

قواعد تلازمی

مفاهیم و الگوریتمها



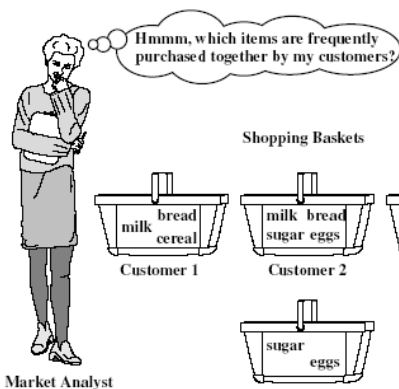
سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

1

قواعد تلازمی



- تحلیل وابستگیها یک حالت غیر نظارتی داده کاوی می باشد که به جستجو برای یافتن ارتباط در مجموعه داده ها می پردازد. به عبارتی دیگر تحلیل وابستگیها مطالعه ویژگیها یا خصوصیات می باشد که با یکدیگر همراه هستند. این روش در سال 1993 توسط Agrawal مطرح شد.



- این الگوریتم در وهله اول در تحلیل سبد خرید و اینکه چقدر اقلام خریداری شده بوسیله مشتریان با هم ارتباط دارند، مورد استفاده قرار گرفت.

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

2

تعاریف و مفاهیم اصلی



قوانین وابستگی به شکل اگر و آنگاه به همراه معیارهای پشتیبان و اطمینان مربوط به قوانین می باشند.

$I = \{ i1 , i2 , \dots , im \}$: مجموعه ای از کل ایتهمهای خریداری شده است

T : هر زیرمجموعه ای از I می باشد که از آن بعنوان تراکنش یاد می کنیم.

D : مجموعه تراکنشهای موجود در T است

TID : شناسه منحصر به فرد و یکتایی است که به هریک از تراکنشها اختصاص می یابد.

نمای کلی یک قانون وابستگی به فرم زیر می باشد:

$X \Rightarrow Y [support , Confidence]$

می باشد بطوریکه داریم:

$$X \subset I, Y \subset I, X \cap Y = \emptyset$$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

3

تعاریف و مفاهیم اصلی



● **پشتیبان یا Support**: نشان دهنده درصد یا تعداد مجموعه تراکنشهایی در D است که شامل هر دوی X و Y ($X \cup Y$) باشند.

● **اطمینان یا Confidence**: میزان وابستگی یک قلم کالای خاص را به دیگری بیان می کند و مطابق فرمول زیر محاسبه می شود:

$$confidence = support(X \cup Y) / support(X)$$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

4

مثال



مجموعه ای از کل اقلام خریداری شده در آمده است.

$I = \{\text{خیار، جعفری، پیاز، گوجه فرنگی، نمک، نان، زیتون، پنیر، کره}\}$

• مجموعه D که شامل تک تک تراکتهای و خریدهایی می باشد، به فرم زیر است:

• $D = \{T1, T2, T3, T4, T5, T6, T7, T8\}$

• $T1$ {جعفری، پیاز، زیتون، خیار، گوجه فرنگی}،

• $T2$ {جعفری، خیار، گوجه فرنگی}،

• $T3$ {نان، نمک، گوجه فرنگی، پیاز، جعفری، خیار}،

• $T4$ {نان، پیاز، خیار، گوجه فرنگی}

• $T5$ {پیاز، نمک، گوجه فرنگی}،

• $T6$ {پنیر، نان}،

• $T7$ {خیار، پنیر، گوجه فرنگی}،

• $T8$ {کره، نان}

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

5

مثال



فرض کنیم قانون وابستگی به فرم زیر را داریم:

$$X \Rightarrow Y [\text{support, Confidence}]$$

$\{ \text{خیار، گوجه فرنگی} \} \Rightarrow \{ \text{پیاز، جعفری} \}$

$X = \{ \text{گوجه فرنگی، خیار} \}$

$Y = \{ \text{جعفری، پیاز} \}$

$$X \Rightarrow Y = \{ \text{گوجه فرنگی، خیار، جعفری، پیاز} \} = \{T1, T3\}$$

$$\text{Support}(X \Rightarrow Y) = 2/8 = 0.25$$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

6

مثال



زیرا مجموعه XUY ، 2 عضو و مجموعه D ، 8 عضو دارد و بدین معنی است که الگوی خرید "گوجه فرنگی، خیار، جعفری، پیاز" در 25% کل سبد خرید رخ می دهد.

$$X = \{T3, T4, T7, T2, T1\}$$

یعنی {خیار، گوجه فرنگی} در $T3, T4, T7, T2, T1$ خریداری شده اند. بنابراین:

$$\text{Support}(x) = 8/5 = 0.62$$

$$\text{Confidence} = \text{support}(X \Rightarrow Y) / \text{support}(X)$$

$$= (2/8) / (5/8) = 2/5 = 0.40$$

یعنی هنگامی که افراد "خیار و گوجه فرنگی" را خریداری می کنند، در 40% اوقات، "جعفری و پیاز" را هم می خرند.

