

一种新的基于投影的频繁模式树构造算法^{*}

李陶深¹ 李新仕^{1,2}

(广西大学计算机与电子信息学院 南宁 530004)¹ (广西财经学院计算机与信息管理系 南宁 530003)²

摘要 本文分析 FP-growth 算法存在的主要问题,提出了一种新的基于投影的频繁模式树构造算法。该算法充分利用大型数据库的投影运算能力,按层来构造频繁模式树(FP-tree),有效地解决了传统的 FP-tree 构造中存在的问题。实验结果表明,本文的算法与传统的频繁模式树的构造算法相比,具有比较好的时间和空间的伸缩性。

关键词 数据挖掘,关联规则,频繁模式树,投影后插式频繁模式树

1 引言

数据挖掘指的是从大量的、不完全的、有噪声的、模糊的、随机的数据中提取默认在其中的、事先不知道的、潜在具有价值的有用信息和知识的过程^[1]。关联规则的挖掘问题于 1993 年由 Rabesh Agrawal 等人首先提出^[2],它作为数据挖掘中的一个重要领域,用于从大量数据中发现项集之间的有趣关联或相关联系^[1]。目前,关联规则挖掘方法的研究大多集中在基于频繁模式树的关联规则挖掘算法的研究中,主要的算法是 Apriori 算法^[3]和 FP-growth 算法^[4],以及它们的改进算法。

FP-growth 算法是一个本质上不同于 Apriori 算法的挖掘频繁模式的有效方法,它不生成候选频繁项集且只需要两次扫描数据库,有效克服了 Apriori 算法的缺点。在 FP-growth 算法研究领域,基于 FP-tree 的频繁闭包项目集挖掘算法主要是 Jian Pei 等提出的 Closet 算法^[5]和 Closet 算法^[6],之后的闭包项目集挖掘算法主要都是在此基础上进行改进的,例如文[7,8]中提出的算法。除了挖掘最大频繁项目集和挖掘频繁闭包项目集之外,许多研究人员对频繁模式树的构造和挖掘过程也进行了大量的研究。文[9]基于共同事务频繁项目树提出的 OFI-tree Mining 算法,通过有效的剪枝节约大量的内存空间;文[10]提出了一种基于被约束子树挖掘频繁项集的有效算法,通过引入了被约束子树,在挖掘频繁模式时不生成条件 FP 树,从而提高频繁模式挖掘的时空效率;文[11]中的 CFP-tree 是对 FP-tree 的进一步压缩存储,并对交易数据库进行投影分割,然后分别进行关联规则的挖掘。

虽然 FP-growth 算法克服了 Apriori 系列算法的缺陷,取得了很好的效果,但它本身仍然存在着不足。该算法是通过逐步生成条件模式基和条件频繁

模式树来挖掘频繁项目集的,在频繁项目集和事务量比较大的情况下,需要动态地产生和释放成千上万的条件频繁模式树,因此占用大量的时间和空间。另外,条件模式基的生成需要不断地搜索 FP-tree,原因在于:FP-tree 是自顶而下构造的,而条件模式基的生成是自底而上的顺序来产生的。如果 FP-tree 采用双向指针结构,则需要占用了更多的存储空间。本文提出了一种新的基于投影的频繁模式树后插式构造方法,支持自底而上的指针链,以满足产生条件模式基的需要,可有效地解决了上述问题。

2 基于投影的频繁模式树后插式构造方法

本文对频繁模式树的构造方法进行改进,以减少搜索父结点的时间,同时又能克服上述问题。我们把改进后的构造方法生成的频繁模式树称为投影后插式频繁模式树(Projection Rear-inserting Frequent Pattern Tree, PRIFP-tree)。

2.1 构造 PRIFP-tree 的基本思想

由构造 FP-tree 的过程中可知,传统的 FP-tree 的构造是一个严格的串行计算的过程,当事务数据库中的数据量很大时,会造成性能的瓶颈。本文提出的基于投影的频繁模式树后插式构造方法的基本思想主要有两个方面:

