

▪ ذخیره داده‌ها به صورت $O(n)$ یا $O(N)$

▪ هر کدام کوچکتر

احتمال بالای تنگی و نامتوازن‌ی چنین درخت‌های دودویی

▪ $\Leftarrow$ ناکارآمدی ذخیره آن به صورت خام بدون فشردن‌سازی

تلخیص چ

پیشنهاد روش فشرده‌سازی و ذخیره فشرده چنین درخت تقسیم دودویی

داده‌ساختار تلخیص چ

▪ کدگذاری اطلاعات مربوط به توزیع مقادیر

▪ نمایش نسخه‌ای از درخت دودویی شامل سطل‌های v با ویژگی زیر

$$\begin{cases} v_{count} \leq \lfloor n\sigma \rfloor, \text{ به جز سطل‌های برگ,} \\ v_{count} + v_{count}^p + v_{count}^s > \left\lfloor \frac{n}{\sigma} \right\rfloor, \text{ به جز ریشه,} \end{cases}$$

▪ v^p والد و v^s همزاد v

▪ n تعداد کل عناصر

▪ $\sigma \in [1, n]$ پارامتر طراحی و مسئول سطح فشرده‌سازی

تلخیص چ

سطل‌های ریشه و برگ
▪ دارای استثنا

گنجاندن ریشه در داده‌ساختار

گنجاندن سطل‌های برگ با شمارنده‌های بزرگتر از مقدار مرزی $\left\lfloor \frac{n}{\sigma} \right\rfloor$ (مقادیر پربسامد)

ویژگی تلخیص سبک سنگینی بین گنجاندن چند سطل سطح بالا و گسترده، و بسیاری از سطل‌های کوچک حاوی اطلاعاتی در مورد چند مقدار کم‌بسامد

تلخیص چ

ویژگی نخست

$$v_{count} \leq [n\sigma]$$

- حذف کننده سطرها مگر گره‌های برگ با شمارنده‌های مقادیر با بسامد بالا
- به دلیل نیاز به دقت بیشتر و ارزش ذخیره‌چنین سطرها برای ارزش ذخیره مقادیر فرزند و داشتن شمارنده‌های دقیق‌تر

ویژگی دوم

$$v_{count} + v_{count}^p + v_{count}^s > \left\lfloor \frac{n}{\sigma} \right\rfloor$$

- در صورت همزادی دو سطر دارای شمارنده‌های کوچک
- در پی اجتناب از داشتن دو شمارنده جداگانه برای هر یک از آنها
- توصیه به ادغام آنها در والد خود
- منجر به نائل شدن به درجه فشردگی مورد نیاز

تلخیص چ

سخن کوتاه

- ایجاد داده ساختار نیازمند ادغام سلسله مراتبی و کاهش سطرها
- گذر از تمام سطرها از پایین به بالا
- بررسی نقض ویژگی در هر یک

در عمل بررسی محدودیت دوم

- به دلیل گذر پائین به بالا

تلخیص چ

به جز سطل ریشه،

هر سطل v ناقض ویژگی تلخیص

- فشردسازی زیردرخت آن
- از طریق فشردسازی شمارنده‌های آن
- ادغام والد v^p و همزاد v^s
- ارتقاء آنها به سطل والد

$$v_{count}^p = v_{count}^p + v_{count} + v_{count}^s$$

حذف سطل v و همزاد آن v^s از تلخیص

تلخیص چ - فشرده سازی

الگوریتم ۴ - فشرده سازی تلخیص چ

ورودی: داده ساختار q -digest بازه $[0, N - 1]$

ورودی: ضریب فشرده سازی σ

خروجی: داده ساختار q -digest فشرده شده

$$1 - \log N \leftarrow \text{سطح}$$

تأزمائی که $(\text{سطح} > 0)$ انجام یده

برای $v \in Q - DIGEST[\text{سطح}]$ انجام یده

$$\text{اگر } v_{count} + v_{count}^p + v_{count}^s \leq \left\lfloor \frac{n}{\sigma} \right\rfloor \text{ آنگاه}$$

$$v_{count}^p \leftarrow v_{count}^p + v_{count} + v_{count}^s$$

$$Q - DIGEST \leftarrow Q - DIGEST \setminus \{v, v^s\}$$

$$1 - \text{سطح} \leftarrow \text{سطح}$$

برگرداندن Q-DIGEST

تلخیص چ-فشرده سازی

زمان فشرده سازی $O(m \cdot \log N)$

▪ $m = |Q - DIGEST|$ تعداد سطرها در داده ساختار

هزینه به روزرسانی به ازای هر مقدار حدود $O(\log N)$

در عمل، مصرف زمان بیشتر در به روزرسانی

▪ به دلیل درج هر مقدار جدید در گره برگ

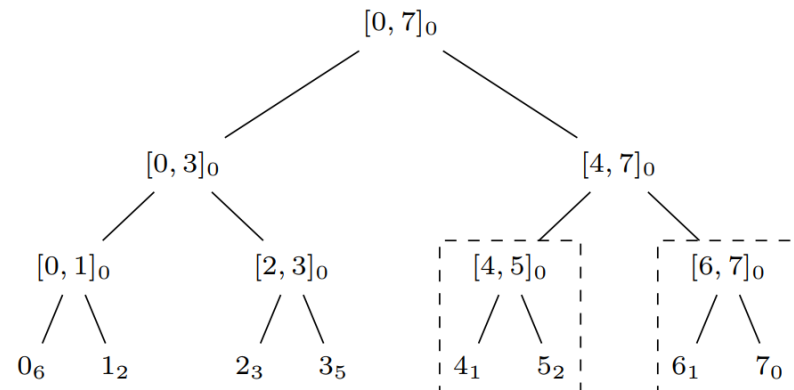
▪ سپس یافتن موقعیت اصلی در طول عملیات فشرده سازی

یافتن موقعیت مناسب خود در تلخیص چ با حرکت تک گامی مقدار در هر زمان

مثال - فشرده سازی درخت با تلخیص چ

مجموعه داده با $n = 20$ مقدار از مثال قبل

بسامدها برای سطلهای مشاهده نشده به طور پیش فرض برابر صفر



مثال - فشرده سازی درخت با تلخیص چ

فشرده سازی با $\sigma = 5$

مقدار مرزی برابر

$$\left\lfloor \frac{n}{\sigma} \right\rfloor = \left\lfloor \frac{20}{5} \right\rfloor = 4$$

حرکت از پایین به بالا،

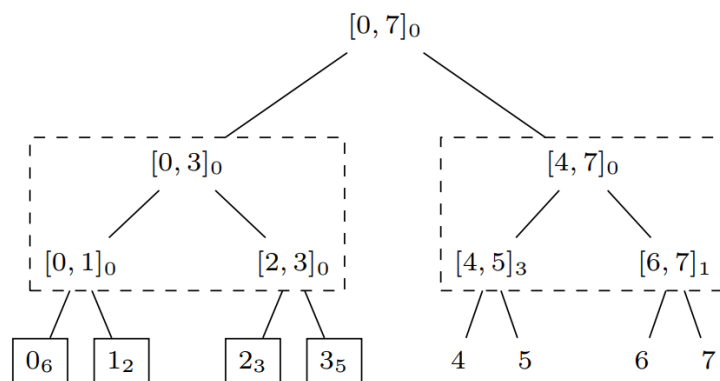
ابتدا سطح چهارم

- برآورده شدن شرط دوم تلخیص در سطلهای ۰ تا ۳
- نقض ویژگی مذکور در فرزندان سطلهای [۵, ۴] و [۷, ۶]
- ادغام در والدین خود
- حذف از داده ساختار