الگوریتم Apriori

```

 $L_1 = \{\text{large 1-itemsets}\}$ 
For (  $k = 2; L_{k-1} \neq \phi; k++$  ) do begin
     $C_k = \text{apriori-gen}(L_{k-1});$ 
    forall transactions  $t \in D$  do begin
         $C_t = \text{subset}(C_k, t)$ 
        forall candidates  $c \in C_t$  do
             $c.\text{count}++;$ 
        end
    end
     $L_k = \{c \in C_k | c.\text{count} \geq \text{minsup}\}$ 
end
Answer =  $\bigcup_k L_k;$ 

```

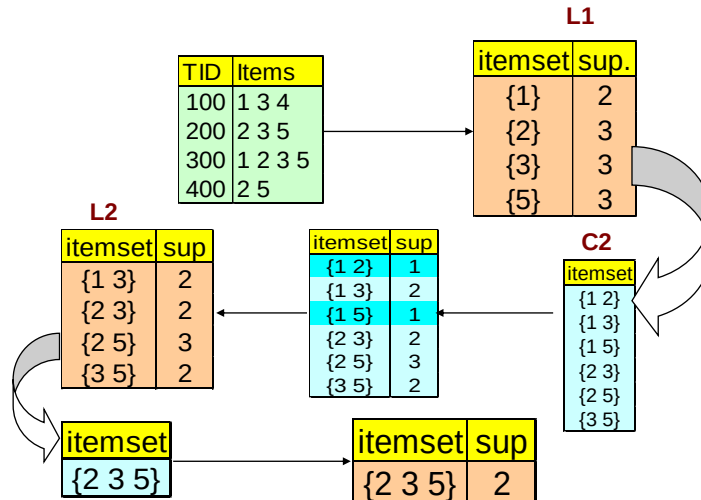
بزرگترین آئتمها را در نظر می گیرد

CK های جدید را مطابق الگوریتم بیان شده تولید می کند

Support هر یک از آئتمهای **CK** را محاسبه می کند

Support فقط آئتمهایی که از حداقل **LK** قرار می دهد بزرگتر یا مساویند را در

الگوریتم Apriori مثال:



سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

9

الگوریتم Apriori

مطابق الگوریتم برای محاسبه C3 بصورت زیر عمل می‌کنیم:

$$L_2 = \{ \{1,3\}, \{2,3\}, \{2,5\}, \{3,4\} \}$$

برای محاسبه C3 مطابق این الگوریتم فقط آنهایی که مولفه اول برابر دارند یعنی $\{2,3\}, \{2,5\}$ هستند و از آنجائیکه برای محاسبه تمامی Ck ها باید قانون لکسیکوگرافیک رعایت شود تنها حالت ممکن $\{2,3,5\}$ می‌باشد. بنابراین

$$C3 = \{2,3,5\}$$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

10

استخراج قوانین



- پروسه استخراج قوانین وابستگی
- ابتدا همه frequent itemsets ها را که دارای حداقل support هستند بیابید.
- برای تمامی frequent itemset ها
 - همه زیرمجموعه‌های آنها را استخراج کنید
 - همه قوانین ممکن را استخراج کنید
 - قوانینی را بپذیرید که از حداقل confidence برخوردار هستند.
- آیتها در هر مجموعه‌ای مطابق با قانون لکسیکوگرافیک چیده می‌شوند به عنوان مثال
- اگر $LK = \{a[1], a[2], \dots, a[k]\}$ باشد ، مطابق این قانون باید رابطه زیر برقرار باشد.

$$a[1] < a[2] < \dots < a[k]$$

مثال دوم



- minimum support $s=30\%$,
minimum confidence $c= 60\%$

Transaction ID	Items Purchased
1	{آب پرتقال, لیموناد}
2	{شیر, آب پرتقال, شیشه پاک کن}
3	{آب پرتقال, پاک کننده, لیموناد}
4	{شیشه پاک کن, لیموناد}
5	{لیموناد, چیپس}



Candidate 1-itemsets (C1)

1-itemset	Support
آب پرتقال	60%
لیموناد	80%
شیر	20%
شیشه پاک کن	40%
پاک کننده	20%
چیپس	20%



Large 1-itemsets (L1)

1-itemset	Support
آب پرتقال	60%
لیموناد	80%
شیشه پاک کن	40%

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

13



Candidate 2-itemsets (C2)

2-itemset	Support
{آب پرتقال, لیموناد}	40%
{آب پرتقال, شیشه پاک کن}	20%
{لیموناد, شیشه پاک کن}	20%



Large 2-itemsets (L2)

