

Discovery of Classification Rules

湯上, 伸弘
九州大学システム情報科学研究科情報理学専攻

<https://doi.org/10.11501/3180411>

出版情報：九州大学，2000，博士（理学），課程博士
バージョン：
権利関係：



①

Discovery of Classification Rules

Nobuhiro Yugami

February 2001

Abstract

This thesis investigates discovery of classification rules from databases for classifying new data and for analyzing properties of the databases. Typical databases involve huge number of rules and the key issues of rule discovery are how to evaluate interestingness of rules and how to find interesting rules efficiently.

In classification task, the desired rules are the ones that are as simple as possible and can classify new instances correctly. Recent researches on supervised learning show learning plural rule sets can significantly improve classification ability while it degrades the simplicity of the total rule sets. This thesis presents a discovery algorithm IGR(Instance Guided Rule induction) that achieves high classification ability without degrading simplicity by inducing one pseudo-optimum rule for each training instance.

In analysis of the properties of the databases, most previous researches use support and accuracy for selecting rules and tend to extract many statistically trivial rules. To resolve this difficulty, this thesis presents new criteria of rules that evaluate a rule by whether every condition in its body is really required to imply its conclusion. Moreover, this thesis presents two algorithms to extract interesting rules based on the criteria, DIG(Discovery of Interesting rules with Grouping attribute values) for nominal databases and IIS(Iterative Interval Specialization) for numeric databases.

The contributions of this thesis are summarized as follows. First, this thesis clarifies what kind of rules should be discovered and gives new criteria of rules to filter out less interesting rules. Second, this thesis presents three efficient discovery algorithms for classifying new instances and for analyzing properties of databases. Finally, this thesis empirically demonstrates the performance of the algorithms.

Acknowledgements

I would like to thank Professor Setsuo Arikawa of Kyushu University for supervising this thesis. He guided my research and gave many important and helpful advice for this thesis.

I am deeply grateful to researchers in Fujitsu Laboratories for thoughtful comments and valuable discussions, especially to Seishi Okamoto, Hiroya Inakoshi and Yuiko Ohta. Many results in this thesis are joint works with them.

I indebted to members of the seminar at department of informatics of Kyushu University for their helpful comments and advice, especially to Hiroki Arimura, Hiroki Ishizaka, Ayumi Shinohara, Takashi Shinohara, Hiroshi Sakamoto and Masayuki Takeda.

I also grateful to Hiromu Hayashi, Kazuhiro Matsuo, Fumihiro Maruyama, Ryusuke Masuoka, Hirotaka Hara for their strong support to my research at Fujitsu Laboratories and for giving a chance to research at Kyushu University.

Finally, all love and gratitude to my wife, Masayo Yugami, for her devotion and invaluable assistance during this research.

Contents

1	Introduction	1
1.1	Rule Discovery for Prediction	1
1.2	Rule Discovery for Extracting New Knowledge	5
1.3	Objective of This Thesis	6
1.4	Overview of This Thesis	6
2	Rule Discovery for Classification	9
2.1	Learning Hypotheses for Classification	9
2.2	Discovery Algorithm:IGR	10
2.2.1	Overview of IGR	10
2.2.2	IGR	12
2.2.3	Condition Selection for Specialization	14
2.2.4	Weights of Rules	16
2.2.5	Example	18
2.3	Experiments	22
2.4	Chapter Summary	28
3	Discovery of Interesting Rules in Nominal Domains	31
3.1	Classification Rules in Nominal Domains	31
3.2	Interestingness of Classification Rules	32

3.2.1	Limitation of Rule Discovery based on Support and Accuracy	32
3.2.2	Interestingness of Primitive Rules	34
3.2.3	Interestingness of Rules with Grouping Attribute Values	37
3.3	Discovery Algorithm: DIG	39
3.4	Time Complexity	42
3.5	Experiments	43
3.6	Chapter Summary	48
4	Discovery of Interesting Rules in Numeric Domains	51
4.1	Classification Rules in Numeric Domains	51
4.2	Difficulty of Rule Discovery in Numeric Domains	52
4.3	Discovery Algorithms	53
4.3.1	SAB: Rule Discovery by Selecting Attribute Boundaries	54
4.3.2	IIS: Rule Discovery by Iterative Interval Specialization	57
4.4	Experiments	61
4.4.1	Time Complexity	61
4.4.2	Discovered Rules	66
4.5	Chapter Summary	67
5	Related Works	70
5.1	Learning Decision Rules	70
5.2	Evaluation of Rules	74
5.3	Discovery of Exception Rules	76
5.4	Rule Discovery in Numeric Domains	79
6	Conclusions	82
6.1	Contributions	82

6.2	Future Works	83
References		85

Chapter 1

Introduction

A key element of an intelligent system is the ability to discover new knowledge. This ability has been researched in both of machine learning and data mining. Knowledge discovery is a complex task consists of many kinds of processes, gathering various kinds of data from distributed resources, cleaning gathered data by editing them (semi-) automatically, combining the data to generate a database to be analyzed, analyzing the database, and presenting the results of analysis to a user[8, 24, 39, 57, 67]. Even if all of these processes play quite important roles in knowledge discovery, the central issue of knowledge discovery is the process of data analysis. Many paradigms have been proposed for data analysis[11, 36] depending on what kind of knowledge is discovered from what kind of data and depending on how the discovered rules are used for. This thesis focuses on discovery of classification rules from labeled instances for predicting classes of new instances accurately and for extracting non-trivial regularities in the given instances.

1.1 Rule Discovery for Prediction

Many applications of knowledge discovery are prediction tasks from prior instances[67]. It is an active subfield of machine learning to learn hypotheses from a set of labeled instances to classify new instances. Many learning paradigms have been proposed to

Table 1.1: An example of training instances. In this example, each instance is represented by its class label and values for three attributes, *Made_in*, *Size* and *Price*.

ID	class	<i>Made_in</i>	<i>Size</i>	<i>Price</i>
1	positive	<i>japan</i>	<i>small</i>	1,000,000
2	positive	<i>japan</i>	<i>small</i>	1,500,000
3	positive	<i>japan</i>	<i>small</i>	2,500,000
4	positive	<i>japan</i>	<i>middle</i>	1,200,000
5	positive	<i>japan</i>	<i>middle</i>	3,000,000
6	negative	<i>japan</i>	<i>big</i>	1,800,000
7	positive	<i>europe</i>	<i>small</i>	1,900,000
8	negative	<i>europe</i>	<i>small</i>	2,800,000
9	negative	<i>europe</i>	<i>middle</i>	1,000,000
10	negative	<i>europe</i>	<i>middle</i>	3,500,000
11	negative	<i>europe</i>	<i>big</i>	1,400,000
12	positive	<i>europe</i>	<i>big</i>	2,000,000

learn various types of hypotheses from various types of instances[41]. As the representation of the instances, this thesis assumes attribute-value representation in which each instance is represented by its class label and values of predefined attributes. This representation is simple but sufficiently powerful in many practical domains and is used in many learning systems to learn decision trees[51, 10], decision rules[38, 15, 17], neural networks[31], and nearest neighbor classifiers[19, 47]. Table 1.1 is an example of the representation of training instances. In this table, each row represents one instance and each column shows the values for one attribute.

This thesis focuses on discovery of decision rules, a set of classification rules, from the training instances with attribute-value representation. A classification rule is a kind of if-then rule that constrains values of attributes in its body and classifies the instances into the class in its conclusion. Following is an example of classification rule learned from the training instances in table 1.1.

$$Size \in \{small, middle\} \wedge Price \leq 1,800,000 \Rightarrow class = positive.$$

This rule means that an instance belongs to class *positive* if it satisfies the two conditions on *Size* and *Price*. This thesis is only concerned with classification rules whose bodies are conjunctions of conditions that constrain values of single attributes.

