



Advanced Approach in Sensitive Rule Hiding

Dr. K. Duraiswamy

K.S. R. College of Technology

Tiruchengode- 637 209, Tamil Nadu, India

E-mail: kduraiswamy@yahoo.co.in

Dr. D. Manjula

Department of Computer Science and Engineering

Anna University

Chennai, Tamil Nadu, India

E-mail: manju@annauniv.edu

N. Maheswari (Corresponding Author)

P.G. Department of Computer Science

Kongu Arts and Science College

Erode-638 107, Tamil Nadu, India

E-mail: mahii_14@yahoo.com

Abstract

Privacy preserving data mining is a novel research direction in data mining and statistical databases, which has recently been proposed in response to the concerns of preserving personal or sensible information derived from data mining algorithms. There have been two types of privacy proposed concerning data mining. The first type of privacy, called output privacy, is that the data is altered so that the mining result will preserve certain privacy. The second type of privacy, called input privacy, is that the data is manipulated so that the mining result is not affected or minimally affected. For output privacy in hiding association rules, current approaches require hidden rules or patterns to be given in advance. However, to specify hidden rules, entire data mining process needs to be executed. For some applications, only certain sensitive rules that contain sensitive items are required to hide. In this work, an algorithm ISSRH (Increase Support Sensitive Rule Hiding) is proposed, to hide the sensitive rules that contain sensitive items, so that sensitive rules containing specified sensitive items on the right hand side of the rule cannot be inferred through association rule mining. Example illustrating the proposed approach is given. The characteristics of the algorithm are discussed.

Keywords: Data Mining, Privacy Preserving, Association Rules, Sensitive Rules, Clustering, Minimum Support, Minimum confidence

1. Introduction

The concept of Privacy-Preserving has recently been proposed in response to the concerns of preserving personal or sensible information derived from data mining algorithms. Successful applications of data mining have been demonstrated in marketing, business, medical analysis, product control, engineering design, bioinformatics and scientific exploration, among others. The current status in data mining research reveals that one of the current technical challenges is the development of techniques that incorporate security and privacy issues. The main reason is that the increasingly popular use of data mining tools has triggered great opportunities in several application areas, which also requires special attention regarding privacy protection. There have been two types of privacy concerning data mining. The first type of privacy, called output privacy, is that the data is minimally altered so that the mining result will preserve certain privacy (Evfimievski, 2002, Oliveira, Zaiane, 2003 a, Oliveira, Zaiane, 2003 b). The second type of privacy, input privacy, is that the data is manipulated so that the mining result is not affected or minimally affected (Dasseni, Verykios, Elmagarmid, Bertino, 2001).

For example, through data mining, one is able to infer sensitive information, including personal information, or even patterns from non-sensitive information or unclassified data. As a motivating example of privacy issue in data mining discussed in (Yi-Hung Wu, Chia-Ming Chiang, and Arbee L.P. Chen, 2007). Consider a supermarket and two breads suppliers A and B. If the transaction database of the supermarket is released, A (or B) can mine the association rules related to his/her breads and apply the rules to the sales promotion and the goods supply. As a result, a supplier is willing to exchange a lower price of goods for the database with the supermarket. From this aspect, it is good for the supermarket to release the database. However, the conclusion can be opposite if a supplier uses the mining methods in a different way. For instance, if A finds the association rules related to B's breads, saying that most customers who buy cheese also buy B's breads, he/she can run a coupon that gives a 10 percent discount when buying A's breads together with cheese. Gradually, the amount of sales on B's breads is down and B cannot give a low price to the supermarket as before. Finally, A monopolizes the bread market and is unwilling to give a low price to the supermarket as before. From this aspect, releasing the database is bad for the supermarket. Therefore, for the supermarket, an effective way to release the database with sensitive rules hidden is required. This leads to the research of sensitive rule hiding.

In this work, the sensitive rules are given and the algorithm ISSRH is proposed to modify data in database so that sensitive rules containing specified sensitive items on the right hand side of rule cannot be inferred through association rule mining. The proposed algorithm is based on modifying or perturbing the database transactions so that the confidence of the association rules can be reduced.