2-itemset	Support
{آب پرتقال, لیموناد}	40%

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

14



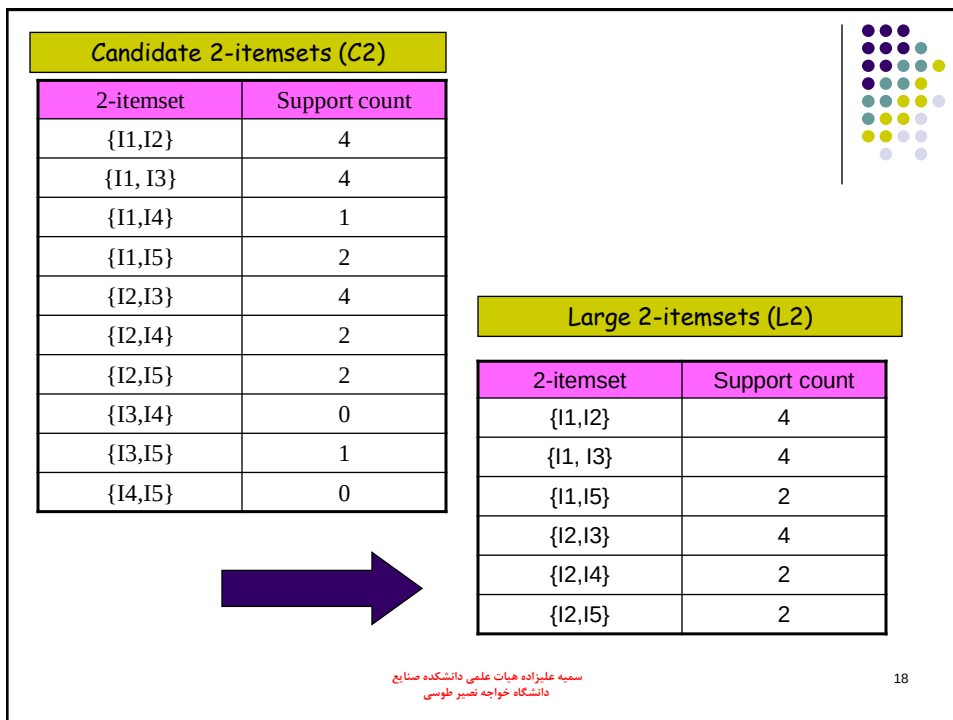
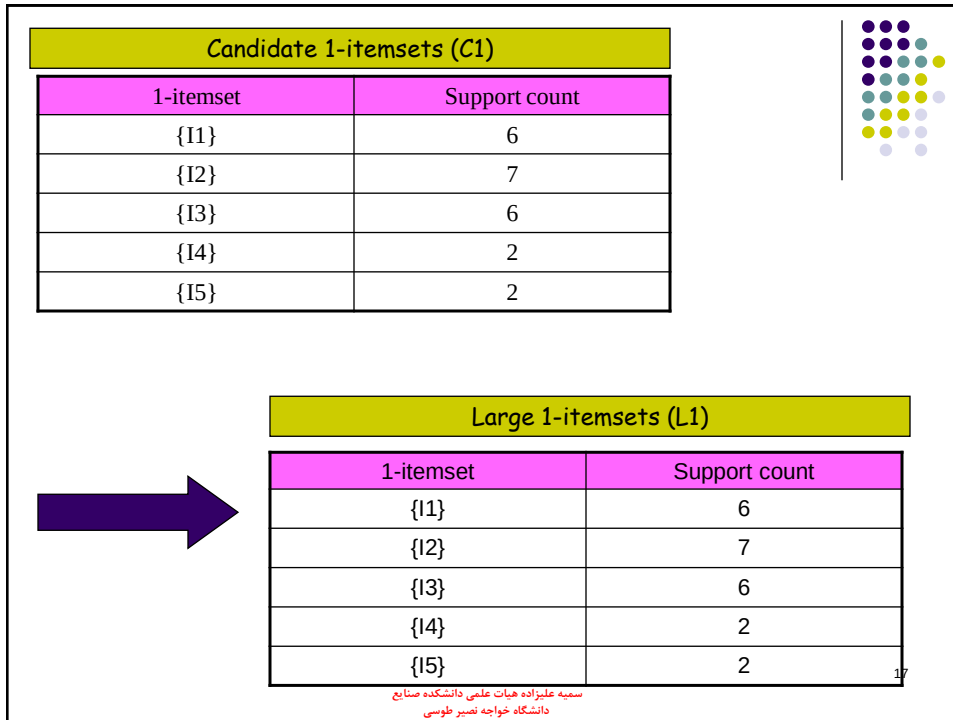
- آب پرتقال $\rightarrow$ لیموناد (confidence = 66.67%)
- لیموناد $\rightarrow$ آب پرتقال (confidence = 50%)

 • آب پرتقال $\rightarrow$ لیموناد (confidence = 66.67%)

مثال سوم

- minimum support count is 2, ($s=2/9=22\%$)

TID	List of Item_IDs
T100	I1, I2, I5
T200	I2, I4
T300	I2, I3
T400	I1, I2, I4
T500	I1, I3
T600	I2, I3
T700	I1, I3
T800	I1, I2, I3, I5
T900	I1, I2, I3





Candidate 3-itemsets (C3)

3-itemset	Support count
{I1,I2,I3}	2
{I1, I2,I5}	2

Large 3-itemsets (L3)

3-itemset	Support count
{I1,I2,I3}	2
{I1, I2,I5}	2



الگوریتم AprioriTid

- همانگونه که قبلاً نیز ذکر شد الگوریتم Apriori همه پایگاه داده را در هر گذر می پیماید تا Support ها را محاسبه کند و پیمودن همه پایگاه داده ممکن است در همه فازها مورد نیاز نباشد. بر مبنای این مشکل، الگوریتم دیگری بنام AprioriTid ابداع شد. این الگوریتم نیز روشی مشابه با الگوریتم Apriori را برای محاسبه CK ها در هر فاز بکار می برد.
- تفاوت عمده ای که این الگوریتم با الگوریتم Apriori دارد در این است که این الگوریتم کل پایگاه داده را برای محاسبه Support بعد از مرحله اول نمی پیماید و از مجموعه C^k برای محاسبه Support استفاده می کند.

الگوریتم AprioriTid

$L_1 = \{large\ 1\text{-itemsets}\}$
 $C_1^{\wedge} = database\ D;$
 For $(k = 2; L_{k-1} \neq \phi; k++)$ do begin
 $C_k = \text{apriori-gen}(L_{k-1});$
 $C_k^{\wedge} = \phi;$
 for all entries $t \in C_{k-1}^{\wedge}$ do begin
 $C_t = \{c \in C_k | (c - c[k] \in t.set - of - items$
 $\wedge (c - c[k - 1]) \in t.set - of - items\};$
 for all candidates $c \in C_t$ do
 $c.count ++;$
 if $(C_t \neq \phi)$ then $C_k^{\wedge} += \langle t.TID, C_t \rangle;$
 end
 end
 $L_k = \{c \in C_k | c.count \geq minsup\}$
 end
 Answer = $\bigcup_k L_k;$

آتمهای با بیشترین فراوانی را محاسبه می کند
 مقداری اولیه $C_1^{\wedge}$ با پایگاه داده موجود
 C_k های جدید را مطابق الگوریتم بیان شده تولید می کند
 یک انبار جدید ایجاد می کند
 itemsetsهایی را در نظر بگیرید که قبلاً وجود داشته اند
 Support را محاسبه می کند
 موجودیتهای خالی را حذف می کند
 فقط آنهایی که از حداقل Support بزرگتر یا مساوند را در L_k قرار می دهد

