

Premises Reduction of Rule-Based Expert Systems Using Association Rules Technique

Dr. Ahmed Tariq Sadiq
Computer Sciences Department
University of Technology
Baghdad, Iraq

Abstract

The heart of expert system is the knowledge base that determines the power of expert system and search engine space. The one important form of this knowledge is the production rule. The premises of these rules are the heart of production rules, therefore, the reduction of these premises is very useful to reduce the time and space in search engine. In this paper an approach will be presented to reduce the premises of production rules using association rules algorithm especially *Apriori* algorithm. Applying this approach gives a logical premise reduction which plays a good role in the search engine.

Keywords: *Expert Systems, Premises Reduction, Rule-Based Expert Systems, Production Rules, Association Rules, Apriori Algorithm Applications.*

1- Introduction

An expert system is a computer program that represents and reasons with knowledge of some specialist subject with a view to solving problems or giving advice [5]. An expert system may completely fulfill a function that normally requires human expertise, or it may play the role of an assistant to a human-maker. The symbolic reasoning of an expert system enables it not only to draw conclusions, through a process similar to the one used by human experts, but also to provide explanations concerning their estimations. Expert system technology is based on the domain knowledge of the problem being addressed. A problem domain defines the objects, properties, tasks and events within which a human expert works and also the heuristics that trained professionals have learned to use in order to perform better [7], user interface provides the mechanism [1].

Figure (1) shows the most important modules that make up a rule-based expert system. The *user interacts* with the expert system through a user interface that simplifies communication and hides much of the system complexity (e.g. the internal structure of the rule base). An expert system

employs a variety of interface styles, including question and answer, menu driven, natural language, or graphics interfaces, where the final decision on interface type is a compromise between user needs and the requirements of the knowledge base and inferencing system [8].

The heart of the expert system is the *general knowledge base*, which contains the problem solving knowledge of the particular application. The *inference engine* applies the knowledge to the solution of actual problems. It is essentially an interpreter for the knowledge base. The program must keep track of *case-specific data*: the facts, conclusions, and other information items relevant to the case under consideration. This includes the data given in a problem instance, partial conclusions, confidence measures of conclusions, and dead ends in the search process. This information is separate from the general knowledge base. The *explanation subsystem* allows the program to explain its reasoning to the user. These explanations include justifications for the system's conclusions (in response to how queries), explanations of why the system needs a particular piece of data (why queries), and, where useful, tutorial explanations or deeper theoretical justifications of the program's actions [8].

Many systems also include a *knowledge-base editor*. Knowledge-base editors help the programmer locate and correct bugs in the program's performance, often accessing the information provided by the explanation subsystem. They also may assist in the addition of new knowledge, help maintain correct rule syntax, and perform consistency checks on the updated knowledge base [8].