In rule discovery for prediction, a hypothesis is a set of discovered rules and a new instance is classified with the whole set of discovered rules. Then, discovered rules are evaluated as a set of the rules.

Learning a hypothesis from a set of training instances is a kind of optimization problem because learning is a task to select the best hypothesis with respect to a certain criterion among huge number of hypotheses[43]. This optimization problem is quite difficult because of the following two reasons. First difficulty is computational intractability. If there are M attributes with D possible values, then the number of distinct instances is D^M and the number of logically distinct hypotheses becomes 2^{D^M} for each class because a hypothesis is defined by a set of instances covered by it. Second difficulty is how to evaluate hypotheses. In many application of machine learning techniques, the primary objective of the learned hypothesis is to classify new, unknown instances correctly and the best hypothesis is the one with the highest classification accuracy. However, learning algorithms only know the given training instances and can not evaluate exact accuracy for new instances. They can only predict the accuracy for new instances from the accuracy for known instances[4]. Strict optimization with such an uncertain evaluation causes over-fitting with the training instances[46, 52].

Most learning algorithms do not try strict optimization because of the above two difficulties. They apply more computationally light methods for the main process of learning and try optimization in pre-processes or post-processes of learning such as feature subset selection[12, 35] and pruning decision trees[18, 42] and decision rules[16]. Some researches tried strict optimization but were effective for only a few databases[29,

46, 70].

In addition to classification performance, simplicity of learned hypotheses is another viewpoint of hypotheses evaluation. Simplicity becomes important when human experts try to understand learned hypotheses and to modify them based on their domain knowledge. In such a case, simple hypotheses with small number of rules are easy to understand and are more preferable than complex hypotheses.

It is well known that if two hypotheses have similar classification ability on the training instances, then the simpler one works better for unknown instances[7]. This heuristic is known as "Occam's razor". Most learning algorithms use Occam's razor explicitly or implicitly. For example, Quinlan and Rivest [54] apply minimum description length (MDL) principle[55] to decision tree induction that minimizes the total description length of a tree itself and descriptions for exceptional training instances miss-classified by the tree. However, even if simpler hypotheses usually achieve high accuracy, these two criteria evaluate different properties of hypotheses and don't need to coincide with each other.

Recent researches show counter examples of Occam's razor. For example, bagging[9] and boosting[58] learn plural hypotheses with weak learners and combine them as a whole hypothesis to classify new instances. Then, the whole hypothesis is more complex than individual hypotheses. However, both of empirical[25, 53] and theoretical analyses[59] show their abilities to improve classification accuracy of the existing learning algorithms. A simple hypothesis is more preferable than complex ones because of its simplicity itself, not because it has higher classification ability[21].

1.2 Rule Discovery for Extracting New Knowledge

The purpose of rule discovery is sometimes not to predict for new instances but to extract interesting regularities to analyze the characteristics of the given databases. Association rules discovery[1] is an example of this type of rule discovery.

In discovery of association rules, each training instance is a set of items and the purpose is to analyze relationships between items. Basket analysis is a typical application of association rule discovery. In this analysis, an instance is a set of items bought by one person and the purpose of the analysis is to generate rules such as "if a customer buys beer then he/she also buys potato chips with high probability".

As in rule discovery for prediction, evaluation of rules is an important issue in discovery for analysis. However, there is a big difference between these two kinds of discovery on what should be evaluated. In prediction tasks, the purpose of discovery is to discover a set of rules, not a specific rule, to achieve high prediction performance and the discovered rules are evaluated as a set of them. Instead, in discovery for analysis, the purpose is to discover rules representing useful regularities and an individual rule is evaluated as itself.

In usual, discovery for analyzing databases is a kind of constraint satisfaction problems[65, 69, 72] to enumerate all rules that satisfy given constraints such as the maximum length of rule's body and thresholds on given criteria to evaluate rules. In association rules mining, most algorithms constrain rules to be discovered with lower bounds of their support and accuracy. However, it is not easy to set appropriate thresholds on support and accuracy. High thresholds lead only trivial rules that human experts already know and low thresholds lead too many rules. In addition, even if large support and high accuracy are important conditions of good rules, they are not sufficient conditions for good rules. Then, it is quite difficult to select useful rules

with only constraints on them and it is required to restrict rules with other criteria that represent usefulness or interestingness of rules directly[5, 37, 49].

1.3 Objective of This Thesis

The primary objective of this thesis is to construct practical discovery systems for classification rules. This objective consists of two problems. First problem is to clarify what kind of classification rules should be discovered. The answer to this question strongly depends on what the discovered rules used for. This thesis deals with two types of rule discovery, discovery for classifying new instances accurately, and discovery for extracting non-trivial regularities. The evaluation of rules is quite important in the second situation and this thesis proposes new criteria of rules for the situation.

Second problem is how to discover the desired classification rules effectively. In usual, there are huge number of possible rules even in a small database and the crucial problem of rule discovery is how to search desired rules in practical time. This thesis proposes effective discovery algorithms for many situations, discovery of rules to achieve high classification performance, discovery of interesting rules in nominal domains, and discovery of interesting rules in numeric domains.

1.4 Overview of This Thesis

This section provides the content of the remainder of this thesis.

Chapter 2 presents a rule discovery algorithm IGR(Instance Guided Rule induction) for classifying new instances correctly. The objective of this algorithm is to achieve high classification accuracy without increasing both in learning time and the number of learned rules. IGR achieves these requirements by learning one pseudo-optimum rule for each training instance. The pseudo-optimum rule means the rule with 100%

accuracy that covers as many positive instances as possible. The rule represents the widest area of the concluding class around the corresponding training instance. If a test instance to be classified is sufficiently near to a certain training instance, we can expect that the rule for the training instance is also pseudo-optimum for the test instance and the rule classifies the test instance correctly. This chapter also reports the experimental evaluation using databases from UCI machine learning repository[6] and compares IGR with other learning algorithms with respect to classification performance, learning speed, and the complexity of learned hypotheses.

Chapter 3 and 4 describe rule discovery for extracting non-trivial regularities from a given database and the purpose of discovered rules is not for classifying new instances but to analyze characteristics of the given database. As described above, huge number of rules are embedded in typical databases and the most important problem of rule discovery is how to filter out less interesting rules. This problem is not so serious in discovery of rules to classify new instances because discovery algorithms can reject rules that don't contribute to improve classification performance. However, even if a rule overlaps with other rules and doesn't improve classification accuracy at all, it may represent a different regularity from other rules by constraining values of different attributes and may be worth to be discovered for analyzing databases.

Chapter 3 discusses how to evaluate rules and presents new criteria of classification rules that requires a rule to be discovered is not a trivial conclusion from simpler rules. Chapter 3 also applies the criteria to rule discovery in nominal domains and presents a new rule discovery algorithm, DIG(Discover Interesting rules with Grouping attribute values), to discover rules with grouping attribute values[22, 44], i.e. rules that permit plural values for each attribute in their bodies. Attribute value grouping increases the representation power of rules and makes it easy to understand the meaning of

discovered rules.

Chapter 4 applies the criteria presented in Chapter 3 to numeric domains in which values of all attributes are numeric. This chapter focuses on a classification rule in which each condition in its body requires that a value of a certain attribute belongs to a specific interval. Because a slight difference of boundaries of the intervals does not affect the meaning of a discovered rule, this chapter assumes values of each attribute are discretized into dozens of values and tries rule enumeration from the discretized instances. Even if the discretization decreases the number of possible rules, one of the most important problems is also how to reduce the time complexity for rule enumeration. This chapter presents two discovery algorithms to enumerate rules in numeric domains and compares them empirically.