سمیه علیزاده هیات علمی دانشکده صنایع
 دانشگاه خواجه نصیر طوسی

Database

TID	Items
100	1 3 4
200	2 3 5
300	1 2 3 5
400	2 5

L_1

Itemset	Support
{1}	2
{2}	3
{3}	3
{5}	3

$C_1^{\wedge}$

TID	Set-of-itemsets
100	{{1},{3},{4}}
200	{{2},{3},{5}}
300	{{1},{2},{3},{5}}
400	{{2},{5}}

C_2

itemset
{1 2}
{1 3}
{1 5}
{2 3}
{2 5}
{3 5}

$C_2^{\wedge}$

TID	Set-of-itemsets
100	{{1 3}}
200	{{2 3},{2 5},{3 5}}
300	{{1 2},{1 3},{1 5},{2 3},{2 5},{3 5}}
400	{{2 5}}

L_2

Itemset	Support
{1 3}	2
{2 3}	2
{2 5}	3
{3 5}	2

C_3

itemset
{2 3 5}

$C_3^{\wedge}$

TID	Set-of-itemsets
200	{{2 3 5}}
300	{{2 3 5}}

L_3

Itemset	Support
{2 3 5}	2

22

سمیه علیزاده هیات علمی دانشکده صنایع



مقایسه AprioriTid و Apriori

- همانگونه که قبلاً نیز ذکر شد الگوریتم Apriori همه پایگاه داده را در هر گذر می‌پیماید تا Support ها را محاسبه کند و پیمودن همه پایگاه داده ممکن است در همه فازها مورد نیاز نباشد.
- تفاوت عمده ای که این الگوریتم با الگوریتم Apriori دارد در این است که این الگوریتم کل پایگاه داده را برای محاسبه Support بعد از مرحله اول نمی‌پیماید و از مجموعه C^k برای محاسبه Support استفاده می‌کند.



مقایسه AprioriTid و Apriori

- در الگوریتم AprioriTid مقادیر C^k بجای پایگاه داده در نظر گرفته می‌شوند. اگر C^k بتواند در حافظه Fit شود، این الگوریتم سریعتر از Apriori عمل خواهد کرد.
- زمانیکه C^k خیلی بزرگ باشد نمی‌تواند در حافظه جای بگیرد و در نتیجه زمان محاسبه بسیار بالا می‌رود، بنابراین الگوریتم Apriori سریعتر از الگوریتم AprioriTid عمل خواهد کرد.

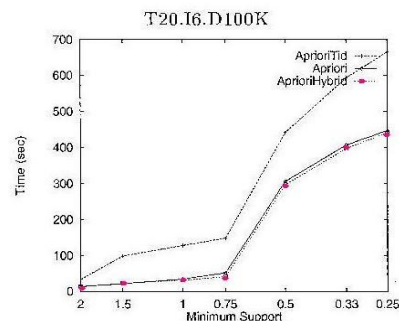
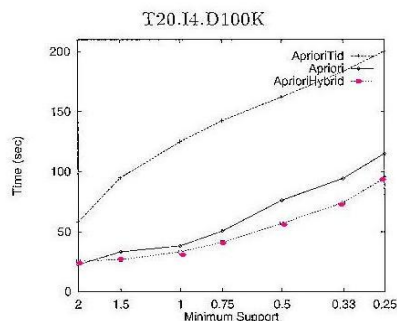
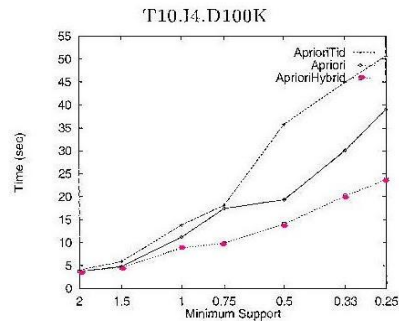
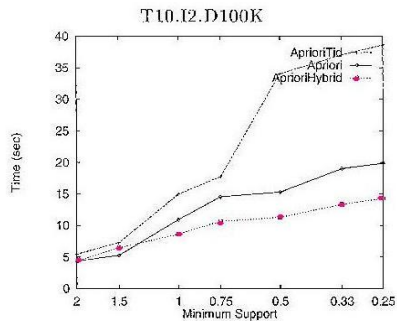
الگوریتم Apriori Hybrid



خصوصیات این الگوریتم به ترتیب زیر هستند:

- این الگوریتم در فازهای اولیه اجرا مطابق الگوریتم Apriori عمل می‌کند.
- اندازه تخمینی C^k را بصورت زیر در نظر می‌گیرد:
- وقتی که انتظار می‌رود C^k مناسب حافظه شود به الگوریتم AprioriTid سوئیچ کرده و مطابق این الگوریتم پیش می‌رود.
- اگرچه تغییر از Apriori به AprioriTid زمان بر است ، اما در بسیاری از موارد نتایج مثبتی دارد.