مثال - فشرده سازی درخت با تلخیص چ

داده ساختار پس از اعمال مراحل

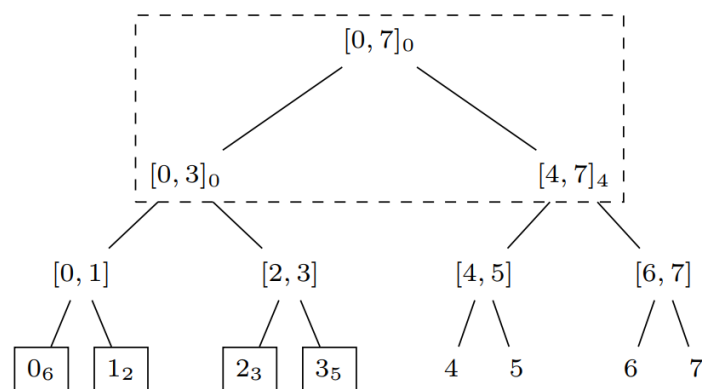
▪ قرارگیری سطل‌های موجود در جعبه‌ها در داده ساختار فشرده



مثال - فشرده‌سازی درخت با تلخیص چ

سطح سوم،

- نقض ویژگی‌های تلخیص در همه سطرها
- فشرده کردن آنها در والدین خود و حذف از داده‌ساختار



مثال - فشرده سازی درخت با تلخیص چ

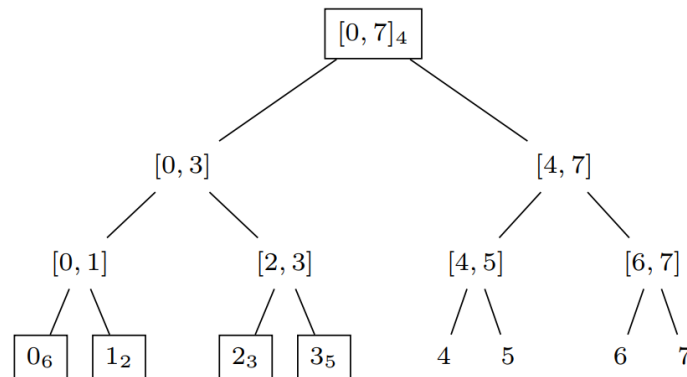
سطح دوم،

- نقض ویژگی تلخیص در دو فرزند سطل ریشه
- عدم تجاوز تعداد کل آنها از مقدار مرزی
- در نتیجه، نیاز به ادغام در والد

عدم نیاز به بررسی ریشه

- گنجانده شدن ریشه در داده ساختار اگر شمارش های غیر صفر مرتبط داشته باشد.

نسخه نهایی داده ساختار فشرده



مثال - فشردہ سازی درخت با تلخیص چ

داده ساختار نیازمند ذخیره فقط پنج سطل با شمارش های غیر صفر

تلخیص چ-فشرده سازی

حرکت از پایین به بالا

یکبار تصمیم گیری درباره ادغام سطرها در طول رویه

عدم لزوم رعایت ویژگی تلخیص در همه سطرها چ فشرده پس از فشرده سازی

تغییرات (به عنوان مثال، ادغام در والد) در برخی از سطرها در سطوح بالای درخت

- امکان نقض ویژگی با سطرها قبلاً گنجانده شده

- در عمل، عدم کاهش دقت با این رفتار

- در بدترین حالت، تولید داده ساختاری با حدی کمتر از مطلوب بودن

- مصرف حافظه بیشتر نسبت به نظر

الگوریتم تلخیص چ

الگوریتم ۵- الگوریتم تلخیص چ

ورودی: مجموعه داده D با مقادیری از بازه $[0, N - 1]$

ورودی: ضریب فشرده سازی σ

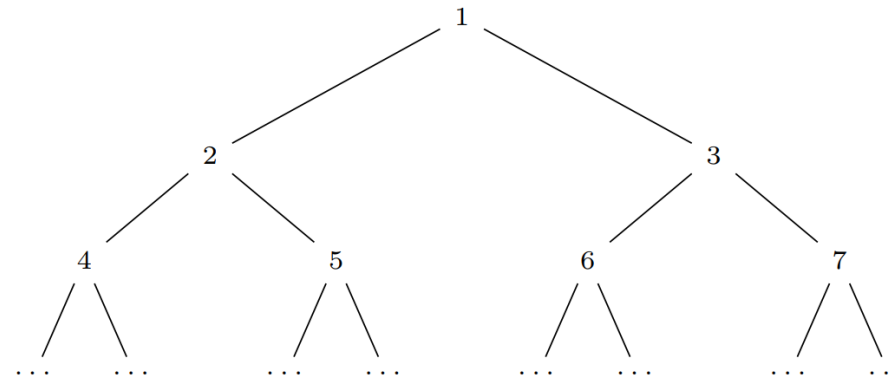
خروجی: داده ساختار q-digest فشرده

$Q - DIGEST \leftarrow BinaryPartitionTree(D, [0, N - 1])$
برگرداندن $Compress(Q - DIGEST, N, \sigma)$

الگوریتم تلخیص چ

بهینه‌سازی نمایش داده‌ساختار

شماره‌گذاری سطرها در درخت دودویی مرتبط با روشی از راست به چپ، از بالا به پایین



الگوریتم تلخیص چ

شماره گذاری همه سطرها

▪ سادگی بازیابی بازه متناظر $[v_{کم}, v_{بیش}]$ با وجود اطلاع صرف از نمایه و اندیس i متناظر

الگوریتم ۶- بازیابی محدوده سطر $[v_{کم}, v_{بیش}]$

ورودی: نمایه سطر i

خروجی: بازه سطر

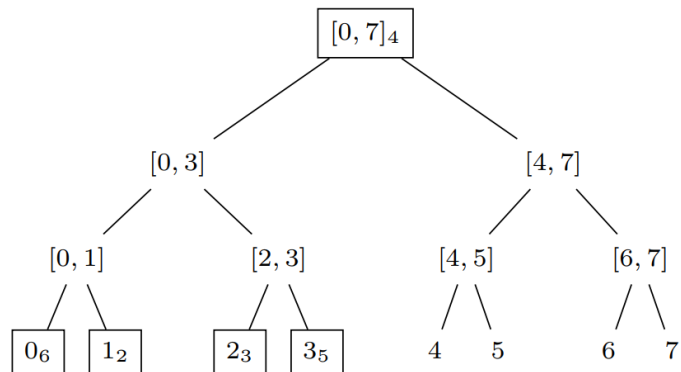
$\left\lfloor \log(i) \right\rfloor \leftarrow \text{سطح}$

$n \leftarrow 2^{\text{سطح}}$ // تعداد سطرها در سطح

$m \leftarrow i - n$ // موقعیت سطر در سطح

برگرداندن $\left\lceil \frac{(N+1)m}{n} \right\rceil, \left\lceil \frac{(N+1)(m+1)}{n} \right\rceil - 1$

نمایش خطی از داده ساختار

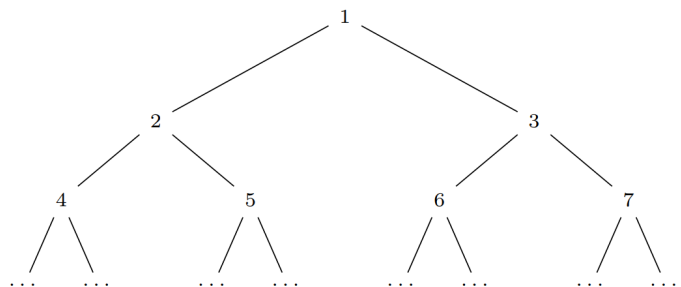


امکان ایجاد نمایش خطی از داده ساختار

- آرایه‌ای از سطرها،
- هر سطر یک دو-تایی از شماره و شمارنده‌های مرتبط است.