Chapter 5 surveys the literatures about knowledge discovery algorithms and discusses the relationships between them and the criteria of rules and the algorithms presented in this thesis.

Chapter 6 summarizes the contributions of this thesis and discusses future work.

Chapter 2

Rule Discovery for Classification

2.1 Learning Hypotheses for Classification

This chapter discusses learning hypotheses from pre-classified training instances to classify new instances correctly. This problem has been researched as supervised learning[11, 36] and many algorithms have been proposed. This thesis focuses on discovery of if-then rules, where a hypothesis to classify new instances is a set of all discovered rules. In supervised learning, the primary objective is to classify new instances correctly and the most important criterion of learned hypotheses is classification accuracy for new instances. In addition, simplicity of the learned hypotheses is also an important criterion of the hypotheses. This becomes crucial when human experts want to understand the meaning of learned hypotheses and to modify the hypotheses with their background knowledge.

These two criteria, classification performance and simplicity, don't coincide with each other. For example, bagging[9] and boosting[58] learn plural hypotheses and combine all of them to classify new instances. Both of empirical[25, 53] and theoretical[59] analyses show that bagging and boosting can significantly improve classification performance. However, they degrade simplicity of the total complexity of the hypotheses. The objective of this chapter is to achieve the same classification performance with

Table 2.1: An examples of training instances. Values for *Price* are discretized into three values, *cheap*, *middle* and *expensive*.

ID	class	<i>Made_in</i>	<i>Size</i>	<i>Price</i>
1	positive	<i>japan</i>	<i>small</i>	<i>cheap</i>
2	positive	<i>japan</i>	<i>small</i>	<i>middle</i>
3	positive	<i>japan</i>	<i>small</i>	<i>expensive</i>
4	positive	<i>japan</i>	<i>middle</i>	<i>cheap</i>
5	positive	<i>japan</i>	<i>middle</i>	<i>expensive</i>
6	negative	<i>japan</i>	<i>big</i>	<i>middle</i>
7	positive	<i>europe</i>	<i>small</i>	<i>middle</i>
8	negative	<i>europe</i>	<i>small</i>	<i>expensive</i>
9	negative	<i>europe</i>	<i>middle</i>	<i>cheap</i>
10	negative	<i>europe</i>	<i>middle</i>	<i>expensive</i>
11	negative	<i>europe</i>	<i>big</i>	<i>cheap</i>
12	positive	<i>europe</i>	<i>big</i>	<i>middle</i>

bagging and boosting without loss of the simplicity of the discovered hypotheses.

2.2 Discovery Algorithm:IGR

2.2.1 Overview of IGR

This section presents a new discovery algorithm IGR(Instance Guided Rule induction) to learn a set of if-then rules to classify new instances correctly. The input of this algorithm is a set of training instances represented by class labels and values of pre-defined attributes. IGR can only deal with nominal attributes and requires values of numeric attributes to be discretized before learning. Table 2.1 shows an example of a set of training instances. The instances in this table are the same instances in Table 1.1 in Chapter 1 but the values of the last attribute, *Price*, are discretized into three values, *cheap*, *middle* and *expensive*. The output is a set of if-then rules with weights for majority voting such as

$$Made_in = japan \wedge Size \in \{small, middle\} \Rightarrow class = positive : 3.0,$$

$Size = small \wedge Price \in \{cheap, middle\} \Rightarrow class = positive : 1.5,$

$Price \in \{cheap, expensive\} \Rightarrow class = negative : 4.0,$

When classifying a new instance, IGR checks whether the new instance satisfies a body of each rule and returns the class that maximizes the sum of the weights of satisfied rules. For example, the instance with $Made_in = japan$, $Size = small$ and $Price = cheap$ satisfies bodies of all of the above three rules. Then, the score for class *positive* becomes $3.0 + 1.5 = 4.5$ and the score for *negative* is 4.0. As a result, the instance is classified into *positive*.

As discussed in the previous section, the objective of this algorithm is to achieve high classification performance without increasing the number of learned rules. The basic idea of IGR is to learn the optimal rule for each training instance. The optimal rule for a instance i belonging to class c is the rule that covers as many instances in class c including i as possible and rejects all instances in other classes. The optimal rule represents the widest area of class c around the instance i and we can expect that the rule can classify new instances near i correctly. Because it is computationally difficult to induce the strict optimal rules, IGR induces a pseudo-optimal rule for each training instance by greedy search and achieves high classification performance by combining such locally pseudo-optimal rules.

When inducing a rule for a target instance i , IGR applies general-to-special greedy search. It starts from a rule with no condition and iteratively specializes it by adding a new condition with the highest information gain among all of possible conditions that accept the target instance. IGR doesn't try backtracking and never removes the selected condition. IGR stops specialization when the rule covers only positive instances. Positive instances are the training instances belong to the same class of the

target instance and negative instances are the training instances in other classes.

The trivial method to learn one rule for each instance is to learn each rule one by one. However, IGR specializes rules for many target instances at once to accelerate learning by sharing condition evaluation processes to select the best conditions to specialize the rules. For example, when selecting the first condition of the rules, every condition has the same score for the rules for all target instances in the same class and IGR doesn't need to iterate evaluation of the conditions for the first specialization for each rule. If the rules for plural target instances select the same condition at the first specialization, then the rules can share the condition evaluation process to select the second conditions for specialization. The complexity of condition evaluation process is the number of conditions times the number of instances covered by the current body because the process requires to count a frequency of instances for each combination of an attribute value and a class label in the instances that satisfy the current body. This process must be done for every condition for every rule and is the most time consuming process in learning classification rules.

2.2.2 IGR

Figure 2.1 shows a pseudo-code of IGR that is a recursive procedure to learn pseudo-optimum rules for a certain class. The inputs of the algorithm are a set of target instances for which pseudo optimal rules should be induced, *Targets*, a current body of the rules for the target instances in *Targets*, *Body*, and the weights of instances to avoid to induce many similar rules. The output of the algorithm is a set of pseudo-optimum rules for the training instances in *Targets*, *Rules*. IGR is first called with all positive instances in *Targets* and with no condition in *Body* ($Body = true$) to induce rules for all positive instances belong to the target class c . The weights of instances are initialized to 1 for all instances. IGR iteratively specializes the body *Body* by adding an

procedure IGR(*Targets*, *Body*, *Weight*)

inputs:

Targets : set of target instances belong to a target class *c*.

Weight : weights of training instances.

Body : current body (conjunction of conditions) for target instances.

output:

Rules : set of discovered rules.

1. *Rules* := ϕ .
2. *Pos* := set of positive instances that satisfy *Body* and belong to a target class *c*.
3. *Neg* := set of negative instances that satisfy *Body* and belong to a target class *c*.
4. **If** *Neg* = ϕ **then**
5. *Rules* := $\{Body \Rightarrow c\}$.
6. **else**
7. Generate all conditions to specialize *Body* and evaluate them with weighted information gain in the set of instances $Pos \cup Neg$.
8. **Foreach** target instance $i \in Targets$
9. Select a condition that accepts i and maximizes information gain.
10. $I(t) := \{i | \text{condition } t \text{ is selected for } i\}$.
11. **Foreach** t such that $I(t) \neq \phi$ (decreasing order of information gain)
12. *Rules* := *Rules* $\cup$ IGR($I(t)$, $Body \wedge t$, *Weight*).
13. **Foreach** $p \in Pos$
14. $Weight[p] := \alpha \cdot Weight[p]$. ($\alpha \leq 1$)
15. **Return** *Rules*.