25



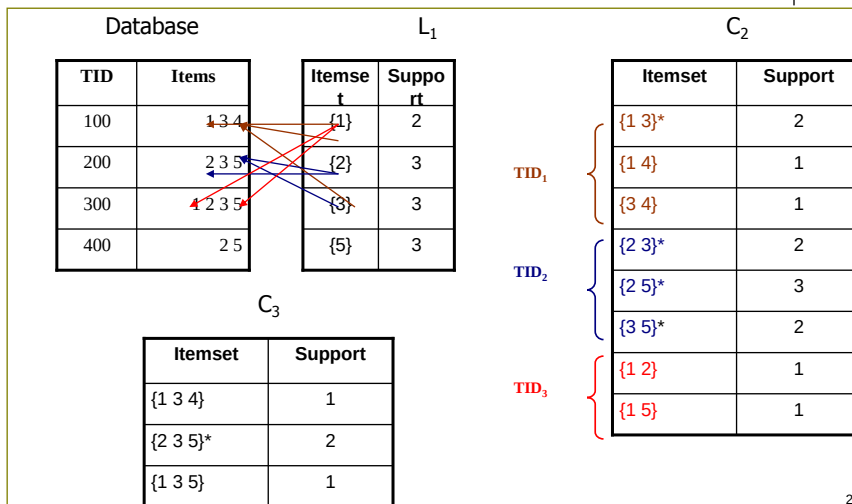


دیگر الگوریتمها

- این الگوریتم از اولین الگوریتمهایی بود که برای استخراج همه آیت‌های بزرگ از پایگاه داده در سال 1993 توسط R.Agrawal, T. Imielinski, and A. Swami ابداع گردید و درحقیقت نام این الگوریتم برگرفته از نامهای اول ابداع کنندگان آن می‌باشد.
- این الگوریتم چندین گذر را بر روی پایگاه داده انجام می‌دهد و در هر گذر همه تراکشن‌ها را می‌پیماید. گامهای این الگوریتم به صورت زیر می‌باشند:
- برای هر یک از تراکشن‌ها بزرگترین آیت‌ها انتخاب می‌شوند.
- آیت‌های کاندیدا (CK) با گسترش هر یک از این آیت‌های بزرگ به سایر آیت‌ها در هر تراکشن ساخته می‌شوند.



الگوریتم AIS





الگوریتم AIS

- از معایب این روش اینستکه در هر گذر تعدادی مجموعه آیتم انتخاب می شوند که از حداقل مقدار پشتیبان (در اینجا 2) برخوردار نیستند و بایستی کنار گذاشته شوند.

- بعنوان مثال در $C2$ مجموعه آیتمهای اضافی عبارتند از $\{1,4\}, \{3,4\}, \{1,2\}, \{1,5\}$



الگوریتم SETM

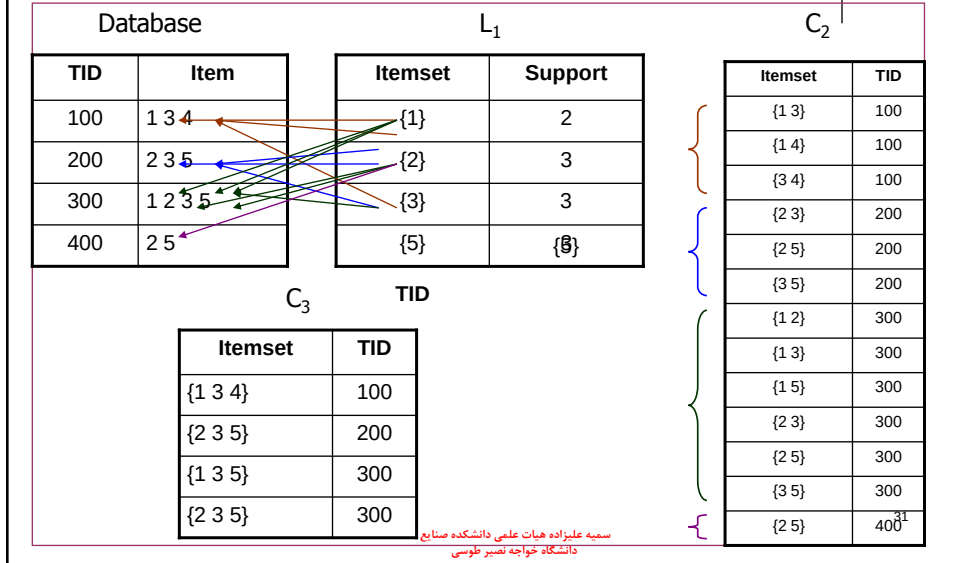
این الگوریتم توسط Houtsma در سال 1995 ابداع شد و در سال 1996 نسخه دوم آن بهمنظور محاسبه Itemset های بزرگ در SQL توسط Srikant مطرح شد

مشابه الگوریتم AIS ، این الگوریتم نیز چندین گذر بر روی پایگاه داده انجام می دهد. گامهای این الگوریتم به قرار زیر می باشند:

- پشتیبان هر یک از آیتمهای مجزا محاسبه و بزرگترین آنها انتخاب می شوند

آیتمهای کاندیدا (CK) با گسترش هر یک از این آیتمهای بزرگ به سایر آیتمها در هر تراکش ساخته می شوند. علاوه بر آن در این مرحله TID های مربوط به هر یک از CK را در یک ساختار ترتیبی نگهداری می کند و سپس support هر یک از CK ها با جمع کردن این ساختار ترتیبی محاسبه می کند.