The rest of the paper is organized as follows. Section 2 gives the view of the previous works. Section 3 presents the statement of the problem. Section 4 presents the framework of the approach. Section 5 presents the proposed algorithm for sensitive rule hiding. Section 6 shows the example of the proposed algorithm. Section 7 analyzes the characteristics of the algorithm. Concluding remarks and future work are described in Section 8.

2. Related Work

In output privacy, given specific rules or patterns to be hidden, many data altering techniques for hiding association, classification and clustering rules have been proposed. For association rules hiding, two basic approaches have been proposed. The first approach (Saygin , Verykios , Clifton, 2001, Verykios, Elmagarmid , Bertino , Saygin , Dasseni ,2004) hides one rule at a time. It first selects transactions that contain the items in a give rule. It then tries to modify transaction by transaction until the confidence or support of the rule fall below minimum confidence or minimum support. Either removing items from the transaction or inserting new items to the transactions does the modification of transaction. The second approach (Oliveira, Zaiane, 2002 a, Oliveira, Zaiane, 2002 b, Oliveira, Zaiane, 2003 a, Oliveira, Zaiane, 2003 b) deals with groups of restricted patterns or association rules at a time. It first selects the transactions that contain the intersecting patterns of a group of restricted patterns. Depending on the disclosure threshold given by users, it sanitizes a percentage of the selected transactions in order to hide the restricted patterns. However, both the above approaches require hidden rules or patterns been given in advance.

The work presented here differs from the related work in some aspects are as follows: First, database indexing is performed. Second, correlations among the sensitive rules are considered. Third, avoids the modification in transactions unnecessarily, if the confidence of the sensitive rule gets reduced. Fourth, alter the transactions in the cluster and finally changes can be updated in the database, which reduces the time period of database updating.

3. Problem Statement

The problem of mining association rules was introduced in (Agrawal, Imielinski, Swami, 1993). Let $I = \{i_1, i_2, \dots, i_m\}$ be a set of literals, called items. Given a set of transactions D , where each transaction T in D is a set of items such that $T \subseteq I$, an association rule is an expression $X \Rightarrow Y$ where $X \subseteq I$, $Y \subseteq I$, and $X \cap Y = \Phi$. The X and Y are called respectively the body (left hand side) and head (right hand side) of the rule. An example of such a rule is that 90% of customers buy hamburgers also buy Coke. The 90% here is called the confidence of the rule, which means that 90% of transaction that contains X (hamburgers) also contains Y (Coke). The confidence is calculated as $|X \cup Y| / |X|$, where $|X|$ is the number of transactions containing X and $|X \cup Y|$ is the number of transactions containing both X and Y . The notation $\cup$ here is not the set union operator. The support of the rule is the percentage of transactions that contain both X and Y , which is calculated as $|X \cup Y| / N$, where N is the number of transactions in D . In other words, the confidence of a rule measures the degree of the correlation between item sets, while the support of a rule measures the significance of the item sets. A typical association rule-mining algorithm first finds all the sets of items that appear frequently enough to be considered significant and then it derives from them the association rules that are strong enough to be considered interesting. The problem of mining association rules is to find all rules that are greater than the user-specified minimum support and minimum confidence.

The objective of data mining is to extract hidden or potentially unknown but interesting rules or patterns from databases. However, the objective of privacy preserving data mining is to hide certain sensitive information so that they cannot be discovered through data mining techniques (Agrawal, Imielinski, Swami, 1993, Evfimievski, Gehrke , Srikant, 2003).

In this work, an algorithm ISSRH (Increase Support Sensitive Rule Hiding) is proposed, to hide the sensitive rules that contain sensitive items, so that sensitive rules containing specified sensitive items on the right hand side of rule cannot be inferred through association rule mining. More specifically, given a transaction database D , a minimum support, a minimum confidence and a set of sensitive items Y , the objective is to minimally modify the database D such that no sensitive rules containing sensitive items Y on the right hand side of the rule will be discovered.

4. Framework of the approach

Figure 1 shows the framework of the approach that consists of six processes. Initially, indexing is performed in the database. Then association rules are mined from the database. Sensitive items are identified to find the sensitive rules. Then the sensitive rules are generated. Clustering is performed on the sensitive rules to group the similar items. The rule hiding process is performed and the transactions are updated in the transaction table and finally it is updated in the original database. The main challenge of rule hiding is how to select the items and transactions to modify. The proposed framework hides the sensitive rules.