مثال،

- نمایش خطی داده ساختار فشرده مثال قبل
- $\langle (1,4), (8,6), (9,2), (10,3), (11,5) \rangle$

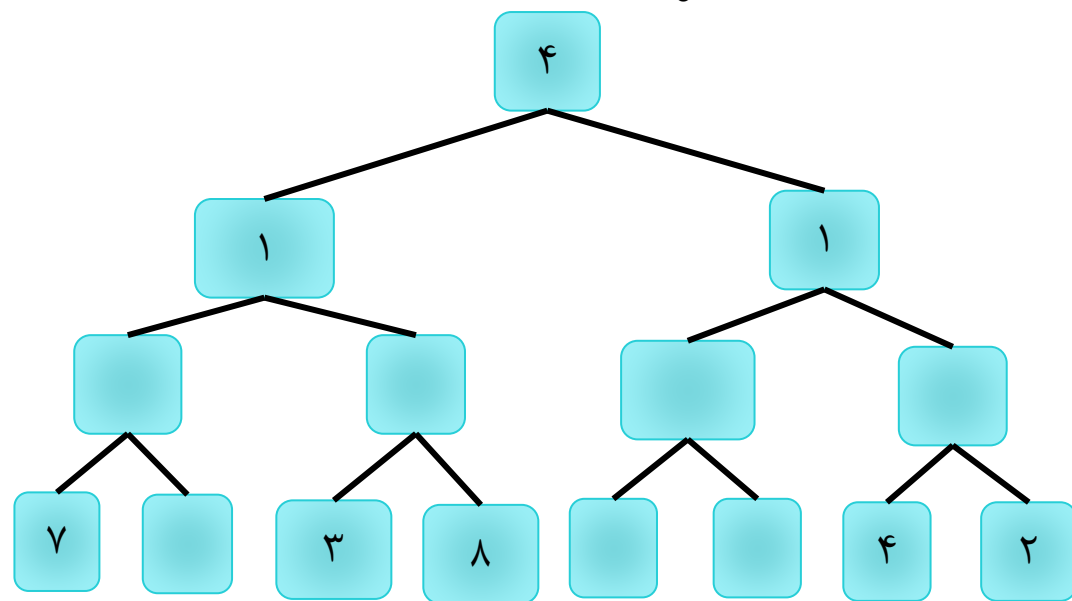


ادغام دو تلخیص چ

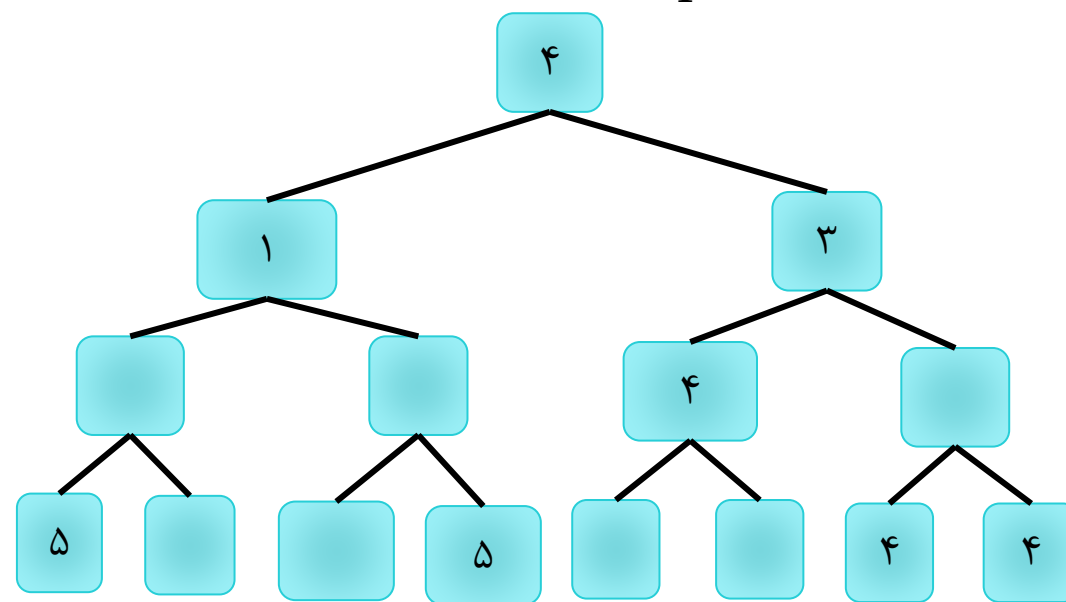
امكان ادغام دو تلخیص با ضریب فشردہ سازی σ و محدودہ های مقدار یکسان
فراہم سازی امکان پردازش توزیعی جریان های داده بزرگ
محاسبه اتحاد مجموعه های سطل های ذخیره شده آنها
شمارش های سطل های با محدودہ یکسان
جمع تعداد کل مقادیر
اجرای الگوریتم فشردہ سازی

مثال-ادغام دو تلخیص چ

$$n_0 = 30, \sigma_0 = 6, \frac{n_0}{\sigma_0} = 5$$

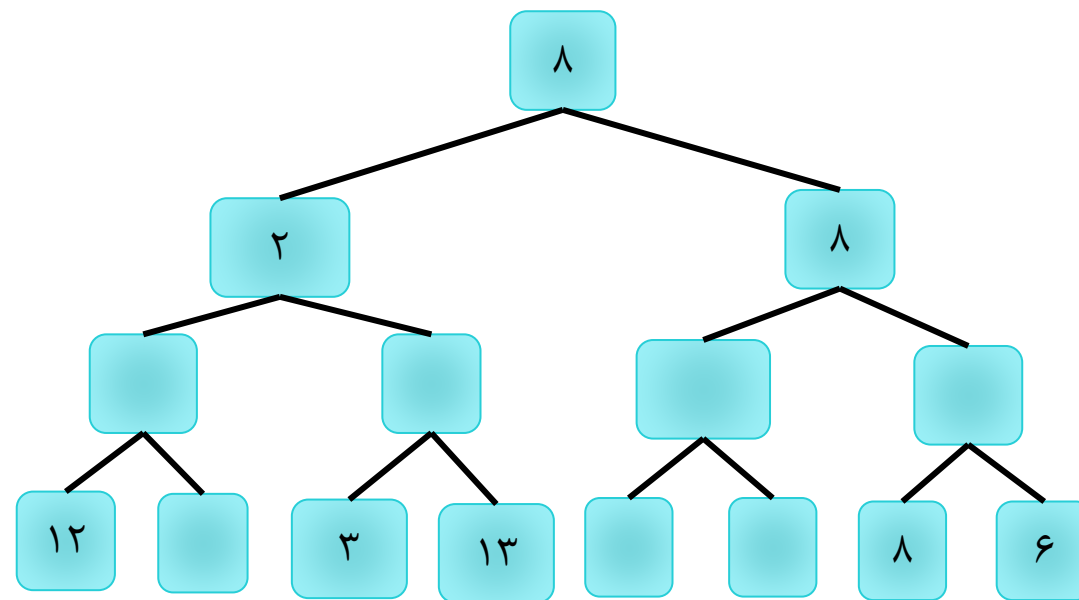
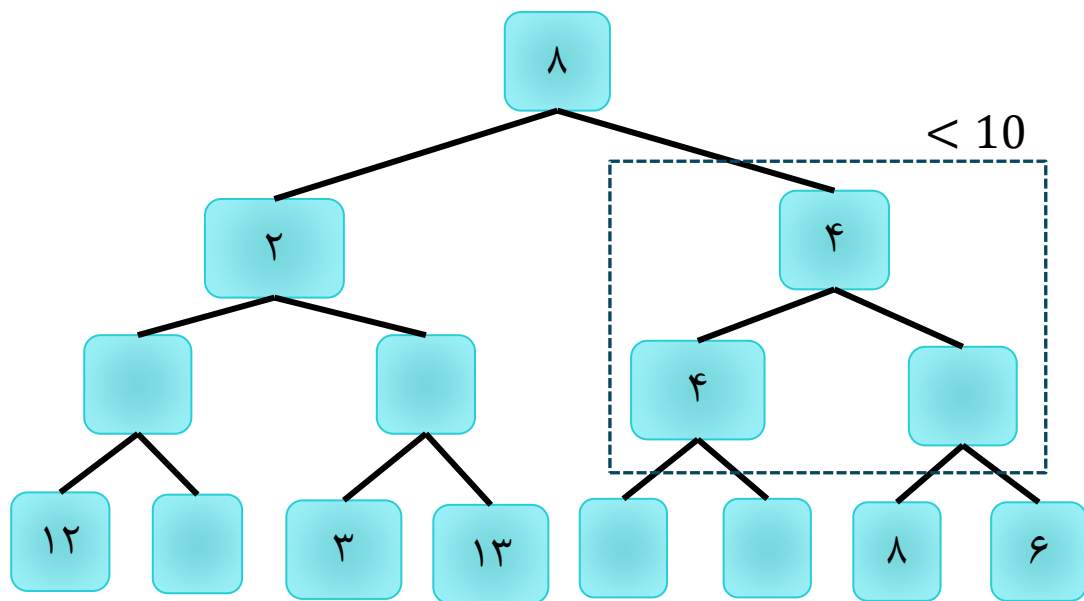