Figure 2.1: Discovery algorithm IGR. IGR learns rules for the training instances in *Targets* and IGR is first called with all positive instances in *Targets*.

appropriate condition until *Body* covers only positive instances (step 3). If *Body* rejects all negative instances, IGR returns a rule $Body \Rightarrow c$. Otherwise, IGR evaluates all of possible conditions to specialize *Body* and selects the best condition for each target instance in *Targets*. IGR uses weighted information gain for evaluating conditions. The next subsection will explain this criterion for condition selection. For all target instances that select the same condition t , $I(t)$, IGR can evaluate the conditions for the next specialization at once and calls IGR recursively with $Targets = I(t)$. IGR updates the weights of instances to avoid to induce many similar rules. IGR decreases

the weights of instances accepted by the selected condition and focuses on the instances rejected by the condition in the discovery process with other conditions (step 10~12).

The number of rules induced by IGR is at most the number of training instances and is usually smaller than that because plural training instances share one rule as their pseudo-optimal rule. This happens when the condition in step 2 in Figure 2.1 is satisfied with plural target instances in *Targets*.

IGR is similar to AQ family that also learns one rule for one seed instance. The most important difference between them is that IGR induces redundant rules to classify training instances correctly. After learning one rule, AQ removes all positive instances covered by the rule and induces the next rule for the remaining instances. The resulting rule sets cover all of positive instances and can classify all training instances correctly. Instead, IGR does not remove instances covered by the generated rules and induces one rule for every training instance. As a result, IGR learns more rules than AQ and some of the rules from IGR may not be required to explain training instances correctly. However, these redundant rules improve classification performance for new, unknown instances as the experimental results in section 2.3 will show. Another difference is that AQ induces rules one by one but IGR induces plural rules at once to reduce running time.

2.2.3 Condition Selection for Specialization

For general-to-special rule induction like IGR, one of key issues of the algorithm is how to select conditions for specializing bodies of rules. When specializing a rule, the simplest and the most traditional way is to add a condition that accepts instances with a certain attribute value. However, this type of condition tends to generate too specific rules when each attribute has many possible values. To avoid over specialization, IGR divides values of one attribute into two groups and uses a condition that accepts values

in one group. When inducing rules for class c , IGR divides values based on whether each value increases the probability of the target class c or not as follows.

$$G^+(x, c, Body) = \{v | v \in Domain(x) \wedge P(c|x = v \wedge Body) > P(c|Body)\},$$

$$G^-(x, c, Body) = \{v | v \in Domain(x) \wedge P(c|x = v \wedge Body) \leq P(c|Body)\},$$

where x and $Domain(x)$ is an attribute and a set of its possible values. Therefore, IGR generates two conditions, $x \in G^+(x, c, Body)$ and $x \in G^-(x, c, Body)$, as candidates for specialization for each attribute x at each specialization step. For simplicity, I use $x = v$ to denote $x \in G^+(x, c, Body)$ ($x \in G^-(x, c, Body)$) if $G^+(G^-)$ involves only one value v . IGR fixes the value groups at each specialization step and doesn't modify the combination of the permitted values for the selected conditions afterwards. Instead, IGR may change the permitted range of the attribute by adding a new condition on the values of the same attribute.

IGR uses weights of instances to avoid inducing many similar rules and counts probabilities by replacing frequencies of instances with the sums of the weights of the instances. For example, $P(c|Body)$ becomes the ratio of the sum of the weights of the instances that satisfy both of $class = c$ and $Body$, on the sum of the weights of the instances that satisfy $Body$. IGR updates the weights of training instances with a dumping factor α at step 12.

IGR evaluates each condition with information gain[51] with small modification. Information gain of a condition t is defined as

$$\begin{aligned} IG(t, c, Body) &= E(P(c|Body)) \\ &- P(t \wedge Body) \cdot E(P(c|t \wedge Body)) \\ &- P(\neg t \wedge Body) \cdot E(P(c|\neg t \wedge Body)), \end{aligned} \quad (2.1)$$

where

$$E(p) = -p \cdot \log_2 p - (1 - p) \cdot \log_2 (1 - p).$$

The probabilities in (2.1) are also counted with weighted frequencies of the instances. Information gain becomes maximum if the condition can completely divide instances into the set of positive instances and the set of negative instances. Oppositely, information gain becomes zero when the condition does not change the probability of the target class at all.

Information gain evaluates a division of instances into two sets and gives the same score to a condition t and its negation $\neg t$. In other words, it can't distinguish t and $\neg t$ and can't judge which is better for specialization. It is not important in decision tree learning because sub trees grow for both accepted and rejected instances in decision trees. However, decision rule learner focuses on only accepted instances and requires to distinguish the conditions and their negations. For example, if t permits most positive and few negative instances and gets high information gain, then $\neg t$ rejects most of positive instances and remains many negative instances. In such a case, it is good to specialize a current body with t but is stupid to use $\neg t$. Therefore, IGR checks whether the condition increases the probability of target class c and uses negative of information gain as a score of the condition if the probability decreases. As a result, IGR uses the following evaluation function for a condition t to specialize a body $Body$ when inducing rules for a class c .

$$Est(t, c, Body) = \begin{cases} IG(t, c, Body) & \text{if } P(c|t \wedge Body) \geq P(c|Body) \\ -IG(t, c, Body) & \text{otherwise} \end{cases} \quad (2.2)$$

2.2.4 Weights of Rules

IGR learns a set of if-then rules and classifies new instances with the rules. There are two common problems to classify a new instance by a set of rules, how to classify

an instance that satisfies no rule and how to classify an instance that satisfies plural rules. In the former case, IGR returns the majority class in the training set. In the later case, IGR selects a class by majority voting with the following weights of rules.

$$weight(r) = \sum_{i \in I(r)} \frac{1}{|S(i)|}, \quad (2.3)$$

where $I(r)$ is the set of training instances covered by a rule r and $S(i)$ is the set of rules that cover an instance i . This definition of weights assigns large weight to an isolated rule that covers many training instances not covered by other rules. Oppositely, even if a rule has large support and covers many instances, its weight may become small if it covers only instances covered by many other rules.

From a statistical viewpoint, a rule becomes more reliable as it covers more positive instances and less negative instances. Because IGR generates rules that cover only positive instances, the simplest weights of rules are their support, the number of positive instances covered by them. The recent version of CN2[14] sets weights of rules based on this approach. When classifying a new instance, it sums up class distributions of rules satisfied by the instance to be classified and classifies the instance into the class with maximum score. Because IGR induces rules covering only positive instances, using the support of rules as their weight coincide with the rule weighting in CN2. If the two rules with different class labels cover an instance to be classified, then it is natural to classify the instance into the class of the rule with larger support. However, using support as weights of rules causes a problem when the instance is covered by plural rules for each class. Intuitively, when plural rules with one class covers the instance to be classified, classification of the instance into the class is supported by the training instances covered by at least one of the matched rules. Then, if the learned rules don't overlap at all and each training instance is covered by only one rule, it is adequate to use the sum of the support of the rules as the score of the class. However, there may

be many rules that cover quite similar sets of the training instances and conclude the same class. If many similar rules cover an instance to be classified, the sum of their support becomes much larger than the number of training instances covered by them and the score becomes too high. In such a case, the training instances covered by many rules affect the results of classification of new instances much stronger than the instances covered by a few rules. A solution to this problem is to keep a set of covered training instances for each rule and to select the class that maximizes the size of the union of the instance sets of the matched rules for the class. However, this makes classification too slow when training sets are huge. Instead of calculating the union, IGR first counts how many rules cover each training instance and assigns a weight to a rule by summing the inverse of the number of rules for each training instance covered by the rule as in equation (2.3). This definition of the weight directly leads the following property.

$$\forall i, \sum_{r \in S(i)} weight(r) = 1.$$

This property means that every training instance gives the same contribution to the weights of rules.