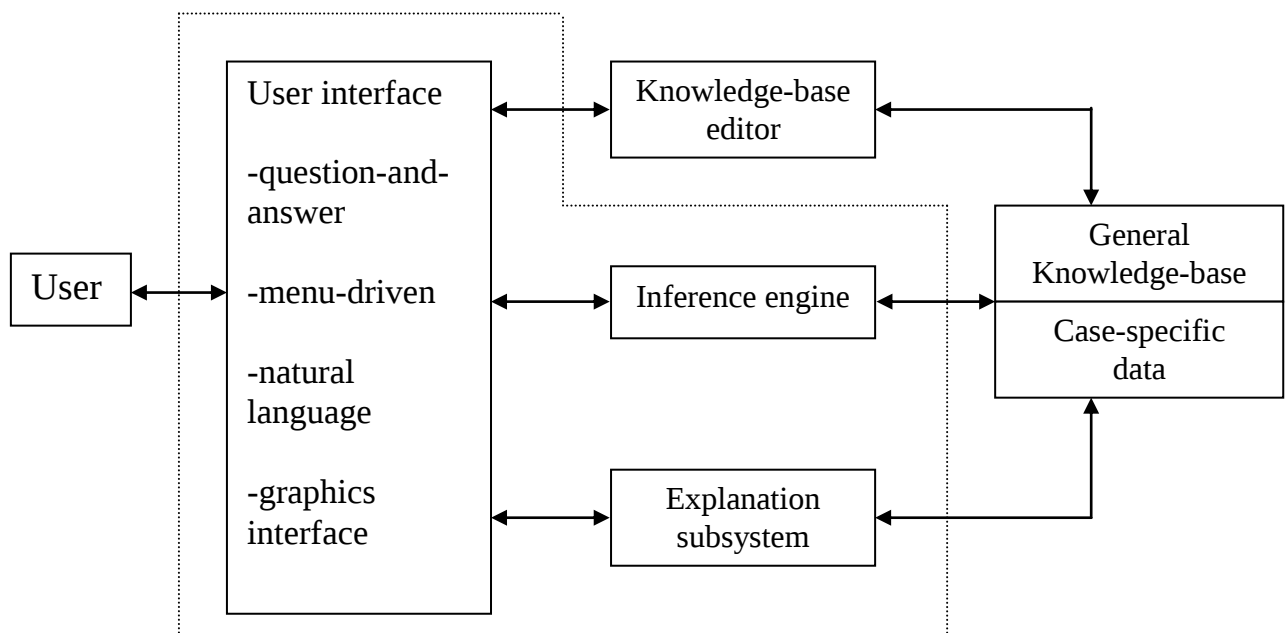


Figure (1) : Architecture of a Typical Expert System

2- Production Rules

Knowledge in an expert system can be represented in various ways. Some of the well-known techniques for representation of knowledge include Production System, Logical Calculus and Structured Models [2].

The production system is the simplest and one of the oldest techniques for knowledge representation. A production system consists of three items: 1) a set of production rules (PR), which together forms the knowledge base, 2) one (or more) dynamic database(s), called the working memory, and 3) a control structure/interpreter, which interpret the database using the set of production rules [2, 5, 9]. The production system, which has wide applications in automata theory, formal grammars and the design of programming languages, however, entered into knowledge engineering by Baughman and Feigenbaum [3] only a few years ago.

The structure of production rules can be formally stated as follows:

$$\text{PR1: } P_1(X) \wedge P_2(Y) \wedge \dots \wedge P_n(X,Y) \longrightarrow Q_1(Y) \vee Q_2(Z) \vee \dots \vee Q_m(Z,X)$$

where P and Q are predicates; X,Y and Z are variables; “ $\wedge$ ”, “ $\vee$ ” and “ $\longrightarrow$ ” denote the logical AND, OR and IF-THEN operators respectively. The left-hand side of a production rules is called the *premises/antecedent/conditional* part and the right-hand side is called the *consequent/conclusion* part [2, 8].

Many of the expert systems were production systems. This probably came about because production systems were invented for cognitive modeling, and they seem to give quite an impressive account of some mental processes. The creators of expert systems believed that they have to model the mental apparatus of their experts, so they naturally turned to production systems [4].

3- Knowledge Base Optimization in a Production System

The performance of production system depends largely on the organization of its knowledge base. The inference derived by a production system per unit time, also called time efficiency, can be improved by reducing the matching time of the premises of production rules with the data in the working memory. Further, if the rules are constructed in a manner so that there is no conflict resolution to can be avoided. Another important issue of rule-base design is to select the rules so that the resulting state-space for rule firing does not contain any cycles. The last issue is to identify the concurrently friable rules that do not have conflict in their action parts [2]. This, if realized for a rule-base system, will improve the

performance to high extent. For optimization of rules in a rule-based system, Zupan [14] suggested the following points:

- 1- Construct by backward reasoning a state-space graph from the desired goal nodes (states) up to the nodes, which cannot be expanded further in a backward manner. Each goal node, also called fixed point, is thus reachable (has connectivity) from all possible starting states. It may be noted that some of the connectivity from the starting nodes to the goal nodes may pass through the cycles. It should also be noted that the resulting state-space will not miss the shortest paths from the goal to any other state, as predecessor states of each state are found by an exhaustive breadth first search.
- 2- The common states in the graph are replaced by a single vertex and the parallel paths are identified.
- 3- Do not generate an existing state.

This paper suggests an approach to reduce the premises of production rules as an optimization approach for production rules.

4- Rule-Based Expert Systems

Rule-based expert systems represent problem-solving knowledge as *If...Then...* rules. This approach lends itself to the architecture of Figure (1), and is one of the techniques for representing domain knowledge in an expert system. It is also one of the most natural and remains widely used in practical and experimental expert systems [8].