5. Proposed Algorithm

In order to hide an sensitive rule, $X \Rightarrow Y$, it can be either decrease its supports, $(|X|/N$ or $|X \cup Y|/N)$, to be smaller than pre-specified minimum support or its confidence $(|X \cup Y|/|X|)$ to be smaller than pre-specified minimum confidence. In the transactions that do not contain both X and Y , to increase the support of X only, the left hand side of the rule, it would reduce the confidence of the rule. In order to hide sensitive rules, when considering hiding sensitive rules with 2 items, $x \Rightarrow z$, where z is a sensitive item and x is a single large one item. In theory, association rules may have more specific rules that contain more items, e.g., $xY \Rightarrow z$, where Y is a large item set. However, for such rule to exist, its confidence must be greater than the confidence of $x \Rightarrow z$, i.e., $\text{conf}(xY \Rightarrow z) > \text{conf}(x \Rightarrow z)$ or $|xYz| > \text{conf}(x \Rightarrow z) * |xY|$. For higher confidence rules, such as $\text{conf}(x \Rightarrow z) = 1$, there will be no more specific rules. In addition, once the more general rule is hidden, the more specific rule might be hidden as well.

The algorithm tries to decrease the support of the right hand side of the rule.

Algorithm ISSRH

Input:

1. source database D
2. min support
3. min conf
4. sensitive items Y

Output:

a transformed database D' , where rules containing Y on RHS will be hidden.

Step 1:

Indexing the transactional database.

Step 2:

Generate the association rules.

Step 3:

Selecting the Sensitive rules with single antecedent and consequent with the sensitive item in the consequent. ($x \rightarrow y$)

Step 4:

Constraint based clustering - Clustering the rules with the right hand side has the common item and indexing the rules.

Step 5:

Check all the rules in the cluster.

for cluster $i = 1$ to n

```
{
do
{
1    select the rule with high confidence
2    find  $K = \text{mincount}$ .
```

```

3      |Tr| = no of sensitive rules that are correlated
Such that  $k < |Tr|$ 
4      Tk = find first k transactions that do not contains the items in the correlated rules
5      sort Tk in ascending order by the number of items.
6      For i= 1 to k
      {
7          Choose the first transaction t from Tk.
8          Modify t to support x.  $\text{count}(x) = \text{count}(x) + 1$ .
9          Find the confidence of every rule
10         If confidence of all the rules < min conf
            {
11             Update the transactions in the clusters
12             Exit
13             Else
14             Update the transactions in the clusters
            }
        }
15     unchecked the rules < minconf
16     }while(rules in cluster[i]==checked)
17     }

```

Step 4: update database D as transformed database D'.

The algorithm tries to generate the association rule using Agrawal, Imielinski, Swami, 1993). Then it selects the sensitive rules with the sensitive items in the right hand side. Cluster the rules with the common item in the right hand side of the rule (Han, Kamber, 2001) and index the rule using (Oliveira, Zaiane, 2003 a).

The mincount= $[\text{count}(xUy/\text{min conf}) - \text{count}(x) + 1$.

(Yi-Hung Wu, Chia-Ming Chiang, and Arbee L.P. Chen, 2007)

Therefore it would take maximum of k no of executions to hide the rule. The rules in every cluster will be hidden.

6. Example

This section shows the example to demonstrate the proposed algorithm to hide the sensitive rules.

The items in database can be represented as a bit vector 1 and 0. (Saygin Y, Verykios V, Clifton C, 2001)

<u>TID</u>	<u>Items</u>	<u>Items</u>
T1	ABC	111
T2	ABC	111
T3	ABC	111
T4	AB	110
T5	A	100
T6	AC	101

Frequent item sets are generated with minimum support 0.33. Association rules with minimum confidence 0.70 are generated. The rules and the corresponding confidence values are as follows:

$C \rightarrow B 0.75$, $B \rightarrow C 0.75$, $C \rightarrow A 1.0$, $B \rightarrow A 1.0$, $BC \rightarrow A 1.0$, $AC \rightarrow B 0.75$, $AB \rightarrow C 0.75$, $C \rightarrow AB 0.75$, and $B \rightarrow AC 0.75$.