(1)充分利用了大型数据库的投影运算能力,解决传统 FP-tree 的构造算法中随着数据库中记录的增加,算法性能急剧下降的瓶颈问题。利用数据库的投影运算,按层次来构造频繁模式树,不会造成随着数据库记录增加算法性能急剧下降的缺点,算法的时间复杂度主要与投影的次数有关系。首先创建一个临时数据库用于存放所有排序后的频繁项,这个临时数据库我们称之为投影数据库(PDB),可以被用来运用投影技术来构造树的结点,而不像传统的 FP-tree 那样逐个地计算每个频繁项。然后对

^{*} 基金项目:广西“新世纪十百千人才工程”专项基金项目(桂人字 2001213 号)和广西自然科学基金项目(桂科自 0229008)联合资助。

PDB 进行投影计算,一次投影两列, $j-1$ 列和 j 列,第 j 列被称为当前结点列,用来统计所有不同频繁项的出现次数,第 $j-1$ 列被称之为父结点列,用以判断当前结点的父结点。这样,一次就可以计算出树的一层结点并把它们链接到对应的父结点上,不必要逐个地计算每个频繁项。

(2)在构造频繁模式树时,采用后插式的方法,以便解决随着频繁模式树深度增加、搜索插入点(父结点)的时间不断增加的问题。在重构 FP-tree 时,将 FP-tree 形成一个链表,在链表中的结点按支持度的降序排列。数据库中的频繁事务项是严格按照支持度从高到低排列的,新一次投影得到的分枝结点在一般情况下会比前一次投影得到的分枝结点的支持度低,所以当按层构造 FP-tree 时,新一层结点一般都插入到链表的尾部,而其父结点为上一层,搜索其父结点也不会花费很大的代价。

2.2 PRIFP-tree 的构造算法

假设算法使用的树结构中的每个结点有四个域,分别为: item_name(存放项目值)、count(支持度计数)、parent(指向父结点)、node_link(指向下一个结点)。下面给出 PRIFP-tree 的构造算法。

输入:事务数据库 DB,最小支持度阈值 ξ

输出:PRIFP-tree

Step1:扫描事务数据库 DB 一次,收集频繁项的集合 F 和它们的支持度。对 F 按支持度降序排序,结果为频繁项表 L 。

Step2:(2)根据 L ,对 DB 的每个事务中的频繁项按 L 中的顺序进行排序,排序后的结果存放在临时投影数据库 PDB 中。

Step3:创建根结点。假设 j 为列号,对 PDB 中的每一列顺序执行以下操作,直到投影完 PDB 中所有的列:

if $j=1$ then {

① 对第 1 列作投影,并计算该列中每个频繁项的出现次数,结果形为 $[root, q:n]$ (其中 q 为频繁项名, n 为累计出现次数, $root$ 代表该结点的父结点是根结点);

② 将这些频繁项 $[q:n]$ 作为 $root$ 的子结点加入到 PRIFP-tree 中;

}

else {

① 对第 $j-1$ 列和第 j 列两列进行投影,统计所有父结点和当前列结点相同的二元频繁项组的出现次数,结果形为二元频繁项组 $[p, q:n]$ (n 为累计出现的次数),并对二元频繁项组按 q 的支持度升序排序;

② 对上面的二元频繁项组进行比较,若存在二元频繁项组中的 q 相同,则将它们进行区别,即在 q

的后面加上后缀(顺序号,从 1 开始编号),并修改 PDB 相关的项目;

③ 对二元频繁项目组 $[p, q:n]$ 的每一个二元组进行以下操作:从表尾(rear)开始向前搜索直到某个结点的值与 q 相等(或者某一个结点的支持度大于 q)为止,产生新结点,将新结点插入到该结点之后,将 item_name 域设置为 q ,将 count 域设置为 n ,然后继续往前搜索直到某一个结点值为 p 为止,该结点即为新结点的父结点,将新结点的 parent 域指向父结点。

}

$j=j+1$;