If we regard the expert system architecture in Figure (1) as a production system, the knowledge base is the set of production rules. In a rule-based system, these condition-action pairs are represented as *If...Then...* rules, with the premises of the rules, the *if* portion, corresponding to the condition, and the conclusion, the *then* portion, corresponding to the action: when the condition is satisfied, the expert system takes the action of asserting the conclusion as true. Case-specific data are kept in the working memory. The inference engine implements the recognize-act cycle of the production system; this control may be either data-driven or goal driven [8].

In a goal-driven expert system, the goal expression is initially placed in the working memory. The system matches rule *conclusions* with the goal, selecting one rule and placing its *premises* in the working memory. This corresponds to a decomposition of the problem's goal into simpler sub-goals. The process continues in the next iteration of the production system, with these premises becoming the new goals to match against rule conclusions. The system thus works back from the original goal until all the sub-goals in working memory are known to be true, indicating that the hypothesis has been verified. Thus, backward search in an expert system

corresponds roughly to the process hypothesis testing in human problem solving [8].

In the data-driven reasoning, the algorithm for this is simple: compare the contents of working memory with the conditions of each rule in the rule base using the ordering of the rule base. If the data in working memory supports a rule's firing the result is placed in working memory and control moves on to the next rule. Once all rules have been considered search starts again at the beginning of the rule set [8].

However, the advantages of the rule-based approach include [4, 8]:

- 1- the ability to use, in a very direct fashion, experimental knowledge acquired from human expert.
- 2- The modularity of rules ease construction and maintenance.
- 3- Good performance is possible in limited domain.
- 4- Good explanation facilities.
- 5- Rules map naturally into state space search.
- 6- Rule chaining is fairly easy to trace and debug.
- 7- Steps within the problem solution process are open to inspection.
- 8- The separation of knowledge from control further simplifies development of expert systems by enabling an iterative development process where the engineer acquires, implements and tests individual rules.

On the other hand, the rule-based expert systems have disadvantages that are [8]:

- 1- Often the rules obtained from human experts are highly heuristic in nature, and lack a deeper, functional knowledge of the domain.
- 2- Missing information or unexpected data values can not be handled.
- 3- Tendency to degrade rapidly near "edges" of the domain knowledge.
- 4- Explanations function at the descriptive level only, omitting deeper, and theoretical explanations.
- 5- The knowledge tends to be very dependent.

5- Association Rules Concept

5-1 Association Rule Algorithms Definition

An association rule is a rule, which implies certain association relationships among a set of objects (such as those which occur together or one which implies the other), in a database. Given a set of transaction, where each transaction is a set of literal (called items), an association rule is an expression of the form XY , where X and Y are sets of items. The intuitive meaning of such a rule is that transactions of the database, which contain X , tend to contain Y [11].

Association rules identify relationships between attributes and items in database such as the presence or absence of one pattern implies the presence or absence of another pattern. Mining of such rules is one of the most popular pattern discovery methods in KDD, an association rule is an expression $X \Rightarrow Y$ where $X=\{x_1, x_2 \dots x_n\}$ and $Y=\{y_1, y_2 \dots y_n\}$ are set of items with left hand side (LHS) and right hand side (RHS). The meaning of such rules is quite intuitive: given database D of transaction T where each transaction $T \in D$ is a set of items $X \Rightarrow Y$ which expresses that whenever a transaction T contains X, the T probably contains Y. Also the probability of rule strength is defined as the percentage of transactions containing Y in addition to X. The prevalence for rule is the percentage of transactions that hold all the items in the union. If prevalence is low, it implies that there is no overwhelming evidence that items in $X \cup Y$ occur together.

The rule $X \Rightarrow Y$ has support S in D if the fractions of the transactions in D contain $(X \cup Y)$ [13].

The problem of mining association rules is to generate all association rules that have certain user- specified minimum support (called min-sup) and confidence (called min-conf) [12]. The important measures for association rules, support (S) and confidence (C) can be defined as: The support (S) of an association rule is the ratio (in percent) of the records that contain $(X \cup Y)$ to the total number of records in database. Therefore, if we say that the support of a rule is 5% then it means that 5% of the total records contains $(X \cup Y)$. Support is the statistical significance of a rule [12].

$$\text{Support}(X \Rightarrow Y) = P(X \cup Y)$$

$$\text{Support}(X \Rightarrow Y) = \text{frequent}(X \cup Y) / \text{total number of records in database.}$$

For a given number of records, confidence (C) is the ratio (in percent) of the numbers of records that contain $(X \cup Y)$, to the number of records that contain X. thus, if we say that a rule has a confidence of 5% it means that 85% of the records containing X also contain Y. The confidence of a rule indicates the degree of correlation in the database between X and Y.

$$\text{Confidence}(X \Rightarrow Y) = P(Y \cup X)$$

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

Confidence is also a measure of rules strength. Mining of database consists of finding all rules that meet the user-specified threshold support and confidence.