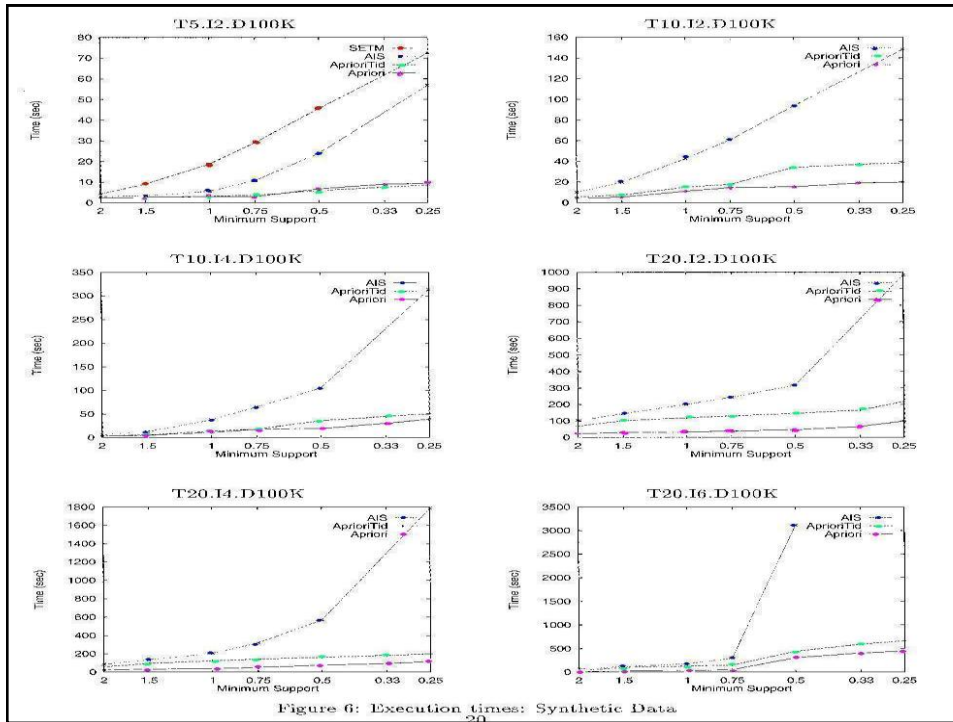
الگوریتم SETM



معایب الگوریتمهای SETM و AIS



- این الگوریتمها خیلی کند هستند
- Itemset های زیادی با Support/Confidence پائینتر از مینیمم Support/Confidence در نظر گرفته شده توسط کاربر تولید می کنند.



Statistical-based Measures

- Measures that take into account statistical dependence

$$Lift = \frac{P(Y | X)}{P(Y)}$$

$$Interest = \frac{P(X, Y)}{P(X)P(Y)}$$

$$PS = P(X, Y) - P(X)P(Y)$$

$$\phi - coefficient = \frac{P(X, Y) - P(X)P(Y)}{\sqrt{P(X)[1 - P(X)]P(Y)[1 - P(Y)]}}$$

سمیه علیزاده هیات علمی، دانشکده صنایع
دانشگاه خواجه نصیر طوسی



Example: Lift/Interest



	Coffee	<u>Coffee</u>	
Tea	15	5	20
<u>Tea</u>	75	5	80
	90	10	100

Association Rule: Tea $\rightarrow$ Coffee

Confidence= $P(\text{Coffee}|\text{Tea}) = 0.75$

but $P(\text{Coffee}) = 0.9$

$\Rightarrow \text{Lift} = 0.75/0.9 = 0.8333 (< 1, \text{therefore is negatively associated})$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

Drawback of Lift & Interest



	Y	<u>Y</u>	
X	10	0	10
<u>X</u>	0	90	90
	10	90	100

	Y	<u>Y</u>	
X	90	0	90
<u>X</u>	0	10	10
	90	10	100

$$\text{Lift} = \frac{0.1}{(0.1)(0.1)} = 10$$

$$\text{Lift} = \frac{0.9}{(0.9)(0.9)} = 1.11$$

Statistical independence:

If $P(X,Y)=P(X)P(Y) \Rightarrow \text{Lift} = 1$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

Which Measures Should Be Used?



- **lift and χ^2** are not good measures for correlations in large transactional DBs
- **all-conf or coherence** could be good measures (Omiecinski@TKDE'03)
- Both **all-conf** and **coherence** have the downward closure property
- Efficient algorithms can be derived for mining (Lee et al. @ICDM'03sub)

May 20, 2012

symbol	measure	range	formula
ϕ	ϕ -coefficient	-1 ... 1	$\frac{P(A,B) - P(A)P(B)}{\sqrt{P(A)P(B)(1-P(A))(1-P(B))}}$
Q	Yule's Q	-1 ... 1	$\frac{P(A,B)P(\bar{A}\bar{B}) - P(A,\bar{B})P(\bar{A},B)}{P(A,B)P(\bar{A}\bar{B}) + P(A,\bar{B})P(\bar{A},B)}$
Y	Yule's Y	-1 ... 1	$\frac{\sqrt{P(A,B)P(\bar{A}\bar{B})} - \sqrt{P(A,\bar{B})P(\bar{A},B)}}{\sqrt{P(A,B)P(\bar{A}\bar{B})} + \sqrt{P(A,\bar{B})P(\bar{A},B)}}$
k	Cohen's	-1 ... 1	$\frac{P(A,B) + P(\bar{A}\bar{B}) - P(A)P(B) - P(\bar{A})P(\bar{B})}{1 - P(A)P(B) - P(\bar{A})P(\bar{B})}$
PS	Piatetsky-Shapiro's	-0.25 ... 0.25	$P(A,B) - P(A)P(B)$
F	Certainty factor	-1 ... 1	$\max(\frac{P(B A) - P(B)}{1 - P(B)}, \frac{P(A B) - P(A)}{1 - P(A)})$
AV	added value	-0.5 ... 1	$\max(P(B A) - P(B), P(A B) - P(A))$
K	Klogsen's Q	-0.33 ... 0.38	$\frac{\sqrt{P(A,B)} \max(P(B A) - P(B), P(A B) - P(A))}{\sum_j \max_k P(A_j, B_k) + \sum_k \max_j P(A_j, B_k) - \max_k P(A_j) - \max_k P(B_k)}$
g	Goodman-kruskal's	0 ... 1	$\frac{2 - \max_j P(A_j) - \max_k P(B_k)}{\sum_i \sum_j P(A_i, B_j) \log \frac{P(A_i, B_j)}{P(A_i)P(B_j)}}$
M	Mutual Information	0 ... 1	$\min(-\sum_i P(A_i) \log P(A_i) \log P(A_i), -\sum_i P(B_i) \log P(B_i) \log P(B_i))$
J	J-Measure	0 ... 1	$\max(P(A,B) \log(\frac{P(A B)}{P(B)}) + P(\bar{A}\bar{B}) \log(\frac{P(\bar{A} \bar{B})}{P(\bar{A})}), P(A,B) \log(\frac{P(A B)}{P(A)}) + P(\bar{A}\bar{B}) \log(\frac{P(\bar{A} \bar{B})}{P(\bar{A})}))$
G	Gini index	0 ... 1	$\max(P(A)[P(B A)^2 + P(\bar{B} A)^2] + P(\bar{A})[P(B \bar{A})^2 + P(\bar{B} \bar{A})^2] - P(B)^2 - P(\bar{B})^2, P(B)[P(A B)^2 + P(\bar{A} B)^2] + P(\bar{B})[P(A \bar{B})^2 + P(\bar{A} \bar{B})^2] - P(A)^2 - P(\bar{A})^2)$
s	support	0 ... 1	$P(A,B)$
c	confidence	0 ... 1	$\max(P(B A), P(A B))$
L	Laplace	0 ... 1	$\max(\frac{NP(A,B)+1}{NP(A)+2}, \frac{NP(A,B)+1}{NP(B)+2})$
IS	Cosine	0 ... 1	$\frac{P(A,B)}{\sqrt{P(A)P(B)}}$
γ	coherence(Jaccard)	0 ... 1	$\frac{P(A,B)}{P(A)+P(B)-P(A,B)}$
α	allconfidence	0 ... 1	$\frac{\max(P(A), P(B))}{P(A,B)}$
o	odds ratio	0 ... ∞	$\frac{P(A,B)P(\bar{A}\bar{B})}{P(A,\bar{B})P(\bar{A},B)}$
V	Conviction	0.5 ... ∞	$\max(\frac{P(A)P(\bar{B})}{P(\bar{A}B)}, \frac{P(B)P(\bar{A})}{P(\bar{B}A)})$
λ	lift	0 ... ∞	$\frac{P(A,B)}{P(A)P(B)}$
S	Collective strength	0 ... ∞	$\frac{P(A,B) + P(\bar{A}\bar{B})}{P(A)P(B) + P(\bar{A})P(\bar{B})} \times \frac{1 - P(A)P(B) - P(\bar{A})P(\bar{B})}{1 - P(A,B) - P(\bar{A}\bar{B})}$
χ^2	χ^2	0 ... ∞	$\sum_i \frac{(P(A_i) - E_i)^2}{E_i}$