2.2.5 Example

This subsection describes how IGR learns classification rules by using the training set shown in Table 2.1 as an example. To discover rules for a class positive, IGR is called with a set of all positive instances $\{1, 2, 3, 4, 5, 7, 12\}$ as *Targets* and *Body* = *true*. Then, *Pos* and *Neg* in step 1 in Figure 2.1 are the sets of all positive and all negative instances respectively. The weights of instances are set to 1 for all instances. Because *Neg* includes four negative instances and doesn't satisfy the condition in step 2, IGR counts class distribution for each attribute value and generates possible conditions by

grouping attribute values (step 5). For example, class distributions for values of *Size* are as follows.

	positive	negative
<i>Size</i> = <i>small</i>	4	1
<i>Size</i> = <i>middle</i>	2	2
<i>Size</i> = <i>big</i>	1	2
total	6	4

IGR divides values of each attribute into two sets by whether the probability of positive instances increases or not. Because only *Size* = *small* increases the probability of positive instances and other two values decreases the probability, IGR generates two conditions for the attribute, *Size* = *small* and *Size* $\in$ {*middle*, *big*}. Information gain is 0.10 for dividing instances in *Pos* and *Neg* by whether *Size* = *small* or *Size* $\in$ {*middle*, *big*}, and the scores for these two conditions become 0.10 and -0.10 respectively. IGR also generates conditions on other two attributes. As a result, IGR generates the following six conditions for specialization.

condition	score
<i>Made_in</i> = <i>japan</i>	0.20
<i>Size</i> = <i>small</i>	0.10
<i>Price</i> = <i>middle</i>	0.04
<i>Size</i> $\in$ { <i>middle</i> , <i>big</i> }	-0.04
<i>Price</i> $\in$ { <i>cheap</i> , <i>expensive</i> }	-0.10
<i>Made_in</i> = <i>europe</i>	-0.20

For each target instance in *Targets*, IGR selects the condition with the highest score among the conditions that accept the target instance (step 7). The condition *Made_in* = *japan* has the highest score and IGR selects this condition for the target instances that satisfy it, i.e. for instances 1, 2, 3, 4 and 5. Target instances 7 and 12 don't satisfy this condition and IGR selects *Size* = *small* for the instance 7 and *Price* = *middle* for the instance 12. Then, IGR calls IGR recursively for these three conditions in descending

order of their scores (step 9 and 10).

IGR first selects the condition *Made_in* = *japan* and calls IGR with *Targets* = {1, 2, 3, 4, 5} and *Body* = (*Made_in* = *japan*). The new body *Made_in* = *japan* covers the following five positive instances and one negative instance.

ID	class	<i>Made_in</i>	<i>Size</i>	<i>Price</i>
1	positive	<i>japan</i>	<i>small</i>	<i>cheap</i>
2	positive	<i>japan</i>	<i>small</i>	<i>middle</i>
3	positive	<i>japan</i>	<i>small</i>	<i>expensive</i>
4	positive	<i>japan</i>	<i>middle</i>	<i>cheap</i>
5	positive	<i>japan</i>	<i>middle</i>	<i>expensive</i>
6	negative	<i>japan</i>	<i>big</i>	<i>middle</i>

Because the value of *Made_in* is fixed to *japan*, IGR generates four conditions on other two attributes to reject the negative instance and evaluates them as follows.

condition	estimation
<i>Size</i> $\in$ { <i>small</i> , <i>middle</i> }	0.65
<i>Price</i> $\in$ { <i>cheap</i> , <i>expensive</i> }	0.32
<i>Size</i> = <i>big</i>	-0.32
<i>Price</i> = <i>middle</i>	-0.65

All of the target instances in *Targets* satisfy the best condition *Size* $\in$ {*small*, *middle*} that covers only positive instances. IGR is called with *Body* = (*Made_in* = *japan* $\wedge$ *Size* $\in$ {*small*, *middle*}) and discovers the following rule.

$$\textit{Made_in} = \textit{japan} \wedge \textit{Size} \in \{\textit{small}, \textit{middle}\} \Rightarrow \textit{class} = \textit{positive}$$

As a result, IGR with *Body* = (*Made_in* = *japan*) returns a rule set with only this rule.

After the recursive call with *Body* = (*Made_in* = *japan*), IGR updates weights of instances with the given parameter α . Let us assume $\alpha = 0.5$. Then, the weights become 0.5 for instances 1~5 and 1.0 for others.

Next, IGR is called with $Targets = \{7\}$ and $Body = (Size = small)$. Because the instances 1,2,3,7 and 8 satisfy $Body$, the weighted class distribution for each attribute value becomes

	positive	negative
$Made_in = japan$	1.5	0.0
$Made_in = europe$	1.0	1.0
$Price = cheap$	0.5	0.0
$Price = middle$	1.5	0.0
$Price = expensive$	0.5	1.0

and IGR generates the following four conditions

condition	estimation
$Price \in \{cheap, middle\}$	0.46
$Made_in = japan$	0.16
$Made_in = europe$	-0.16
$Price = expensive$	-0.46

Because there is only one target instance 7 and this instance satisfies the first condition, $Price \in \{cheap, middle\}$, the recursive call occurs only for this condition and returns the following rule.

$$Size = small \wedge Price \in \{cheap, middle\} \Rightarrow class = positive.$$

At the end, IGR is called for the target instance 12 with $Body = (Price = middle)$ and returns

$$Made_in = europe \wedge Price = middle \Rightarrow class = positive.$$

As a result, IGR learns the following three rules for class *positive* from the training instances shown in Table 2.1.

$$R_1 : Made_in = japan \wedge Size \in \{small, middle\} \Rightarrow class = positive,$$

$$R_2 : Size = small \wedge Price \in \{cheap, middle\} \Rightarrow class = positive,$$

$$R_3 : Made_in = europe \wedge Price = middle \Rightarrow class = positive.$$

Table 2.2: Characteristics of databases.

	classes	attributes	instances		classes	attributes	instances
breast	2	9	638	promoters	2	57	106
crx	2	15	653	satellite	7	36	6,435
dna	3	60	3,186	segment	7	19	2,310
glass	7	10	214	soy-large	19	35	307
hayes-roth	4	4	160	tic-tac-toe	2	9	958
iris	3	4	351	vehicle	4	18	846
krkp	2	36	3,196	voting	2	16	435
monks-1	2	6	432	waveform	3	21	5,000
pima	2	8	768	wine	3	13	178

After inducing rules, IGR assigns weights of rules. IGR requires a set of instances covered by a rule r , $I(r)$ and a set of rules that cover an instance i , $S(i)$ to calculate the weight of r . The first rule R_1 covers five instances, 1~5 and the instance 1 and 2 are also covered by R_2 but the remaining three instances are covered by only R_1 . Then, the weight for R_1 is calculated as follows.

$$\begin{aligned} weight(R_1) &= \sum_{i \in \{1,2,3,4,5\}} \frac{1}{|S(i)|} \\ &= \frac{1}{2} + \frac{1}{2} + \frac{1}{1} + \frac{1}{1} + \frac{1}{1} \\ &= 4.0. \end{aligned}$$

Similarly, the weights for R_2 and R_3 are

$$weight(R_2) = 2.0,$$

$$weight(R_3) = 1.5.$$

2.3 Experiments

This section compares IGR with AQ, C4.5 and two bagging learners with C4.5 by using eighteen databases in UCI machine learning repository[6]. AQ has many variants

to induce a rule for one seed instance. To show how IGR improves classification performance by learning more rules than AQ, this section uses AQ with the same conditions for specialization and the same condition selection criterion used in IGR. The bagging learners use C4.5 as their weak learner and learn five and twenty decision trees respectively.