5-2-2 Apriori Algorithm

Apriori is an influential algorithm for frequent itemsets for association rules was introduced by Agrawal [10]. The name of the algorithm is based on the fact that the algorithm uses prior knowledge of frequent itemset properties. Apriori employs an iterative approach known as a level-wise search, where k -itemsets are used to explore $(k+1)$ -itemsets. First, the set of frequent 1-itemsets is found. This set is denoted L_1 . L_1 is used to find L_2 , the set of frequent 2-itemsets, which is used to find L_3 , and so on, until no more frequent k -itemsets can be found. The finding of each L_k requires one full scan of the database [6].

To improve the efficiency of the level-wise generation of frequent itemsets, an important property called Apriori property, is used to reduce the search space. **Apriori property** is based on this concept : If an itemset I does not satisfy the minimum support threshold (minsup), then I is not frequent, that is, $P(I) < \text{min-sup}$. If an item A is added to the itemset I then the resulting itemset (i.e. $I \cup A$) cannot occur more frequently than I . Therefore, $I \cup A$ is not frequent either, that is, $P(I \cup A) < \text{min-sup}$ this property belongs to a special category of properties called anti-monotone in the sense that if a set cannot pass a test, all of its supersets will fail the same test as well. It is called anti-monotone because the property is monotonic in the context of failing a test. In general to find L_k two step process consisting of **join** and **prune** actions :

1- The **Join Step**: to find L_k is generated by joining L_{k-1} with itself. This set of candidates is denoted C_k , let $l1$ and $l2$ be itemsets in L_{k-1} . The notation $li[j]$ refers to j th item in li ($l1[k-2]$ refers to the second to the last item in $l1$). Apriori assumes that items with a transaction or itemset are sorted in lexicographic order. The join, $L_{k-1} \text{ join } L_{k-1}$, is performed, where members of L_{k-1} are joinable if their first $(k-2)$ -items are in common, that is members $l1$ and $l2$ of L_{k-1} are joined if $(l1[1] = l2[1]) \wedge (l1[2] = l2[2]) \wedge \dots \wedge (l1[k-2] = l2[k-2]) \wedge (l1[k-1] = l2[k-1])$. The condition $(l1[k-1] = l2[k-1])$ simply ensure that no duplicates are generated. The result itemset formed by joining $l1$ and $l2$ is $l1[1] l1[2] \dots l1[k-1] = l2[k-1]$.

2- The **Prune Step** : C_k is a superset of L_k , that is, its members may or may not be frequent, but all of the frequent k -itemsets are included in C_k . A scan of the database to determine the count of each candidate in C_k would result in the determination of L_k (i.e. all candidate having a count no less than the minimum support count are frequent by definition, and therefore belong to L_k). C_k , however, can be huge, and so this could involve heavy computation. To reduce the size of C_k , the Apriori property is used as follows. Any $(K-1)$ -itemset that is not frequent cannot be a subset of a

frequent k -itemset. Hence, if any $(k-1)$ _subset of a candidate k _itemsets is not in L_{k-1} , then the candidate cannot be frequent either and so can be removed from C_k . This subset testing can be done quickly by maintaining a hash tree of all frequent itemsets. The following steps illustrate the Apriori algorithm [11].

Algorithm: Apriori finds frequent itemset using an iterative level-wise approach based on candidate generation.

Input: Database (D) of transaction;
Minimum support (min-sup)
Threshold

Output: Association Rules.

Begin

$L_1 = \text{find_frequent_1-itemsets}(D);$

For ($k = 2; L_{k-1} \neq \Phi; k++$)

Begin

$C_k = \text{apriori_gen}(L_{k-1}, \text{min_sup});$

For each transaction $t \in D$

Begin

$C_t = \text{subset}(C_k, t);$

For each candidate $c \in C_t$

$C.\text{count}++;$

End;

$L_k = \{c \in C_k \mid c.\text{count} \geq \text{min_sup}\}$

End;

Return $L = \bigcup_k L_k$;

End.