The item C is considered as sensitive item. The sensitive rule with single antecedent and consequent is $B \rightarrow C$. The rule is clustered. In the fifth transaction the item B will be added and placed as 1.

Then the database will be updated as

<u>TID</u>	<u>Items</u>	<u>Items</u>
T1	ABC	011

T2	ABC	011
T3	ABC	111
T4	AB	110
T5	AB	110
T6	AC	101

After updating the database the rules $B \rightarrow C$ will have the confidence as 0.6, which is less than minimum confidence and hidden. While hiding the rule the rules $AB \rightarrow C$, $B \rightarrow AC$ are also hidden and the rule $A \rightarrow B$ is generated as side effect.

7. Analysis

This section analyses some of the characteristics of the proposed algorithm. The first characteristic is the database indexing. The indexing helps in reducing the number of scanning of the database. The second characteristic is the time effect. The time taken to scan the database to search the sensitive rules in the database is reduced because of clustering the sensitive rules. The third characteristic is the database effect. The minimum numbers of transactions are modified because of correlation among the sensitive rules. The fourth characteristic is the efficiency of the algorithm. The database is updated after all the rules are hidden that saves the updating time. The fifth characteristic is the transaction effect. The alteration in the transactions are stopped when the confidence of the sensitive rules are reduced than the minimum confidence.

8. Conclusion

In this work, the database privacy problems caused by data mining technology are discussed and the algorithm for hiding sensitive rules is presented. The proposed algorithm here can automatically hide sensitive rule sets. The previous works does not consider the characteristics that are discussed here. Example illustrating the proposed algorithm is given and the characteristics of the algorithm are analyzed. Further the efficiency of the algorithm will be analyzed and improved by reducing the side effects.

References

- Agrawal R, Imielinski T, Swami A (1993) Mining association rules between sets of items in large databases. *In: Proceedings of ACM SIGMOD International Conference on Management of Data*, Washington DC
- Dasseni E, Verykios V, Elmagarmid A, Bertino E (2001) Hiding association rules by using confidence and support. *In: Proceedings of 4th Information Hiding Workshop*, Pittsburgh, PA, pp 369–383
- Evfimievski A (2002) Randomization in privacy preserving data mining. *SIGKDD Explorations* 4(2), Issue 2:43–48
- Evfimievski A, Gehrke J, Srikant R (2003) Limiting Privacy Breaches in Privacy Preserving Data Mining. *PODS 2003*, San Diego, CA
- Evfimievski A, Srikant R, Agrawal R, Gehrke J (2002) Privacy preserving mining of association rules. *In: Proceedings of the 8th ACM SIGKDD Int'l Conference on Knowledge Discovery and Data Mining*, Edmonton, Canada
- Han, J., Kamber, M, (2001) *Data Mining : Concepts and Techniques.*, Morgan Kaufmann Publishers.
- Oliveira S, Zaiane O (2002 a) Privacy preserving frequent item set mining. *In: Proceedings of IEEE International Conference on Data Mining*, November 2002, pp 43–54.
- Oliveira S, Zaiane O (2002 b) A framework for enforcing privacy in mining frequent patterns. *Technical report, TR02-13*, Computer Science Department, University of Alberta, Canada
- Oliveira S, Zaiane O (2003 a) Algorithms for balancing privacy and knowledge discovery in association rule mining. *In: Proceedings of 7th International Database Engineering and Applications Symposium (IDEAS03)*, Hong Kong
- Oliveira S, Zaiane O (2003 b) Protecting sensitive knowledge by data sanitization. *In: Proceedings of IEEE International Conference on Data Mining*
- Saygin Y, Verykios V, Clifton C (2001) Using unknowns to prevent discovery of association rules. *SIGMOD Record* 30(4):45–54
- Verykios V, Elmagarmid A, Bertino E, Saygin Y, Dasseni E (2004) Association Rules Hiding. *IEEE Trans Knowledge and Data Eng* 16(4):434–447
- Yi-Hung Wu, Chia-Ming Chiang, and Arbee L.P. Chen, (2007), Hiding Sensitive Association Rules with Limited Side Effects, *IEEE Trans Knowledge and Data Eng*, Vol. 19, No. 1.