Table 2.2 shows the domain characteristics of the databases used in the experiments. Because IGR can't deal with numeric attribute directly, ranges of numeric attributes were discretized into five intervals with same population for small databases with less than 500 instances and were discretized into ten intervals for larger databases with more than 500 instances. AQ also learned from discretized instances because it used the same conditions for specialization with IGR. Other learning algorithms, C4.5 and bagging with it, can deal with numeric attributes directly and learned rules from original, non-discretized instances. All experiments in this section rejected instances with missing attribute value and used only instances in which values for all attributes are given.

Before comparing IGR with other learning algorithms, I first investigate how weighting of instances with α affects discovered rule sets and how weights of rules in majority voting affects classification accuracy. Table 2.3 reports the effects of α on classification accuracy. The results in this table are the average of ten runs of five fold cross validation. In most databases, classification accuracy slightly decreases as α decreases but the differences were quite small when $\alpha \geq 0.05$. Then, from the viewpoint of classification accuracy, there is no reason to assign weights to the training instances.

However, weighting instances with α is important to learn smaller hypotheses with shorter learning time. Table 2.4 shows how α affects the hypothesis size, i.e. the number of rules, and learning time. For example, IGR with $\alpha = 0.1$ extracted about

Table 2.3: Dependency of classification accuracy on α . All results are the average of ten runs of 5-fold cross validation and this table also reports 95% confidence intervals of the average accuracy.

	$\alpha=1.00$	$\alpha=0.20$	$\alpha=0.10$	$\alpha=0.05$	$\alpha=0.00$
breast	95.9 $\pm$ 0.2	95.8 $\pm$ 0.1	95.7 $\pm$ 0.2	95.7 $\pm$ 0.2	95.2 $\pm$ 0.3
crx	86.3 $\pm$ 0.4	85.5 $\pm$ 0.5	85.3 $\pm$ 0.3	85.6 $\pm$ 0.4	83.8 $\pm$ 0.6
dna	96.0 $\pm$ 0.1	96.1 $\pm$ 0.1	95.9 $\pm$ 0.1	95.9 $\pm$ 0.1	94.9 $\pm$ 0.1
glass	69.2 $\pm$ 1.3	68.1 $\pm$ 1.3	67.9 $\pm$ 0.9	68.0 $\pm$ 1.1	68.2 $\pm$ 1.1
hayes-roth	82.8 $\pm$ 0.9	82.6 $\pm$ 0.9	82.6 $\pm$ 0.9	82.7 $\pm$ 1.0	82.3 $\pm$ 1.1
iris	90.7 $\pm$ 0.9	90.6 $\pm$ 0.8	90.7 $\pm$ 0.8	90.8 $\pm$ 0.9	91.1 $\pm$ 0.9
krkp	99.0 $\pm$ 0.1	99.2 $\pm$ 0.1	99.2 $\pm$ 0.1	99.3 $\pm$ 0.1	99.4 $\pm$ 0.1
monks-1	97.6 $\pm$ 1.1	96.6 $\pm$ 0.8	96.5 $\pm$ 1.1	96.5 $\pm$ 1.1	96.8 $\pm$ 0.9
pima	74.4 $\pm$ 0.6	74.9 $\pm$ 0.4	74.6 $\pm$ 0.5	74.1 $\pm$ 0.7	72.4 $\pm$ 0.8
promoters	87.5 $\pm$ 1.4	86.6 $\pm$ 1.3	84.9 $\pm$ 1.3	83.3 $\pm$ 1.7	80.1 $\pm$ 2.1
satellite	85.6 $\pm$ 0.1	86.7 $\pm$ 0.1	86.1 $\pm$ 0.2	85.7 $\pm$ 0.3	83.0 $\pm$ 0.2
segment	95.8 $\pm$ 0.2	96.1 $\pm$ 0.1	95.9 $\pm$ 0.1	95.8 $\pm$ 0.1	94.3 $\pm$ 0.2
soy-large	85.8 $\pm$ 0.8	85.8 $\pm$ 0.9	85.7 $\pm$ 0.7	85.3 $\pm$ 0.7	84.7 $\pm$ 0.8
tic-tac-toe	96.3 $\pm$ 0.5	97.8 $\pm$ 0.2	97.6 $\pm$ 0.2	97.4 $\pm$ 0.3	97.4 $\pm$ 0.3
vehicle	69.8 $\pm$ 0.6	69.4 $\pm$ 0.8	68.5 $\pm$ 0.9	68.3 $\pm$ 1.0	64.5 $\pm$ 1.2
voting	94.6 $\pm$ 0.3	94.7 $\pm$ 0.3	94.8 $\pm$ 0.3	94.7 $\pm$ 0.3	94.9 $\pm$ 0.2
waveform	82.6 $\pm$ 0.2	82.0 $\pm$ 0.2	81.0 $\pm$ 0.2	80.1 $\pm$ 0.2	74.9 $\pm$ 0.3
wine	93.3 $\pm$ 0.6	93.0 $\pm$ 0.6	92.7 $\pm$ 0.4	92.4 $\pm$ 0.5	91.0 $\pm$ 0.6

30% less rules in about 30% shorter learning time than IGR without weighting, i.e. $\alpha = 1$. This difference is not so large on average but the effects of the weights became significant in the databases in which IGR discovered many rules such as satellite and waveform. In these databases, IGR with weighting instances induced less than a half rules and was about three times faster than IGR without weighting. Then, weighting with small α is useful from the viewpoints of simplicity of learned hypotheses and the complexity for learning. In the following experiments in this section, α is fixed at 0.1.

Table 2.5 shows the effect of the weights of rules in majority voting defined in equation (2.3). This table compares three kinds of weights for rules in majority voting, constant weight for every rule, using support of rules as the weights of rules and the

Table 2.8: Average learning time (cpu seconds) of the three algorithms, C4.5, bagging with 20 runs of C4.5 and IGR.

	C4.5	bagging(20)	IGR($\alpha=0.1$)
breast	0.07	1.18	0.16
crx	0.67	8.72	0.34
dna	2.69	48.16	7.45
glass	0.40	5.48	0.11
hayes-roth	0.02	0.27	0.03
iris	0.02	0.33	0.02
krkp	1.31	26.03	2.18
monks-1	0.04	0.81	0.10
pima	1.55	18.91	0.37
promoters	0.09	1.39	0.08
satellite	57.79	875.45	23.21
segment	5.46	85.66	1.45
soy-large	0.30	5.36	0.33
tic-tac-toe	0.16	2.97	0.27
vehicle	2.51	36.76	0.77
voting	0.07	1.13	0.14
waveform	82.64	940.06	10.28
wine	0.13	2.04	0.04

portant because the simple rules are easy to understand. This chapter presented a new discovery algorithm IGR that achieves high classification performance without increasing the number of rules. IGR learns the pseudo-optimum rule for each training instance that covers as many positive instances as possible and rejects all negative instances.

Empirical results show that IGR can learn much more accurate rule sets than AQ and C4.5 in many databases. IGR is comparable to bagging with C4.5 with respect to classification accuracy but the size of rule sets from IGR is quite smaller than the sets from bagging. In addition, IGR is much faster than bagging learners.

Even if IGR achieved high accuracy, there are some remaining problems to be im-

proved. The most important problem is pruning[16, 28] learned rules. The current IGR doesn't apply pruning and generates only 100% accurate rules. The lack of pruning tends to lead over specialized rules especially in noisy databases and degrades the ability of classification. In addition, it increases the number of rules and degrades readability of learned rules.