$$n_1 = 30, \sigma_1 = 6, \frac{n_1}{\sigma_1} = 5$$



مثال-ادغام دو تلخیص چ

$$n = n_0 + n_1 = 60, \sigma = 6, \frac{n}{\sigma} = 10$$



تلخیص چ -

استفاده از الگوریتم تلخیص چ در پاسخ به پرس و جوی چندک و یافتن q - چندک دنباله مرتب شده S با مرتب سازی سطرها به ترتیب صعودی مقادیر بیش v آنها هنگام تساوی: مقادیر کوچکتر اول

پویش دنباله S از ابتدا

افزودن شمارنده های سطرها دیده شده

شرط پایان: بزرگتر شدن مقدار سطر v^* مجموع از qn ،

▪ تخمین رتبه سطر

بیش v^* به عنوان تخمین q - چندک

تلخیص چ - پرسش چندک

الگوریتم ۷- پاسخ به پرسوئوهای چندک با تلخیص چ

ورودی: داده ساختار q-digest

ورودی: مقدار $q \in [0, 1]$

خروجی: q - چندک

$S \leftarrow \text{sort}(Q - \text{DIGEST})$

$\text{rank} \leftarrow 0$

برای $(v, \text{count}) \in S$ انجام بده

$\text{rank} \leftarrow \text{rank} + \text{count}$

اگر $\text{rank} \geq q \cdot n$ آنگاه

برگرداندن $v_{\max}$

تلخیص چ - پرسش چندک

وجود حداقل $q \cdot n$ سطل با مقدار حداکثر کمتر از بیش v^*

- در نتیجه رتبه سطل v^* حداقل $q \cdot n$

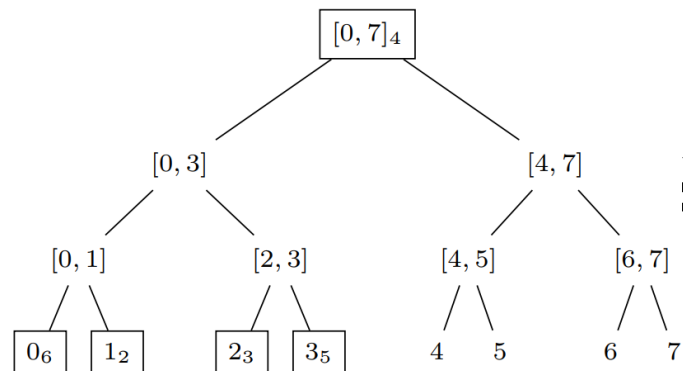
امکان خطا در محاسبه تقریب ϵ -چندک q

- در صورت وجود مقادیر کمتر از بیش v^* در اجداد سطل v^*
- به دلیل شمرده نشدن در الگوریتم γ

کران چنین خطایی با $\epsilon \cdot n$

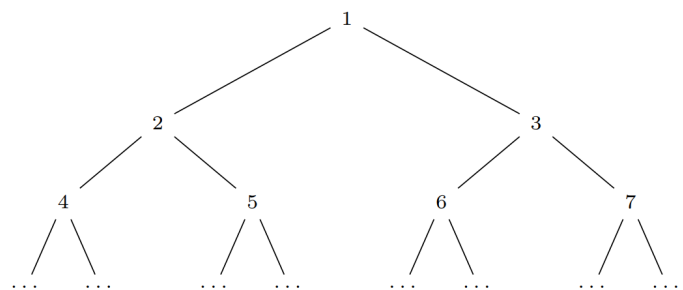
گزارش رتبه‌ای در بازه $[q \cdot n, (q + \epsilon) \cdot n]$ در الگوریتم

عدم تخمین مقدار دقیق q - چندک کمتر از حد تخمین



مثال - پرس وجوی چندک با تلخه

پرس وجوی چندک با محاسبه چندک ۰,۶۵ از داده ساختار تلخیص چ مثال قبلی با نمایش خطی



$\langle (1,4), (8,6), (9,2), (10,3), (11,5) \rangle$

بنابراین، دنباله مرتب شده سطرها

$S = \langle (8,6), (9,2), (10,3), (11,5), (1,4) \rangle$

مثال - پرس وجوی چندک با تلخیص چ

با شروع از ابتدا

جمع شمارش‌های سطل‌ها تا هنگام بزرگتر شدن مجموع از $0.65 \times n = 13$

فراتر رفتن تلخیص چ از مقدار مرزی در سطل (۱۱،۵)

▪ مربوط به مقدار برگ ۳

تخمین چندک ۰,۶۵ (یا صدک شصت و پنجم) مجموعه داده مثال قبل برابر مقدار ۳

تلخیص چ- پرسش چندک معکوس

به روشی مشابه

ایجاد دنباله مرتب شده S از سطرها

گذر از آنها با نگهداری مجموع جاری شمارش‌های سطرها مشاهده شده

رتبه مقدار داده شده x

▪ مجموع شمارش‌های سطرها v با بیش $x > v$

تلخیص چ - پرسش چندک معکوس

الگوریتم ۸- پاسخ به پرسوئوهای چندک معکوس با تلخیص چ

ورودی: مقدار x

ورودی: داده ساختار q -digest

خروجی: رتبه عنصر

$S \leftarrow \text{sort}(Q - \text{DIGEST})$

$\leftarrow$ رتبه

برای $(v, \text{count}) \in S$ انجام بده

اگر $x > v_{\max}$ آنگاه

$\text{count} + \text{رتبه} \leftarrow \text{رتبه}$

برگرداندن رتبه

تلخیص چ - پرسش چندک بازه

الگوریتم ۹ - پاسخ به پرسوئوهای بازه با تلخیص چ

ورودی: بازه $[a, b]$

ورودی: داده ساختار q-digest

خروجی: تعداد مقادیر در محدوده

$r_a \leftarrow \text{InverseQuantileQuery}(a, Q - \text{DIGEST})$

$r_b \leftarrow \text{InverseQuantileQuery}(b, Q - \text{DIGEST})$

برگرداندن $r_b - r_a$

تلخیص چ - پرسش چندک معکوس

رتبه حاصل در بازه $[rank(x), rank(x) + \epsilon \cdot n]$

پاسخ به پرسوجوی بازه

- کفایت محاسبه دو پرسوجوی چندک معکوس برای یافتن رتبه‌ها و تفاوت بین مرزهای محدوده a و b
- حداکثر خطا برای پرسوجوی محدوده در تلخیص چ $2\epsilon \cdot n$