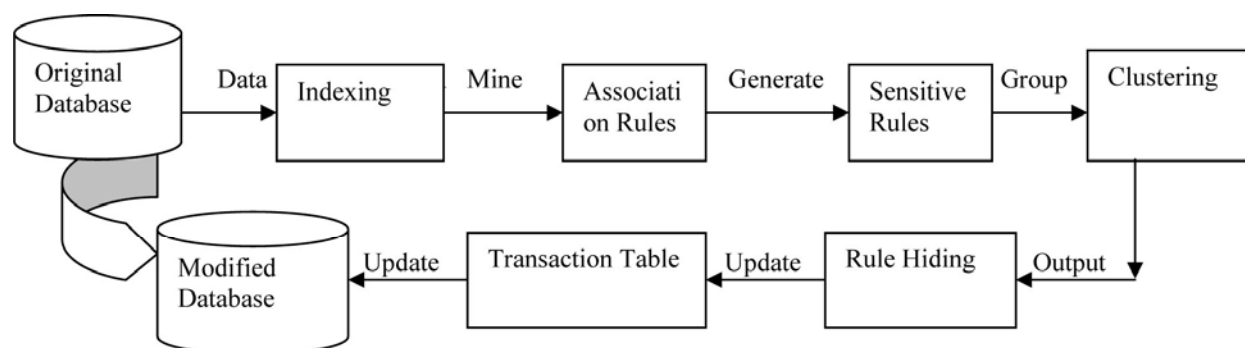


Figure 1. Framework of the approach

The framework describes the proposed approach of rule hiding.



Boundedness of Commutators on Generalized Morrey Spaces

Yuze Cai

Shazhou Polytechnical Institute Of Technology
Jiangsu 215600 China

Hao Zhang

Wuxi Institute of Technology
Jiangsu 214121, China

Abstract

In this paper, we establish the boundedness of strongly singular integrals operators T and commutators T_b on generalized Morrey spaces, where T_b are generated by $BMO(R^n)$ functions b and the strongly singular integrals operators T .

Keywords: Strongly singular integrals, Generalized Morrey spaces, Commutators, BMO functions

Let $\nu \in C_c^\infty(R^n)$, $\text{supp } \nu \subset \{x \in R^n : |x| \leq 2\}$. We define strongly singular integrals kernels $K(x) = \frac{e^{i|x|^{-s'}}}{|x|^n} \nu(x)$,

where $0 < s < 1$, $s' = \frac{s}{1-s}$, and the corresponding strongly singular integrals

$$Tf(x) = p.v. \int_{R^n} K(x-y)f(y)dy.$$

Let $b \in BMO(R^n)$. We define the commutators T_b generated by functions b and operators T as the following:

$$T_b f(x) = p.v. \int_{R^n} (b(x) - b(y)) K(x-y)f(y)dy.$$

Wainger S.etc (Wainger, S. 1965.) had studied the boundedness of the operators T on $L^q(R^n)$. Chanillo (Chanillo, S. 1984.) developed weighted $L^q(R^n)$ theory by virtue of a basal lemma which will be mentioned later. Garcia-Cuerva J. etc (Garcia-Cuerva, J., Harboure, E., Segovia, C. and Torrea, J. L.) obtained the boundedness of higher-order commutators on weighted $L^q(R^n)$. Morrey (Morrey, C. B. 1938.) proposed the classical Morrey spaces when he studied the properties of local solutions of second order elliptic equations. In Ref. (Mizuhara T., 1991.) Mizuhara introduced the following generalized Morrey spaces.

Let Φ be positive increasing functions on $(0, \infty)$ satisfying $\Phi(2r) \leq D\Phi(r)$ for any $r > 0$, where $D \geq 1$ is a constant independent of r .

Definition 1 (Mizuhara T. 1991.) For $1 \leq p < \infty$, we define the generalized Morrey spaces as the following

$$L^{p,\Phi}(R^n) = \left\{ f \in L_{loc}^p : \|f\|_{L^{p,\Phi}} < \infty \right\},$$

Where $\|f\|_{L^{p,\Phi}}^p = \sup_{t \in R^n, r > 0} \frac{1}{\Phi(r)} \int_{B(t,r)} |f(x)|^p dx$, $B(t,r)$ is a sphere r in radius whose center is at point t .