سمیه علیزاده هیات علمی دانشکده صنایع
دانشگاه خواجه نصیر طوسی

symbol	measure	range	formula
ϕ	ϕ -coefficient	-1 ... 1	$\frac{P(A,B) - P(A)P(B)}{\sqrt{P(A)P(B)(1-P(A))(1-P(B))}}$
Q	Yule's Q	-1 ... 1	$\frac{P(A,B)P(\bar{A}\bar{B}) - P(A,\bar{B})P(\bar{A},B)}{P(A,B)P(\bar{A}\bar{B}) + P(A,\bar{B})P(\bar{A},B)}$
Y	Yule's Y	-1 ... 1	$\frac{\sqrt{P(A,B)P(\bar{A}\bar{B})} - \sqrt{P(A,\bar{B})P(\bar{A},B)}}{\sqrt{P(A,B)P(\bar{A}\bar{B})} + \sqrt{P(A,\bar{B})P(\bar{A},B)}}$
k	Cohen's	-1 ... 1	$\frac{P(A,B) + P(\bar{A}\bar{B}) - P(A)P(B) - P(\bar{A})P(\bar{B})}{1 - P(A)P(B) - P(\bar{A})P(\bar{B})}$
PS	Piatetsky-Shapiro's	-0.25 ... 0.25	$P(A,B) - P(A)P(B)$
F	Certainty factor	-1 ... 1	$\max(\frac{P(B A) - P(B)}{1 - P(B)}, \frac{P(A B) - P(A)}{1 - P(A)})$
AV	added value	-0.5 ... 1	$\max(P(B A) - P(B), P(A B) - P(A))$
K	Klogsen's Q	-0.33 ... 0.38	$\frac{\sqrt{P(A,B)} \max(P(B A) - P(B), P(A B) - P(A))}{\sum_j \max_k P(A_j, B_k) + \sum_k \max_j P(A_j, B_k) - \max_k P(A_j) - \max_k P(B_k)}$
g	Goodman-kruskal's	0 ... 1	$\frac{2 - \max_j P(A_j) - \max_k P(B_k)}{\sum_i \sum_j P(A_i, B_j) \log \frac{P(A_i, B_j)}{P(A_i)P(B_j)}}$
M	Mutual Information	0 ... 1	$\min(-\sum_i P(A_i) \log P(A_i) \log P(A_i), -\sum_i P(B_i) \log P(B_i) \log P(B_i))$
J	J-Measure	0 ... 1	$\max(P(A,B) \log(\frac{P(A B)}{P(B)}) + P(\bar{A}\bar{B}) \log(\frac{P(\bar{A} \bar{B})}{P(\bar{A})}), P(A,B) \log(\frac{P(A B)}{P(A)}) + P(\bar{A}\bar{B}) \log(\frac{P(\bar{A} \bar{B})}{P(\bar{A})}))$
G	Gini index	0 ... 1	$\max(P(A)[P(B A)^2 + P(\bar{B} A)^2] + P(\bar{A})[P(B \bar{A})^2 + P(\bar{B} \bar{A})^2] - P(B)^2 - P(\bar{B})^2, P(B)[P(A B)^2 + P(\bar{A} B)^2] + P(\bar{B})[P(A \bar{B})^2 + P(\bar{A} \bar{B})^2] - P(A)^2 - P(\bar{A})^2)$
s	support	0 ... 1	$P(A,B)$
c	confidence	0 ... 1	$\max(P(B A), P(A B))$
L	Laplace	0 ... 1	$\max(\frac{NP(A,B)+1}{NP(A)+2}, \frac{NP(A,B)+1}{NP(B)+2})$
IS	Cosine	0 ... 1	$\frac{P(A,B)}{\sqrt{P(A)P(B)}}$
γ	coherence(Jaccard)	0 ... 1	$\frac{P(A,B)}{P(A)+P(B)-P(A,B)}$
α	allconfidence	0 ... 1	$\frac{\max(P(A), P(B))}{P(A,B)}$
o	odds ratio	0 ... ∞	$\frac{P(A,B)P(\bar{A}\bar{B})}{P(A,\bar{B})P(\bar{A},B)}$
V	Conviction	0.5 ... ∞	$\max(\frac{P(A)P(\bar{B})}{P(\bar{A}B)}, \frac{P(B)P(\bar{A})}{P(\bar{B}A)})$
λ	lift	0 ... ∞	$\frac{P(A,B)}{P(A)P(B)}$
S	Collective strength	0 ... ∞	$\frac{P(A,B) + P(\bar{A}\bar{B})}{P(A)P(B) + P(\bar{A})P(\bar{B})} \times \frac{1 - P(A)P(B) - P(\bar{A})P(\bar{B})}{1 - P(A,B) - P(\bar{A}\bar{B})}$
χ^2	χ^2	0 ... ∞	$\sum_i \frac{(P(A_i) - E_i)^2}{E_i}$