2.3 PRIFP-tree 的构造过程

下面通过一个实例,说明 PRIFP-tree 的构造过程。所用的事务交易数据库如表 1 所示。

表 1 事务交易数据库

TID	事务项	排序后的频繁事务项(PDB)	投影修改后的 PDB
100	f,a,c,d,g,l,m,p	f,c,a,m,p	f,c,a,m,p
200	a,b,c,f,l,o	f,c,a,b,o	f,c,a,b,o
300	b,f,h,j,m,p	f,b,m,p	f,b1,m1,p
400	b,c,k,m,o,s	c,b,m,o	c,b2,m2,o1
500	a,f,c,e,l,n,o,p	f,c,a,o,p	f,c,a,o2,p

由上述算法,可以看出 PRIFP-tree 的构造思路是:

(1) 第一遍扫描数据库,得到各项目的出现频率(支持度)如下: $a:3, b:3, c:4, d:1, e:1, f:4, g:1, h:1, i:1, j:1, k:1, l:2, m:3, n:1, o:3, p:3, s:1$ (“:”后的数字表示支持度)

(2) 取最小支持度阈值 $\xi=3$,得到频繁项目集 L 并按支持度由高到低排列如下(支持度相同按先后顺序): $L=[f:4, c:4, a:3, b:3, m:3, o:3, p:3]$

(3) 根据 L ,对 DB 的每一个事务中的频繁项按 L 中的顺序进行排序,排序后的结果存放在临时的投影数据库 PDB 中。

(4) 最后利用 PDB 进行投影,最后得到如图 1 所示的 PRIFP-tree。

3 性能分析

3.1 PRIFP-tree 构造算法的性质

从以上 PRIFP-tree 的构造过程,可以得出 PRIFP-tree 的构造算法具有以下性质:

引理 1 给定一个事务数据库 DB 和它的最小支持度阈值 ξ ,它对应的 PRIFP-tree 包含了数据库中用于挖掘所有的频繁项的信息。

证明:在 FP-tree 的构造过程中,事务数据库中的每一个事务都对应着 FP-tree 的一条路径,而每一个事务中频繁模式项的信息完全保存在 FP-tree 中。另外,在 FP-tree 中的每个路径中频繁项集有可能被多个事务所共享,因为所有的事务路径都必

须从 root 结点开始,它们有可能共享前缀路径。因此 FP-tree 包含了数据库中用于挖掘所有频繁项的信息。而 PRIFP-tree 仅仅对 FP-tree 的指针方向进

行改变,因此与 FP-tree 一样,数据库中的所有频繁项模式信息被完整地保存到了对应的 PRIFP-tree 中。证毕。

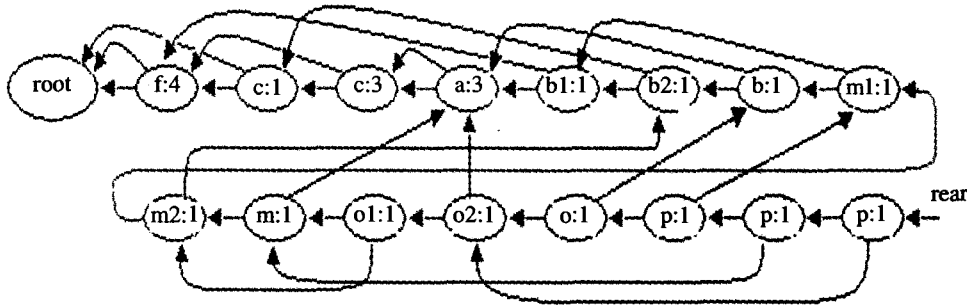


图3 PRIFP-tree

引理 2 对第 $j-1$ 和第 j 列进行投影得到的序列如果按照第 j 列的支持度由小到大进行排列,然后按顺序插入到 FP-tree 中,不会产生异常。

证明:由于频繁项目集中的每个事务中的项目是按支持度从大到小排列的,因此对第 $j-1$ 列和第 j 列进行投影时得到的分枝集合中的任一个分枝(设为 $[x, y]$), x 的支持度比 y 的支持度大。若分枝集合按第 j 列的支持度从小到大排列,则对于任一个分枝 $[x, y]$, x 是 y 的父结点,而 x 作为孩子的分枝(设为 $[z, x]$),其必排列在分枝 $[x, y]$ 之后,因此分枝 $[x, y]$ 优先插入,即插入时往前搜索到的第一个 x 结点不可能是同层结点而是上一层结点,因此不会产生异常。证毕。