ویژگی‌ها

الگوریتم با ائتلاف

فشرده‌سازی اطلاعات مربوط به مقادیر با بسامد پایین

نگهداری اطلاعات مربوط به مقادیر با بسامد بالا

طرح تقریب خوب

▪ با رخ دادن تغییرات گسترده در بسامد مقادیر مختلف

الگوریتم عرضه اطلاعاتی در مورد توزیع مقادیر مقادیر

▪ بدون در اختیار نهادی اطلاعاتی در مورد مکان رخ دادن مقادیر

ویژگی‌ها

سبک سنگینی بین دقت الگوریتم و حافظه مورد نیاز برای ذخیره داده‌ساختار

کنترل با ضریب فشرده‌سازی σ

برای محدوده داده شده $[0, N]$ ،

- انتظار ذخیره حداکثر $3 \cdot \sigma$ سطل
- خطا در محاسبه تقریب ϵ -چندک- q با کران بالای زیر

$$\epsilon \leq \frac{\log N}{\sigma}$$

ویژگی‌ها

ویژگی اصلی q-digest

- تطبیق یافتن با داده‌ها
- ایجاد سطل‌هایی با وزن تقریباً مساوی

در مقابل هیستوگرام سنتی، q-digest

- تلخیص چ اجازه همپوشانی سطل‌ها
- پاسخ به پرس‌وجوهای اجماع (به عنوان مثال، مقادیر پربسامد)

مشکلات اصلی کاربردهای عملی

- مدیریت مقادیر عدد صحیح
- نیاز به دانستن بازه آنها از قبل
- خطاهای بزرگ برای چندک‌های خیلی بزرگ

تلخیص ت T-DIGEST

از روش‌های انباشت دقیق آمار مبتنی بر رتبه به صورت برخط

تد دانینگ و اوتمار ارتل در سال ۱۳۹۳

امکان تخمین چندک در جریان‌های نامحدود با تمرکز بر مقادیر خیلی بزرگ مانند چندک ۰,۹۹

باز بودن تحقیق و وجود به‌روزرسانی‌های جدید

تلخیص ت

- تلخیص جریان داده ورودی D در خوشه‌هایی با اندازه متغیر $\{C_i\}_{i=1}^m$ منجر به حفظ مناسب دقت در محاسبه چندک در حین پردازش حجم زیادی از داده
- هر خوشه C_i نمایش زیرمجموعه‌ای از مقادیر ورودی
- با اندازه‌ای جهت اطمینان حاصل
 - از بزرگ نبودن برای تخمین چندک‌ها با درونیابی
 - خیلی کوچک نبودن برای جلوگیری از داشتن خوشه‌های زیاد

تلخیص ت

- هر خوشه C_i با مرکز جرم c_i ،
 - نقطه داده در مرکز خوشه،
 - میانگین مقادیر ورودی
 - c_i^{count} تعداد عناصر

داده ساختار تلخیص ت آرایه ای از مراکز جرم $\{(c_1, c_1^{count}), (c_2, c_2^{count}), \dots, (c_m, c_m^{count})\}$ مرتب شده با ترتیبی صعودی

- امکان تخمین حداکثر مقدار چندک متناظر با هر مرکز جرم c_i

$$q(c_i) = \frac{1}{n} \left(\sum_{j < i} c_j^{count} + c_i^{count} \right)$$
$$q(c_i) = \frac{1}{n} \left(\sum_{j < i} c_j^{count} + \frac{1}{r} \cdot c_i^{count} \right)$$

▪ $n = \sum_{j=1}^m c_j^{count}$ تعداد کل مقادیر نمایه شده در داده ساختار

تلخیص ت

هر خوشه C_i در داده‌ساختار مسئول بازه خاصی از مقادیر چندک $[q(c_{i-1}), q(c_i)]$

وابستگی طول آن به اندازه خوشه

تعیین اندازه صحیح خوشه تأثیر مستقیم بر دقت

انجام با تابع مقیاس غیرنزولی

تابع $k = k(q, \sigma)$ در نظر گرفتن فشرده‌سازی مطلوب σ و مقادیر چندک q بر اساس فاصله از انتها
مانند $q = 0$ و $q = 1$

اهمیت انتخاب خوب تابع مقیاس

وجود توابع متفاوت از نظر دقت

مثال، از توابع رایج

$$k(q, \sigma) = \frac{\sigma}{2} q$$

$$k(q, \sigma) = \frac{\sigma}{2\pi} \arcsin(2q - 1)$$

پارامتر فشرده‌سازی $\sigma > 1$

مقادیر بزرگتر مربوط به فشرده‌سازی کمتر

تلخیص ت

با توجه به تابع مقیاس انتخاب شده $k = k(q, \sigma)$,

مرتبط بودن هر خوشه C_i با مرکز جرم c_i در داده ساختار

تعریف اندازه k (یا $K(c_i)$) با طول مقیاس گذاری شده بازه چندک برای خوشه

$$K(c_i) = k(q(c_i), \sigma) - k(q(c_{i-1}), \sigma), i = 2 \dots m$$

$$K(c_1) = k(q(c_1), \sigma) \quad \blacksquare$$

تلخیص ت

برای محدود کردن تعداد مقادیر در خوشه

▪ امکان محدودسازی اندازه k

▪ با توابع مقیاس غیرخطی،

▪ امکان وجود خوشه‌های غیریکنواخت، با اندازه‌های خوشه بزرگتر برای چندک‌های محدوده میانی و کوچکتر در نزدیکی انتها

▪ تا خوشه‌های تکی صرفاً حاوی تک مقدار

طراحی الگوریتم برای ساخت داده‌ساختار کاملاً ادغام‌پذیر

▪ همه خوشه‌ها $\{C_i\}_{i=1}^m$ برآورده‌ساز «ویژگی تلخیص»

$$\begin{cases} K(c_i) \leq 1, \text{ به جز خوشه‌های تکی} \\ K(c_i) + K(c_{i+1}) > 1 \end{cases}$$

محدودساز اندازه k هر خوشه

تضمین عدم امکان ادغام بیشتر هر دو خوشه مجاور

تلخیص ت

در عمل، جهت برآورده سازی محدودیت های ویژگی تلخیص

▪ عدم نیاز به محاسبه مجدد اندازه k برای هر خوشه در هر تغییر

▪ به دلیل مرتب بودن همه مراکز جرم

▪ امکان محدود کردن تعداد مقادیر در خوشه G_i با انتخاب مرزی برای مقدار چندک حداکثر تخمینی

$$q_{limit} = k^{-1}(k(q(c_i), \sigma) + 1, \sigma)$$

▪ مقدار آن برای تابع مقیاس قبلی

$$q_{limit} = \frac{1}{2} \left[1 + \sin \left(\arcsin(2 \cdot q(c_i) - 1) + \frac{2\pi}{\sigma} \right) \right]$$

تلخیص

با داشتن قوانین محدود کردن تعداد مقادیر در هر خوشه

- امکان تعریف الگوریتم ادغام
- مشابه رویه خوشه‌بندی منظم

خلاصه کردن دنباله ورودی از نقاط داده وزنی $\{(x_1, x_1^{count}), (x_2, x_2^{count}), \dots, (x_b, x_b^{count})\}$

- مرتب‌سازی آنها همراه با همه مراکز جرم از داده‌ساختار
- با گذری یکباره از دنباله حاصل X ، ادغام آنها در صورت عدم نقض ویژگی تلخیص
- آغاز با چپ‌ترین مرکز جرم
- در نظر گرفتن خوشه متناظر به عنوان خوشه نامزد فعلی
- محاسبه مقدار مرزی آن q_{limit}
- تخمین مقادیر چندک تقریبی آنها با پردازش متوالی همه مراکز جرم از X
- مقایسه با مقدار مرزی
- در صورت عدم تجاوز میانگین جذب مرکز جرم از مقدار مرزی، ادغام آن در خوشه نامزد
- ادامه با مرکز جرم بعدی از دنباله X
- در غیر این صورت، به معنای رسیدن به حداکثر ظرفیت خوشه نامزد
- عدم امکان افزودن مقادیر جدید
- ذخیره خوشه نامزد فعلی در داده‌ساختار
- انتشار خوشه نامزد جدید با مرکز جرم میانگین
- محاسبه دوباره مقدار مرزی چندک q_{limit}
- در پایان، حاصل داده‌ساختار کاملاً ادغام شده