#	Measure	Formula
1	ϕ -coefficient	$\frac{P(A,B) - P(A)P(B)}{\sqrt{P(A)P(B)(1-P(A))(1-P(B))}}$
2	Goodman-Kruskal's (λ)	$\frac{\sum_j \max_k P(A_j, B_k) + \sum_k \max_j P(A_j, B_k) - \max_j P(A_j) - \max_k P(B_k)}{2 - \max_j P(A_j) - \max_k P(B_k)}$
3	Odds ratio (α)	$\frac{P(A,B)P(\bar{A},\bar{B})}{P(A,\bar{B})P(\bar{A},B)}$
4	Yule's Q	$\frac{P(A,B)P(\bar{A},\bar{B}) - P(A,\bar{B})P(\bar{A},B)}{P(A,B)P(\bar{A},\bar{B}) + P(A,\bar{B})P(\bar{A},B)} = \frac{\alpha-1}{\alpha+1}$
5	Yule's Y	$\frac{\sqrt{P(A,B)P(\bar{A},\bar{B})} - \sqrt{P(A,\bar{B})P(\bar{A},B)}}{\sqrt{P(A,B)P(\bar{A},\bar{B})} + \sqrt{P(A,\bar{B})P(\bar{A},B)}} = \frac{\sqrt{\alpha}-1}{\sqrt{\alpha}+1}$
6	Kappa (κ)	$\frac{P(A,B) + P(\bar{A},\bar{B}) - P(A)P(B) - P(\bar{A})P(\bar{B})}{1 - P(A)P(B) - P(\bar{A})P(\bar{B})}$
7	Mutual Information (M)	$\frac{\sum_i \sum_j P(A_i, B_j) \log \frac{P(A_i, B_j)}{P(A_i)P(B_j)}}{\min(-\sum_i P(A_i) \log P(A_i), -\sum_j P(B_j) \log P(B_j))}$
8	J-Measure (J)	$\max \left(P(A, B) \log \left(\frac{P(B A)}{P(B)} \right) + P(\bar{A}\bar{B}) \log \left(\frac{P(\bar{B} \bar{A})}{P(\bar{B})} \right), \right.$ $\left. P(A, \bar{B}) \log \left(\frac{P(\bar{B} A)}{P(\bar{B})} \right) + P(\bar{A}B) \log \left(\frac{P(B \bar{A})}{P(B)} \right) \right)$
9	Gini index (G)	$\max \left(P(A)[P(B A)^2 + P(\bar{B} A)^2] + P(\bar{A})[P(B \bar{A})^2 + P(\bar{B} \bar{A})^2] \right.$ $\left. - P(B)^2 - P(\bar{B})^2, \right.$ $\left. P(B)[P(A B)^2 + P(\bar{A} B)^2] + P(\bar{B})[P(A \bar{B})^2 + P(\bar{A} \bar{B})^2] \right.$ $\left. - P(A)^2 - P(\bar{A})^2 \right)$
10	Support (s)	$P(A, B)$
11	Confidence (c)	$\max(P(B A), P(A B))$
12	Laplace (L)	$\max \left(\frac{NP(A,B)+1}{NP(A)+2}, \frac{NP(A,B)+1}{NP(B)+2} \right)$
13	Conviction (V)	$\max \left(\frac{P(A)P(\bar{B})}{P(A\bar{B})}, \frac{P(B)P(\bar{A})}{P(\bar{A}B)} \right)$
14	Interest (I)	$\frac{P(A,B)}{P(\bar{A})P(\bar{B})}$
15	cosine (IS)	$\frac{P(A,B)}{\sqrt{P(A)P(B)}}$
16	Piatetsky-Shapiro's (PS)	$P(A, B) - P(A)P(B)$
17	Certainty factor (F)	$\max \left(\frac{P(B A) - P(B)}{1 - P(B)}, \frac{P(A B) - P(A)}{1 - P(A)} \right)$
18	Added Value (AV)	$\max(P(B A) - P(B), P(A B) - P(A))$
19	Collective strength (S)	$\frac{P(A,B) + P(\bar{A}\bar{B})}{P(A)P(B) + P(\bar{A})P(\bar{B})} \times \frac{1 - P(A)P(B) - P(\bar{A})P(\bar{B})}{1 - P(A,B) - P(\bar{A}\bar{B})}$
20	Jaccard (ζ)	$\frac{P(A,B)}{P(A) + P(B) - P(A,B)}$
21	Kloggen (K)	$\sqrt{P(\bar{A}, \bar{B})} \max(P(B A) - P(B), P(A B) - P(A))$