Chapter 3

Discovery of Interesting Rules in Nominal Domains

3.1 Classification Rules in Nominal Domains

This chapter investigates rule discovery for analyzing characteristics of given databases. This chapter discusses what kind of rules are useful and presents the interestingness of rules as the criterion of classification rules. This chapter also presents a new discovery algorithm based on the interestingness. This chapter assumes all attributes are nominal and deals with classification rules with grouping attribute values[22, 44] like

$$x_1 \in D_1 \wedge \cdots \wedge x_L \in D_L \Rightarrow class = c, \quad (3.1)$$

where x_i is an attribute and D_i is a subset of possible values for x_i . This rule means if a value of a certain instance for attribute x_i is a member of D_i for every $1 \leq i \leq L$, then the instance belongs to class c . For simplicity, I will use $x_i = v$ instead of $x_i \in D_i$ if D_i includes only one value v . Many learning/discovery algorithms[51, 62] deal with only a rule without grouping attribute values that permits exactly one value for each attribute. This thesis calls such rules primitive rules. Grouping values increases the number of possible rules and makes rule discovery difficult. Instead, it increases readability of discovered rules because one rule with grouping can represent plural primitive rules.

3.2 Interestingness of Classification Rules

3.2.1 Limitation of Rule Discovery based on Support and Accuracy

This subsection shows why it is not sufficient to restrict support and accuracy to extract interesting rules. Support of a classification rule is the probability of instances that belong to the class in its head and satisfy all conditions in its body. The support of the rule 3.1 is

$$P(class = c \wedge x_1 \in D_1 \wedge \cdots \wedge x_L \in D_L). \quad (3.2)$$

Accuracy of a rule is the conditional probability of the class of its head on the condition that its body is satisfied,

$$P(class = c | x_1 \in D_1 \wedge \cdots \wedge x_L \in D_L). \quad (3.3)$$

If support of a rule is quite low, then the rule represents very rare case and it can be applied with very small probability. If accuracy of a rule is low, then we can not believe the results implied with the rule. Both of support and accuracy are fundamental characteristics of classification rules and it is natural to require large support and high accuracy for rules to be discovered.

Many discovery systems extract rules whose support and accuracy are higher than the given lower bounds[1]. However, the fixed lower bounds for support and accuracy are not sufficient to filter out useless rules. Following is an example of a rule discovered from mushroom database in UCI machine learning repository[6]. This database includes about 8,000 mushrooms as instances and each instance is classified into *edible* or *poisonous*.

$$R1 : odor \in \{almond, none\} \wedge cap_color \in \{brown, red\} \Rightarrow edible,$$

where

$$\text{support}(R1) = 23\%,$$

$$\text{accuracy}(R1) = 99\%.$$

Support and accuracy of this rule are 23% and 99% respectively and this rule looks like a quite good rule because of its large support and high accuracy. However, this rule is useless in usual because the following rule can be also discovered from the database.

$$R2 : \text{odor} \in \{\text{almond}, \text{none}\} \Rightarrow \text{edible},$$

$$\text{support}(R2) = 47\%, \text{accuracy}(R2) = 97\%.$$

This rule is a generalization of $R1$ and can explain class labels of all edible mushrooms covered by $R1$ with similar confidence. The specialization of $R2$ to $R1$ by the condition $\text{cap_color} \in \{\text{brown}, \text{red}\}$ rejects a half of positive instances covered by $R2$ and gains only 2% for accuracy. Unless this slight gain of accuracy is quite important for the purpose of rule discovery, $R1$ is a trivial conclusion of $R2$ and it is sufficient to discover only $R2$.

This example is not a rare case. If there exist a rule with large support and high accuracy, specializing it with irrelevant conditions with no relationship to classes causes many useless rules with large support and high accuracy. To reject such rules, it is not sufficient to constrain rules only with their support and accuracy, and different criteria are required to extract interesting rules.

Following is also an example of a rule from mushroom database.

$$R3 : \text{stalk_shape} = \text{enlarging} \wedge \text{habitat} = \text{woods} \Rightarrow \text{poisonous},$$

$$\text{support}(R3) = 9\%, \text{accuracy}(R3) = 93\%.$$

Its generalizations conclude *poisonous* with the following probabilities.

$$\text{Accuracy}(\text{stalk_shape} = \text{enlarging} \Rightarrow \text{poisonous}) = 54\%,$$

$$\text{Accuracy}(\text{habitat} = \text{woods} \Rightarrow \text{poisonous}) = 41\%.$$

Because about 48% of mushrooms in the database are poisonous, these probabilities mean that every condition in the body of $R3$ has only very weak relation to whether *poisonous* or *edible*, but the rule $R3$ shows the combination of the conditions leads poisonous with high accuracy. In other words, both of the conditions are required to conclude poisonous. Even if both of support and accuracy of $R3$ are lower than those of $R1$, this rule is usually more preferable than $R1$ because the accuracy of its generalizations are much lower than the accuracy of $R3$ and $R3$ gives us new information for poisonous mushrooms even if we already know the general rules.

The above example shows that interestingness of a rule strongly depends on more general and simpler rules. If a rule concludes the same class with similar accuracy with its generalizations, then the rule gives no new information and is useless. Oppositely, if accuracy of the rule is much higher than all of its generalizations and all conditions in the rule are required to imply the conclusion, then it is worth to show the rule to the users of discovery systems.

3.2.2 Interestingness of Primitive Rules

This subsection discusses interestingness of primitive rules before rules with grouping attribute values and proposes a criterion of primitive rules that evaluates whether every condition in a rule to be evaluated is really required to lead the conclusion of the rule. The next subsection will modify it to evaluate rules with grouping values.

Let's consider the rule in (3.1) with L conditions. This subsection assumes primitive rules and each D_i involves exactly one value for attribute x_i . For simplicity, let us use B to represent the body of the rule and B_i to represent a conjunction of $L-1$ conditions

in B excluding the i th condition, i.e.

$$B_i = x_1 \in D_1 \wedge \dots \wedge x_{i-1} \in D_{i-1} \wedge x_{i+1} \in D_{i+1} \wedge \dots \wedge x_L \in D_L. \quad (3.4)$$

As described above, I require interesting rules that all conditions in their bodies are required to achieve high accuracy. This leads the following constraint.

$$P(c|B) > P(c|B_i), \quad 1 \leq \forall i \leq L. \quad (3.5)$$

This constraint means every condition $x_i \in D_i$ in B contributes to improve accuracy of rule (3.1). The more the difference of accuracy between the rule (3.1) and its generalization($B_i \Rightarrow c$) is, the more the condition is important to conclude class c .

The constraint (3.5) is not sufficient to reject useless rules that are trivial conclusions of more general rules. When a condition $x_i \in D_i$ positively relates to class c , it is not interesting even if the specialization by the condition increases the probability of c and the constraint (3.5) is satisfied. Assuming independence of $x_i \in D_i$ and B_i , accuracy of the rule (3.1) can be calculated as follows.

$$\begin{aligned} P(c|B) &= P(c|x_i \in D_i \wedge B_i) \\ &= \frac{P(x_i \in D_i \wedge B_i|c)P(c)}{P(x_i \in D_i \wedge B_i)} \\ &= \frac{P(x_i \in D_i|c)P(B_i|c)P(c)}{P(x_i \in D_i)P(B_i)} \\ &= \frac{P(c|x_i \in D_i)P(c|B_i)}{P(c)}. \end{aligned}$$

If the real accuracy is comparable to this expectation, the rule to be evaluated represents only a regularity that can be easily expected from two more general rules, $x_i \in D_i \Rightarrow c$ and $B_i \Rightarrow c$. Then, it is natural to require the following constraint in addition to (3.5).

$$P(c|B) > \frac{P(c|x_i \in D_i)P(c|B_i)}{P(c)}, \quad 1 \leq \forall i \leq L. \quad (3.6)$$

Again, large difference between the left hand side and the right hand side means large contribution of the condition $x_i \in D_i$ to the rule.