Procedure $\text{Apriori_gen}(L_{k-1}:\text{frequent}(k-1)\text{-itemsets};$
 $\text{min_sup: minimum support threshold})$

Begin

For each itemset $l_1 \in L_{k-1}$

For each itemset $l_2 \in L_{k-1}$

If $(l_1[1] = l_2[1]) \wedge (l_1[2] = l_2[2]) \wedge \dots \wedge (l_1[k-2] = l_2[k-2]) \wedge (l_1[k-1] < l_2[k-1])$ **Then**

Begin

$c = l_1 \text{ join } l_2;$

If $\text{has_infrequent_subset}(c, L_{k-1})$ **Then**

Delete c

Else

add c to C_k ;

```

        End
    Return Ck;
End.

Procedure has_infrequent_subset (c:candidate k_itemsets);
Begin
    For each (k-1)_subset s of c
        If s  $\notin$  Lk-1 then
            Return TRUE;
Return FALSE;
End.

```

6- The Suggested Approach for Premises Reduction Using Association Rules

The power of expert systems is concentrated in knowledge base, because the knowledge base represents the human experience, the time and space search in control of expert systems, therefore, the reduction of this knowledge will play a big role in the performance of expert systems. This paper focuses on reduction of “*premises*” of production rules using association rule technique especially applying the *Apriori* algorithm. The standard *Apriori* algorithm is not reasonable because it checks k-item set with only (k-1)-item set. We need to check k-item set with all previous item sets. The algorithm of suggested approach is as follows:

Input: Production rules of an expert system.

Output: Reduction production rules of an expert system.

Begin

Convert the premises of production rules into transaction database;

Apply an improved *Apriori* algorithm to find the association rules of resulting database;

Sort the association rules descending (depending on transaction length of right-hand side then left-hand side);

Repeat

Replace the symbols of first association rules with one symbol in transactions database;

Delete the rest association rules that contain all symbols in the first one;

Until No association rules;

Recover the production rules from transaction database;

End.

Illustrated Example

Assume the following rule-based expert system :

If A and B and C then X
If A and B and C and D then X
If A and B and C and E then Y
If B and C and D then Y
If B and C and D and E then Z
If C and D and E then X
If E and X then Q1
If E and Y then Q2
If X and Y then Q1
If E and X and Y then Q2

Step 1: Convert these rules into transaction tables as follows:

Transaction ID	Transaction
1	A, B, C
2	A, B, C, D
3	A, B, C, E
4	B, C, D
5	B, C, D, E
6	C, D, E
7	E, X
8	E, Y
9	X, Y
10	E, X, Y

Step 2: The association rules from applying improved *Apriori* algorithm are:

A $\longrightarrow$ B
A $\longrightarrow$ C
A $\longrightarrow$ BC
B $\longrightarrow$ C
AC $\longrightarrow$ B
D $\longrightarrow$ C
AB $\longrightarrow$ C
BD $\longrightarrow$ C

Step 3: Sort the above association rules descending:

A $\longrightarrow$ BC

$AC \longrightarrow B$
 $AB \longrightarrow C$
 $BD \longrightarrow C$
 $A \longrightarrow B$
 $A \longrightarrow C$
 $B \longrightarrow C$
 $D \longrightarrow C$

Step 4: Repeat replacing and deletion operations. According to the first association rule $A \longrightarrow BC$, replace each A, B and C with P1, then delete each association rule which contains A, B and C. The resulting transaction database :

Transaction ID	Transaction
1	P1
2	P1, D
3	P1, E
4	B, C, D
5	B, C, D, E
6	C, D, E
7	E, X
8	E, Y
9	X, Y
10	E, X, Y

The rest of association rules are :

$BD \longrightarrow C$
 $D \longrightarrow C$

Now, the first association rule $BD \longrightarrow C$, replace each D, B and C with P2, then delete each association rule which contains D, B and C. The resulting transaction database :

Transaction ID	Transaction
1	P1
2	P1, D
3	P1, E
4	P2
5	P2, E
6	C, D, E
7	E, X
8	E, Y
9	X, Y
10	E, X, Y

The rest of association rules are :

$$D \longrightarrow C$$

Now, the first association rule $D \longrightarrow C$, replace each D and C with P3, then delete each association rule which contains D and C. The resulting transaction database :

Transaction ID	Transaction
1	P1
2	P1, D
3	P1, E
4	P2
5	P2, E
6	P3, E
7	E, X
8	E, Y
9	X, Y
10	E, X, Y

At this step there are no association rules.