تلخیص ت

الگوریتم ۱۰-۱ ادغام مقادیر در تلخیص ت

ورودی: بافر B با مقادیر $\{(x_1, x_1^{count}), (x_2, x_2^{count}), \dots, (x_b, x_b^{count})\}$

ورودی: داده ساختار T-DIGEST

ورودی: پارامتر فشردگی $\sigma > 1$ ، تابع مقیاس k

خروجی: داده ساختار t-digest با بافر ادغام شده

$X \leftarrow \text{sort}(T - \text{DIGEST} \cup B)$

$T - \text{DIGEST} \leftarrow \emptyset$

$m \leftarrow \text{count}(X), n \leftarrow \sum_{x_i \in X} x_i^{count}$

$c \leftarrow x_1, q_c \leftarrow \cdot$

$q_{limit} \leftarrow k^{-1}(k(q_c, \sigma) + 1, \sigma)$

برای i از ۲ تا m انجام بده

$\hat{q} \leftarrow q_c + \frac{1}{n} c^{count} + \frac{1}{n} \cdot x_i^{count}$

اگر $\hat{q} \leq q_{limit}$ آنگاه

$c^{count} \leftarrow c^{count} + x_i^{count}$

$c \leftarrow c + x_i^{count} \cdot \frac{x_i - c}{c^{count} + x_i^{count}}$

continue

$T - \text{DIGEST} \leftarrow T - \text{DIGEST} \cup \{(c, c^{count})\}$

$q_c \leftarrow q_c + \frac{1}{n} c^{count}$

$q_{limit} \leftarrow k^{-1}(k(q_c, \sigma) + 1, \sigma)$

$c \leftarrow x_i$

$T - \text{DIGEST} \leftarrow T - \text{DIGEST} \cup \{(c, c^{count})\}$ // آخرین خوشه معمولاً تکی

برگرداندن T-DIGEST

تلخیص ت

با هر بار انتشار خوشه نامزد جدید

- نیاز به محاسبه دوباره مقدار مرزی $qlimit$
- شامل محاسبه پرهزینه تابع مقیاس و معکوس آن

کم بودن تعداد خوشه‌ها در عمل

معرفی تکنیک‌های مختلف برای بهینه‌سازی تخمین مرزی

- مانند استفاده از تقریب‌های کارآمد برای اجزای توابع مقیاس یا تخمین تقریبی حداکثر تعداد عناصری که می‌توان در هر خوشه خلاصه کرد.

تلخیص ت

الگوریتم کامل برای نمایه‌سازی جریان داده پیوسته بر اساس پردازش داده‌های جریانی با بافرهایی اندازه ثابت و ادغام مداوم آنها در داده‌ساختار

تلخیص ت

الگوریتم ۱۱- اندیس کردن جریان داده با تلخیص ت

ورودی: جریان داده $D = \{x_1, x_2, \dots\}$

ورودی: اندازه بافر b ، پارامتر فشرده‌سازی $\sigma > 1$ ، تابع مقیاس k

ورودی: داده‌ساختار t -digest

$T - DIGEST \leftarrow \emptyset$

تازمانی که D انجام بده

$B \leftarrow \{(x_1, 1), (x_2, 1), \dots, (x_b, 1)\}$

$T - DIGEST \leftarrow Merge(T - DIGEST, B, \sigma, k)$

برگرداندن T-DIGEST

تلخیص ت

تقسیم هزینه‌های زمان اجرای الگوریتم بافر-و-ادغام ۱۱ بین درج‌های مکرر مقادیر ورودی در بافر و فراخوانی‌های نادر الگوریتم ۱۰

درج‌ها ارزان

- هزینه‌های کلی تحت سلطه مرتب‌سازی و فراخوانی‌های تابع مقیاس در زیرالگوریتم ادغام
- سرشکنی در چندین درج

مثال - نمایه کردن جریان داده با تلخیص ت

مجموعه داده با $n = 20$

$\{0,0,3,4,1,6,0,5,2,0,3,3,2,3,0,2,5,0,3,1\}$

پارامتر فشرده سازی $\sigma = 5$ ، اندازه بافر $b = 10$ و تابع مقیاس مطابق تعریف قبلی

$$k(q, \sigma) = \frac{\sigma}{2\pi} \arcsin(2q - 1)$$

مثال - نمایه کردن جریان داده با تلخیص ت

پُر کردن بافر B با ده مقدار اول از ورودی

$$B = \langle (0,1), (0,1), (3,1), (4,1), (1,1), (6,1), (0,1), (5,1), (2,1), (0,1) \rangle$$

▪ با توجه به الگوریتم ۱۱، نیاز به ترکیب مقادیر بافر و مراکز جرم موجود در داده ساختار

▪ داده ساختار t-digest تهی: فهرست نامزدهای مراکز جرم X فقط شامل $n = 10$ مقدار از B

▪ مرتب سازی صعودی آنها

$$X = \langle (0,1), (0,1), (0,1), (0,1), (1,1), (2,1), (3,1), (4,1), (5,1), (6,1) \rangle,$$

انتخاب اولین خوشه نامزد با چپ ترین مرکز جرم $(0,1)$

▪ محاسبه مقدار مرزی چندک آن q_{limit}

$$q_{limit} = \frac{1}{2} \left[1 + \sin \left(\arcsin(-1) + \frac{2\pi}{5} \right) \right] = 0.34549$$

مثال - نمایه کردن جریان داده با تلخیص ت

دریافت مقدار بعدی از X

▪ دوباره $(0,1)$

▪ ارزیابی امکان ادغام آن بدون نقض ویژگی تلخیص در خوشه نامزد

▪ برای تخمین حداکثر مقدار چندک $\hat{q}$ این خوشه ادغام شده نیاز به

$$\hat{q} = \frac{1+1}{10} = 0.2,$$

▪ کمتر از مقدار مرزی چندک q_{limit}

▪ امکان ادغام مقدار $(0,1)$ در خوشه نامزد فعلی

▪ اکنون دارای $c^{count} = 2$ مقدار

▪ به دلیل یکسانی مقادیر، ثابت باقی ماندن مرکز جرم ثابت $c = 0$

به طور مشابه، امکان ادغام مقدار بعدی $(0,1)$ در خوشه نامزد

▪ صرفاً تغییر تعداد مقادیر خلاصه شده به $c^{count} = 3$

مثال - نمایه کردن جریان داده با تلخیص ت

دریافت مقدار چهارم از X

▪ دوباره $(0,1)$

▪ ارزیابی امکان ادغام آن بدون نقض ویژگی تلخیص در خوشه نامزد

▪ برای تخمین حداکثر مقدار چندک $\hat{q}$ این خوشه ادغام شده نیاز به

$$\hat{q} = \frac{3 + 1}{10} = 0.4,$$

▪ بیشتر از مقدار مرزی $q_{limit} = 0.34549$

▪ عدم امکان ادغام مقدار $(0,1)$ در خوشه نامزد فعلی

▪ توقف تلاش برای جذب سایر مراکز جرم در خوشه

▪ ذخیره آن در داده ساختار

$$T - DIGEST = \langle (0,3) \rangle,$$

▪ نگهداری حداکثر مقدار چندک آن به عنوان $q = \frac{c^{count}}{n} = 0.3$

مثال - نمایه کردن جریان داده با تلخیص ت

شروع به ایجاد خوشه نامزد جدید از مقدار در انتظار فعلی $(0,1)$ ، با $c = 0$ و $c^{count} = 1$ ، و مقدار مرزی چندک آن

$$q_{limit} = \frac{1}{2} \left[1 + \sin \left(\arcsin (2 \times 0.3 - 1) + \frac{2\pi}{5} \right) \right] = 0.874025$$