Both of the constraints (3.5) and (3.6) give lower bounds of accuracy and we define interestingness of a primitive rule as the margin of its accuracy to satisfy these constraints as follows.

$$I(B \Rightarrow c) = \frac{P(c|B) - \max_{1 \leq i \leq L} (P(c|B_i), \min(1, \frac{P(c|x_i \in D_i)P(c|B_i)}{P(c)}))}{1 - P(c)}. \quad (3.7)$$

The denominator, $1 - P(c)$, is a normalization factor for class distributions and the interestingness becomes 1 for the rule that concludes a certain class with 100% accuracy even if each its generalization, $x_i \in D_i$ and B_i , is independent of the conclusion. The constraint on the minimum interestingness restricts accuracy of rules with dynamic, not fixed, lower bounds depending on accuracy of their generalizations.

For a rule with only one condition, its generalization is the rule with no condition and its interestingness becomes

$$I(x_1 \in D_1 \Rightarrow c) = \frac{P(c|x_1 \in D_1) - P(c)}{1 - P(c)}, \quad (3.8)$$

and the restriction on the interestingness coincides with the restriction of accuracy with a fixed threshold.

Let me show how this interestingness works for the examples $R1$ and $R3$ from mushroom database. Because 48% of mushrooms in the database are poisonous and the relationships between poisonous and the conditions in $R1$ are

$$P(\text{poisonous}|\text{odor} \in \{\text{almond}, \text{none}\}) = 97\%,$$

$$P(\text{poisonous}|\text{cap.color} \in \{\text{brown}, \text{red}\}) = 50\%.$$

Then the interestingness of $R1$ is

$$\begin{aligned} I(R1) &= \frac{0.99 - \max(0.97, 0.50, \frac{0.97 \times 0.50}{0.48})}{1 - 0.48} \\ &= -0.03. \end{aligned}$$

This means the interestingness I judges $R1$ is not interesting at all. In opposite, the interestingness of $R3$ becomes

$$I(R1) = \frac{0.93 - \max(0.54, 0.41, \frac{0.54 \times 0.41}{0.48})}{1 - 0.48} = 0.75,$$

and $R3$ is interesting with respect to I . These results for $R1$ and $R3$ coincide with the discussion in the previous subsection.

3.2.3 Interestingness of Rules with Grouping Attribute Values

The interestingness I defined in (3.7) works well for primitive rules as shown in the previous subsection but it sometimes gives high scores to inappropriate rules with grouping attribute values. Following is an example from mushroom database again.

$$R4 : cap_color \in \{brown, red\} \wedge stalk_root \in \{bulbous, rooted\} \Rightarrow edible,$$

where

$$Support(R4) = 16\%, Accuracy(R4) = 100\%.$$

This rule looks like interesting because accuracy of its generalizations are

$$Accuracy(cap_color \in \{brown, red\} \Rightarrow edible) = 50\%,$$

$$Accuracy(stalk_root \in \{bulbous, rooted\} \Rightarrow edible) = 53\%,$$

and every condition in $R4$ has very weak relation to the conclusion of the rule but only the edible instances satisfy the conjunction of the two conditions. The interestingness defined in (3.7) gives high score, 0.98 to $R4$ and judges $R4$ is interesting. However, by exploring relationship between *edible* and each attribute value allowed in the rule, we

can find the following rule.

$$stalk_root = rooted \Rightarrow edible,$$

$$Accuracy = 100\%.$$

This rule means that the condition on *cap_color* in $R4$ is redundant for mushrooms with *stalk_root = rooted* and $R4$ is not interesting at all for classifying the mushrooms. Instead, the following rule that prohibits *stalk_root = rooted* is more preferable than $R4$ because each condition in its body is required to classify *all* mushrooms covered by its body with high accuracy.

$$R5 : cap_color \in \{brown, red\} \wedge stalk_root = bulbous \Rightarrow edible.$$

The problem of the interestingness defined in (3.7) is that it evaluates a rule based on only each condition in its body and it does not concern each attribute value allowed in the conditions at all. To resolve this problem, we introduce the interestingness of an attribute value to evaluate how the condition on the attribute contributes to classify the instances with the value accurately. Formally, the interestingness of a value v for an attribute x_i allowed in the rule $x_1 \in D_1 \wedge \dots \wedge x_L \in D_L \Rightarrow c$ is

$$I_{value}(x_i = v, x_1 \in D_1 \wedge \dots \wedge x_L \in D_L \Rightarrow c) = \max_{v_k \in D_k (k \neq i)} I(x_1 = v_1 \wedge \dots \wedge x_i = v \wedge \dots \wedge x_L = v_L \Rightarrow c). \quad (3.9)$$

For example, the interestingness of the attribute values in $R4$ is calculated from the interestingness of the following four primitive rules.

$$cap_color = brown \wedge stalk_root = bulbous \Rightarrow edible (I = 0.93),$$

$$cap_color = brown \wedge stalk_root = rooted \Rightarrow edible (I = 0.00),$$

$$cap_color = red \wedge stalk_root = bulbous \Rightarrow edible (I = 1.02),$$

$$cap_color = red \wedge stalk_root = rooted \Rightarrow edible (I = 0.00).$$

Table 3.2: The number of discovered rules and the running time of DIG. The running time is the total cpu time to discover rules with two attributes ($L = 2$). This table shows the dependency on the lower bound of the interestingness. The lower bounds of support and accuracy are 2%(10% for promoters database only) and 90% respectively.

	# of combs	LB_{int}	# of rules	cpu time(sec)
dna	5,310	$-\infty$	145	155
		0.0	85	22
		0.1	17	20
		0.3	0	19
		0.5	0	20
		0.7	0	20
krkp	1,260	$-\infty$	82	7
		0.0	73	7
		0.1	10	7
		0.3	3	7
		0.5	3	7
		0.7	3	7
german	156	$-\infty$	31	95
		0.0	7	0.3
		0.1	3	0.3
		0.3	2	0.3
		0.5	0	0.3
		0.7	0	0.3
mushroom	460	$-\infty$	17,769	540
		0.0	4,899	11
		0.1	215	7
		0.3	105	7
		0.5	52	7
		0.7	15	7
promoters	3,192	$-\infty$	135	68
		0.0	61	1
		0.1	14	1
		0.3	2	1
		0.5	1	1
		0.7	0	1

Table 3.3: The number of discovered rules and the running time of DIG. The running time is the total cpu time to discover rules with three attributes ($L = 3$). This table shows the dependency on the lower bound of the interestingness. The lower bounds of support and accuracy are 2%(10% for promoters database only) and 90%.

	# of combs	LB_{int}	# of rules	cpu time(sec)
dna	102,660	$-\infty$	4,095,272	972,012
		0.0	508,444	11,114
		0.1	1,052	1,066
		0.3	1	938
		0.5	0	936
		0.7	0	936
krkp	14,280	$-\infty$	4,657	201
		0.0	4,334	174
		0.1	25	173
		0.3	2	173
		0.5	2	173
		0.7	0	173
german	572	$-\infty$	151,291	74,220
		0.0	25,976	11,124
		0.1	415	278
		0.3	31	51
		0.5	0	7
		0.7	0	3
mushroom	3,080	$-\infty$	-	> 10 days
		0.0	9,713,290	236,238
		0.1	175	99
		0.3	26	98
		0.5	10	98
		0.7	1	99
promoters	58,520	$-\infty$	1,462,162	130,030
		0.0	815,513	25,708
		0.1	5,672	2,586
		0.3	489	767
		0.5	9	190
		0.7	0	90