Step 5: Recover the production rules from resulting database transaction:

If P1 then X
If P1 and D then X
If P1 and E then Y
If P2 then Y
If P2 and E then Z
If P3 and E then X
If E and X then Q1
If E and Y then Q2
If X and Y then Q1
If E and X and Y then Q2

*Applying the suggested algorithm reduces the premises **from 30 to 19** !!! .*

7- Conclusion

The premises of rule-based expert system play a big role in the search engine complexity. In this paper, a logical method is presented to reduce these premises using association rules mining in data mining technique. The presented method has proved a good logical tool for reducing the premises because it depends on the logical relations among these premises. Association rules mining using Apriori algorithm will produce logical association rules in the transactions, therefore, this paper applies Apriori algorithm.

References

- [1] : Abbas A. Al-Bakry, “**Expert Systems Development Using Knowledge Agent**”, Ph. D. Thesis, University of Technology, Baghdad, Iraq, 2003.
- [2] : Amit Konar, “ **Artificial Intelligence and Soft Computing : Behavioral and Cognitive Modeling of the Human Brain**”, CRC Press, 2000.
- [3] : Buchanan, B. G. and Feigenbaum, E. A., “**DENDRAL and Meta DENDERAL: Their Applications dimension**”, Artificial Intelligence, Vol. 11, pp. 5-24, 1978.
- [4] : Daniel H. Marcellus, “**Expert Systems Programming in Turbo Prolog**”, Prentice-Hall Inc., 1989.
- [5] : Jackson, P., “**Introduction to Expert Systems**”, 2nd edition, Addison-Wisely, 1990.
- [6] : Jiawei Han and Micheline Kanmber, “**Data Mining Concept and Techniques**”, Academic Press, USA, 2001.
- [7] : Klein, M. and Methlie, L. B., “**Knowledge Based Decision Support Systems with Applications in Business**”, 2nd edition, JohnWiley & Sons, 1995.
- [8] : George F. Luger and William A. Stubblefield, “**Artificial Intelligence Structures and Strategies for Complex Problem Solving**”, 3rd edition, Addison Wesley Longmann Inc., 1998.
- [9] : Nilson, N. J., “**Principles of Artificial Intelligence**”, Morgan Kaufmann, San Mateom, CA, pp 6-7, 1980.

- [10] : Rakesh Agrawal, “**Mining Sequential Patterns**”, Proc. of the Int’l Conference Data Mining Engineering (ICDE95), Taiwan, March 1995.
- [11] : Rakesh Agrawal and Ramakrishnan Strikant, “**Fast Algorithms for Mining Association Rules**”, Proc. of the 20th Int’l Conference on Very Large Databases, Santiago, Chile, 1994.
- [12]: Rakesh Agrawal, Tomasz Imielinski and Arun Swami, “**Mining Association Rules Between Sets of Items in Large Database**”, Washington, U.S.A., May 1994.
- [13] : Rakesh Agrawal, Heikki Mannila, Ramakrishnan Strikant, Hannu Toivonen and Inkeri Verkamo, “**Fast Discovery of Association Rules**”, Proc. of the 20th Int’l Conference on Very Large Databases, Santiago, Chile, 1994.
- [14] : Zupan, B. and Cheng, A. M. K., “**Optimization of Rule-Based Systems Using State-Space Graphs**”, IEEE Trans. On Knowledge and Data Eng., Vol. 10, No 2, March/April 1998.

تقليص المسببات المنطقية في الانظمة الخبيرة باستخدام تقنية القواعد المقترنة

د.احمد طارق صادق

قسم علوم الحاسبات

الجامعة التكنولوجية

الملخص

قاعدة المعرفة تعتبر بمثابة قلب النظام الخبير وهي التي تحدد قوة النظام الخبير ومجال محرك البحث المنطقي. وواحدة من هذه القواعد هي قواعد الانتاج. تعتبر المسببات المنطقية لهذه القواعد بمثابة القلب لها، لذلك فتقليصها يكون مفيد جداً لغرض تقليص وقت ومساحة البحث. هذا البحث يقدم نموذج لتقليص المسببات المنطقية في قواعد الانتاج باستخدام خوارزميات القواعد المقترنة وبالاخص خوارزمية Apriori. تطبيق النموذج أعطى تقليصاً جيداً ومنطقياً المسببات المنطقية والذي سيلعب دوراً كبيراً في محرك البحث المنطقي.