دریافت مقدار $(1,1)$

▪ در صورت ادغام آن در خوشه نامزد فعلی، حداکثر چندک تخمینی $\hat{q}$

$$\hat{q} = 0.3 + \frac{1+1}{10} = 0.5,$$

▪ عدم تجاوز از مقدار مرزی فعلی

▪ ادغام مقدار $(1,1)$ در خوشه فعلی ادغام

▪ افزایش تعداد آن به $c^{count} = 2$ و بدل به مرکز جرم

$$c = 0 + \frac{1-0}{2} = 0.5$$

مثال - نمایه کردن جریان داده با تلخیص ت

به طور مشابه، پردازش بقیه مقادیر از X پردازش می‌کنیم و رشد داده‌ساختار به صورت زیر

$$T - DIGEST = \langle (0,3), (2,5), (5,1), (6,1) \rangle$$

با ادامه پردازش مجموعه داده، ایجاد بافر جدید پر

$$B = \langle (3,1), (3,1), (2,1), (3,1), (0,1), (2,1), (5,1), (0,1), (3,1), (1,1) \rangle,$$

پیوند آن با داده‌ساختار

▪ مرتب‌سازی صعودی بر اساس مراکز جرم دنباله حاصل X

$$X = \langle (0,3), (0,1), (0,1), (1,1), (2,5), (2,1), (2,1), (3,1), (3,1), (3,1), (3,1), (5,1), (5,1), (6,1) \rangle$$

تخیله داده‌ساختار

▪ آغاز از چپ‌ترین مقادیر و ادغام مقادیر به طور متوالی و ذخیره خوشه‌های در داده‌ساختار

▪ ادامه تا عدم امکان ادغام بیشتر

مثال - نمایه کردن جریان داده با تلخیص ت

داده ساختار حاصل دارای $m = 5$ خوشه:

$$T - DIGEST = \langle (0.1667,6),(2.36364,11),(5,1),(5,1),(6,1) \rangle$$

تلخیص ت-پرسش چندک

- به دلیل از دست دادن اطلاعات دقیق در مورد مقادیر خوشه‌بندی شده
- به جز خوشه‌های تکی، که در آن مرکز جرم مقدار اولیه
- داده‌ساختار نمایشی با اتلاف از جریان داده
- نیاز به درونیابی برای تخمین چندک‌ها و پاسخ به پرس‌وجوی چندک،
- با در نظر گرفتن توزیع خوشه‌هایی که توسط تابع مقیاس انتخاب شده تولید شده‌اند

تلخیص ت

الگوریتم ۱۲- پاسخ به پرس وجوهای چندک با تلخیص ت

ورودی: داده ساختار t -digest با m خوشه

ورودی: مقدار $q \in [0, 1]$

خروجی: q - چندک

$$n \leftarrow \sum_{j=1}^m c_j^{count}$$

اگر $1 \cdot q < n \cdot q$ آنگاه

برگرداندن c_1

اگر $n \cdot q > n - \frac{1}{m} c_m^{count}$ آنگاه

برگرداندن c_m

// یافتن i به طوری که $q(c_i) + \frac{1}{n} c_{i+1}^{count} > q$ $\exists i \in [1, m)$

// $\Leftarrow$ چندک جستجو شده جایی بین c_i و c_{i+1}

اگر $1 == c_i^{count} > q(c_i)$ آنگاه

برگرداندن c_i

اگر $1 == c_{i+1}^{count}$ و $q \leq q(c_{i+1}) + \frac{1}{n}$ آنگاه

برگرداندن c_{i+1}

$$\Delta_{چپ} \leftarrow (c_i^{count} == 1) ? 1 : 0$$

$$\Delta_{راست} \leftarrow (c_{i+1}^{count} == 1) ? 1 : 0$$

$$w_{چپ} \leftarrow n \cdot q - n \cdot q(c_i) + \frac{c_i^{count} - \Delta_{چپ}}{2}$$

$$w_{راست} \leftarrow n \cdot q(c_{i+1}) - n \cdot q + \frac{c_{i+1}^{count} - \Delta_{راست}}{2}$$

$$\text{برگرداندن } \frac{c_i \cdot w_{راست} + c_{i+1} \cdot w_{چپ}}{w_{چپ} + w_{راست}}$$

تلخیص ت - پرسش چندک

یافتن q - چندک

- محاسبه رتبه مقدار جستجو شده x در این دنباله مرتب شده
- $n \cdot q$
- n تعداد کل مقادیر خلاصه شده در داده ساختار
- اگر رتبه کمتر از یک گزارش مرکز جرم C_1 به عنوان چندک
- اگر رتبه در نیمه آخرین خوشه آخر از حداکثر تعداد یا حتی بالاتر، گزارش C_m به عنوان حداکثر مقدار در تلخیص

تلخیص ت - پرسش چندک

در غیر این صورت،

▪ یافتن خوشه‌های C_i و C_{i+1} با مقادیر چندک تخمینی با معادله احاطه کننده مقدار چندک داده شده q

اگر خوشه چپ C_i تکی و حداکثر چندک آن بزرگتر از q

▪ برگرداندن مرکز جرم C_i به عنوان چندک

▪ مقدار خلاصه‌شده در خوشه

به طور مشابه، اگر خوشه راست C_{i+1} تکی و مقدار چندک حداقل تخمینی آن کوچکتر از q

▪ این خوشه مسئول چندک جستجو شده

▪ گزارش مرکز جرم C_{i+1} به عنوان بهترین تخمین برای آن

در غیر این صورت،

▪ محاسبه وزن‌ها با ارزیابی سهم هر یک از این خوشه‌ها

▪ محاسبه میانگین وزنی مراکز جرم از هر دو خوشه،

▪ ایجاد درونیابی جهت گزارش به عنوان q - چندک

تلخیص ت- پرسش چندک

وابستگی الگوریتم تخمین چندک به انتخاب تابع مقیاس

با توابع «توزیع شده تر؟» تولیدکننده دم بزرگتری از خوشه‌های تکی در دو سر طیف
▪ بهبود دقت برای چندک‌های دو سر طیف

علاوه بر این،

▪ توصیه بر ذخیره حداقل و حداکثر مقادیر در طول نمایه‌سازی و استفاده در درونیابی

مثال - پرس وجوی چندک با تلخیص ت

پرس وجوی چندک برای محاسبه چندک ۰,۶۵ از داده ساختار حاصل قبلی

$$T - DIGEST = \langle (0.1667, 6), (2.36364, 11), (5, 1), (5, 1), (6, 1) \rangle$$

مثال - پرس و جوی چندک با تلخیص ت

تعداد کل مقادیر مجموع شمارش‌های همه خوشه‌ها

▪ در مورد مثال $n = 20$

▪ رتبه چندک جستجو شده x برابر با $13 = 20 \cdot 0.65 = n \cdot q$

▪ نه کمتر از یک و نه خیلی نزدیک به n ،

▪ جستجو برای دو خوشه متوالی C_i و C_{i+1}

▪ مراکز جرم آنها احاطه کننده چندک جستجو شده

در داده‌ساختار فعلی، برابر C_2 و C_3

▪ زیرا

$$q(c_2) + \frac{1}{2 \times 20} c_3^{count} = \frac{6+11}{20} + \frac{1}{40} = 0.875 > 0.65$$

مثال - پرس و جوی چندک با تلخیص ت

خوشه c_3 تکی
▪ $c_3^{count} = 1$

▪ بررسی حداکثر مقدار چندک آن با مقدار چندک جستجو شده q
▪ جهت تعیین بهترین تناسب بودن مرکز جرم آن

$$q(c_3) + \frac{1}{20} = \frac{6 + 11 + 1}{20} + \frac{1}{20} = 0.95,$$

به طور قابل توجهی بزرگتر از مقدار $q = 0.65$

چندک واقعی در جایی بین مراکز جرم c_2 و c_3
▪ اما خوشه راست با تنظیم $\Delta_{right} = 1$ تکی