3.2 实验分析

本文的测试主要针对频繁模式树的构造及频繁模式基的产生这个阶段。实验环境为:P4 1.4GHZ CPU, 256MB 内存, Windows XP 操作系统, 采用 Delphi 7.0 进行编程。为了分析不同数据库系统对 PRIFP-tree 构造性能的影响, 数据库系统分别采用 SQL server 2000 和 Visual Foxpro 6.0 进行测试。为了测试算法的运行时间和占用存储空间随着支持度减小而变化的情况, 采用的模拟数据库共有 3000 个事务, 200 个项目, 事务平均长度为 10, 运行时间随支持度变化情况如图 2 示。

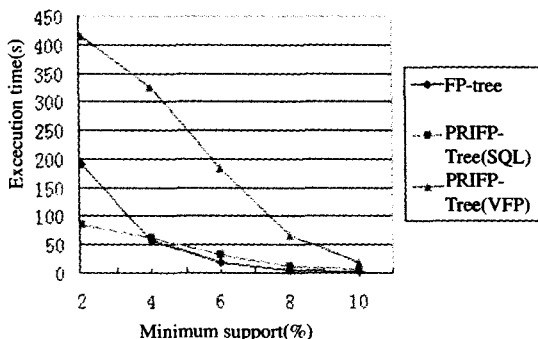
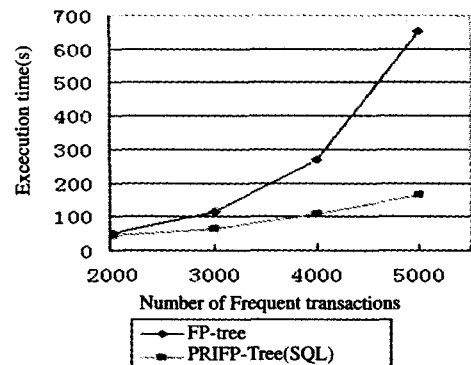


图2 运行时间与支持度的关系

为了测试算法随着事务数增加而变化的情况,

采用的模拟数据库共 200 个频繁项目, 频繁事务的平均长度为 6, 实验结果如图 3 所示。



分析以上实验结果可得出以下结论:随着支持度的减少,无论是频繁项目集和频繁事务长度的增加,还是随着事务数的增加,FP-tree 的结点数都增长比较快,因此搜索 FP-tree 产生条件模式基的时间也相应地增长比较快。而基于 PRIFP-tree 的算法在产生条件模式基时不需要搜索频繁模式树,因此花费的时间主要集中在构造频繁模式树时需要不断地对数据库进行更新,但所花费的时间的增长速度却比较平缓。

结束语 本文提出一种基于投影的频繁模式树后插式构造方法(PRIFP-tree 构造方法),通过对传统的频繁模式树进行重构,解决了频繁模式树的自顶而下构造与频繁模式自底而上挖掘的矛盾。新构造方法主要运用大型数据库的投影技术,充分利用计算机的内存和计算能力,按层来生成频繁模式树,克服了传统频繁模式树构造中的缺点,即传统的频繁模式树的构造是一个严格的串行计算过程。

实验结果表明:基于投影的频繁模式树后插式构造方法与传统的频繁模式树的构造方法相比较,具有比较好的时间和空间的伸缩性。

参考文献

- 1 Han J, Kamber M. Data Mining: Concepts and Techniques. Morgan Kaufman, San Francisco, CA, 2001

(下转第 177 页)