مثال - پرس وجوی چندک با تلخیص ت

بنابراین، تخمین چندک جستجو شده با درونیابی با استفاده از میانگین وزنی مراکز جرم

$$w_{\text{چپ}} = 20 \cdot 0.65 - 20 \cdot q(c_2) + \frac{c_2^{\text{count}} - 0}{2} = 13 - 20 \cdot \frac{6+11}{20} + \frac{11}{2} = 1.5, \quad \blacksquare$$

$$w_{\text{راست}} = 20 \cdot q(c_3) - 20 \cdot 0.65 + \frac{c_3^{\text{count}} - 1}{2} = 20 \cdot \frac{6+11+1}{20} - 13 + \frac{1-1}{2} = 5. \quad \blacksquare$$

مثال - پرس و جوی چندک با تلخیص ت

در نهایت، چندک ۰,۶۵ تخمینی برای مجموعه داده

$$x = \frac{c_2 \cdot w_{\text{راست}} + c_3 \cdot w_{\text{چپ}}}{w_{\text{چپ}} + w_{\text{راست}}} = \frac{2.36364 \times 5 + 5 \times 1.5}{1.5 + 5} = 2.97$$

▪ بسیار نزدیک به مقدار دقیق ۳ برای آن مجموعه داده

پرس و جوی چندک معکوس

یافتن رتبه مقدار داده شده x

مقایسه مقدار x با حداقل و حداکثر مراکز جرم در داده ساختار

- به ترتیب چپ‌ترین و راست‌ترین خوشه‌ها

اگر x خارج از آن بازه

- بسته به سمت ظهور مقدار x خروجی: یک یا تعداد کل مقادیر n برابر مقدار رتبه تخمینی

در غیر این صورت، به دنبال مقدار x در میان مراکز جرم

- در صورت یافتن چنین خوشه‌هایی

- جمع کردن شمارنده‌های آنها

- گزارش رتبه (x) به عنوان رتبه خوشه با کوچکترین نمایه با مقدار تنظیم شده

پرس و جوی چندک معکوس

ناموفق بودن بررسی‌های قبل

- اطمینان از قرار گرفتن مقدار x بین مراکز جرم خوشه‌های متوالی، مثلاً (c_i, c_{i+1}) قرار می‌گیرد و رتبه آن حداقل $n \cdot q(c_i)$
- اگر هر دو خوشه تکی، به این معنی که مراکز جرم آنها دقیقاً مقادیر ورودی
- عدم نیاز به به تصحیح آن مقدار و برگرداندن آن به عنوان رتبه جستجو شده
- زمانی که فقط یکی از آن خوشه‌ها تکی باشد، رتبه را با سهم مقیاس‌گذاری شده خوشه دیگر تنظیم می‌کنیم تا مقدار نهایی را به دست آوریم.
- در غیر این صورت، با استفاده از اندازه هر دو خوشه، ایجاد درونیابی و محاسبه مقدار رتبه

الگوریتم ۱۲- پاسخ به پرس و جوهای چندک معکوس با t-digest

ورودی: مقدار x

ورودی: داده ساختار t-digest با m خوشه

خروجی: رتبه عنصر

$$n \leftarrow \sum_{j=1}^m c_j^{count}$$

اگر $x < c_1$ آنگاه برگرداندن ۱

اگر $x > c_m$ آنگاه برگرداندن n

// بررسی مرکز جرم بودن x

اگر $\exists j: c_j == x$ آنگاه

$$J \leftarrow \{j: c_j == x\}, i^* \leftarrow \min(J)$$

$$\text{برگرداندن } n \cdot q(c_{i^*}) - c_{i^*}^{count} + \frac{1}{\gamma} \sum_{j \in J} c_j^{count}$$

// در این مرحله، می توانیم مطمئن باشیم که $\exists i \in [1, m): x \in (c_i, c_{i+1})$

$$\text{رتبه} \leftarrow n \cdot q(c_i)$$

اگر $c_i^{count} == 1$ و $c_{i+1}^{count} == 1$ آنگاه

برگرداندن رتبه

$$\hat{x} \leftarrow \frac{x - c_i}{c_{i+1} - c_i}$$

اگر $c_i^{count} == 1$ آنگاه

$$\text{برگرداندن } \text{رتبه} + \frac{\hat{x}}{\gamma} \cdot c_{i+1}^{count}$$

اگر $c_{i+1}^{count} == 1$ آنگاه

$$\text{برگرداندن } \text{رتبه} - \frac{1 - \hat{x}}{\gamma} \cdot c_i^{count}$$

$$\text{برگرداندن } \text{رتبه} + \frac{\hat{x}}{\gamma} \cdot c_{i+1}^{count} - \frac{1 - \hat{x}}{\gamma} \cdot c_i^{count}$$

پرس وجوی بازه

شبيه تلخيص چ

- انجام دو پرس وجوی چندک معکوس
- یافتن تفاوت بين رتبه‌های مرزهای بازه

ویژگی‌ها

سبک-سنگینی بین اندازه داده‌ساختار (کنترل شونده با پارامتر فشرده‌سازی σ) و سرعت و دقت تخمین چندک‌ها امکان دستیابی به سرعت بیشتر با استفاده از مقدار ثابت حافظه

- با انتخاب مقدار کوچکتر σ و اندازه بافر بزرگ b

بالاترین دقت،

- ترجیح استفاده از σ بزرگتر برای فشرده‌سازی کمتر و بافر بزرگتر (به عنوان مثال، $10 \times \sigma$)

کوچکترین حافظه

- ترجیح استفاده از بافر کوچکتر و مقادیر بزرگتر پارامتر فشرده‌سازی σ

ویژگی‌ها

بازهٔ تعداد خوشه‌ها (m) در داده‌ساختار و با نمایه‌سازی $n \geq \frac{\sigma}{2}$ مقدار و برآورده‌سازی ویژگی تلخیص

$$\left\lfloor \frac{\sigma}{2} \right\rfloor \leq m \leq \lfloor \sigma \rfloor$$

مثال - تخمین فضای مورد نیاز

در پی نمایه‌سازی حداقل $n = 1000$ مقدار با پارامتر فشرده‌سازی $\sigma = 100$

- انتظار ۵۰ تا ۱۰۰ خوشه کاملاً ادغام شده

در داده‌ساختار

- نمایش هر خوشه با مرکز جرم و تعداد مقادیر نمایه شده
- با داشتن شمارنده‌های ۳۲ بیتی و عدد ممیز شناور با دقت دو برابر ۶۴ بیتی برای مقدار مرکز جرم
- کل مرکز جرم نیازمند ۱۲ بایت حافظه
- جا شدن کل داده‌ساختار در حدود ۱,۲ کیلوبایت حافظه

مثال - تخمین فضای مورد نیاز

برای دقت بالا،

- استفاده از بافری ده برابر بزرگتر از پارامترهای فشرده‌سازی

- با داشتن $b = 10 \cdot \sigma = 1000$

- تخصیص شمارنده‌های ۱۶ بیتی کوچکتر و اعداد ممیز شناور ۶۴ بیتی

- در زمان اجرا به ۱۰ کیلوبایت حافظه اضافی

ویژگی‌ها

دستیابی به دقت ϵ در تخمین q - چندک
▪ متناسب با $q \cdot (1 - q)$

سایر الگوریتم‌ها صرفاً در پی کنترل خطای مطلق ثابت

تلخیص ت در پی کنترل خطای نسبی

▪ مقاوم کردن الگوریتم در برابر خطاهای قابل توجه برای چندک‌های انتهایی طیف

مزیت تلخیص ت به تلخیص چ

▪ مدیریت مقادیر ممیز شناور در مقابل محدود بودن به اعداد صحیح

ویژگی‌ها

ادغام دو داده‌ساختار تلخیص‌ت با استفاده از الگوریتم ادغام

- عدم پیوستگی داده‌ساختار حاصل

- اما تولید تخمین خوبی بر اساس نتایج تجربی

امکان موازی‌سازی برای بخش‌های مختلف جریان داده

- استفاده در پاسخ به پرس‌وجوهای رتبه از آنها

- مفید در عملیات MapReduce و استخراج جریان برای برنامه‌های داده بزرگ