We know that when $\Phi(r) = r^\lambda$, $(0 < \lambda < n)$, $L^{p,\Phi}$ are the classical Morrey spaces.

In this paper we will show the boundedness of strongly singular integrals operators T on generalized Morrey spaces by weighted inequations. Furthermore, we will show the boundedness of the commutators generalized by BMO functions b of operators T by virtue of sharp estimate.

1. Boundedness of strongly singular integrals operators T on generalized Morry spaces

We fix the following notations in Lemma 1 and Theorem 1:

For $\forall x_0 \in R^n, r > 0$, $f(y) = f\chi_{B(x_0, 2r)}(y) + \sum_{k=1}^{\infty} f\chi_{B(x_0, 2^{k+1}r) \setminus B(x_0, 2^k r)}(y) \equiv \sum_{k=0}^{\infty} f_k(y)$.

Lemma 1 For $1 \leq D(\Phi) < 2^n, 1 < p < \infty$, $f \in L^{p, \Phi}$, if T are strongly singular integrals then for $\forall k, (k > 0)$ we have:

$$\int_{B(x_0, r)} |T(f_k)(x)|^p dx \leq C \left(\frac{D}{2^{\theta n}} \right)^k \|f\|_{L^{p, \Phi}}^p \Phi(r),$$

Where $\theta \in \left(\frac{\ln D}{\ln 2^n}, 1 \right)$.

Proof: By the weighted property of A_p , for any $\alpha \in (0, 1), \forall k > 0$, $(M\chi_{B(x_0, r)})^\alpha \in A_1 \subset A_p, (1 < p \leq \infty)$, we have

$$\begin{aligned} \int_{B(x_0, r)} |T(f_k)(x)|^p dx &\leq C \int_{B(x_0, r)} \left| \int_{B(x_0, 2^{k+1}r)} \frac{|f_k(y)|}{|x-y|^n} dy \right|^p dx \\ &\leq C \int_{B(x_0, r)} |M(|f_k|)(x)|^p dx \\ &\leq C \int_{R^n} |M(|f_k|)(x)|^p (M\chi_{B(x_0, r)})^\theta dx \\ &\leq C \int_{R^n} |f_k(x)|^p (M\chi_{B(x_0, r)})^\theta dx \\ &\leq C \frac{1}{(2^{k-1})^{n\theta}} \int_{B(x_0, 2^{k+1}r)} |f_k(x)|^p dx \\ &\leq C \frac{1}{(2^{k-1})^{n\theta}} \|f\|_{L^{p, \Phi}}^p \Phi(2^{k+1}r) \\ &\leq C \left(\frac{D}{2^{\theta n}} \right)^k \|f\|_{L^{p, \Phi}}^p \Phi(r). \end{aligned}$$

Theorem 1 For $1 \leq D(\Phi) < 2^n, 1 < p < \infty$, and T are strongly singular integrals then T are bounded on $L^{p, \Phi}(R^n)$.

Proof: $\forall x_0 \in R^n, r > 0$, by the boundedness of T on $L^p (p > 1)$ (Wainger, S. 1965.), we get

$$\int_{B(x_0, r)} |T(f_0)(x)|^p dx \leq C \int_{B(x_0, 2r)} |f(x)|^p dx \leq C \|f\|_{L^{p, \Phi}}^p \Phi(r).$$

And by Lemma 1 we have

$$\begin{aligned} \left[\int_{B(x_0, r)} |T(f)(x)|^p dx \right]^{\frac{1}{p}} &\leq \sum_{k=0}^{\infty} \left[\int_{B(x_0, r)} |T(f_k)(x)|^p dx \right]^{\frac{1}{p}} \\ &\leq C \left[1 + \sum_{k=1}^{\infty} \left(\frac{D}{2^{n\theta}} \right)^{\frac{k}{p}} \right] \|f\|_{L^{p, \Phi}}^{\frac{1}{p}} \Phi^{\frac{1}{p}}(r) \\ &\leq C \|f\|_{L^{p, \Phi}}^{\frac{1}{p}} \Phi^{\frac{1}{p}}(r). \end{aligned}$$

We used $\theta \in \left(\frac{\ln D}{\ln 2^n}, 1 \right)$ in the last inequation.

2. Boundedness of the commutators on generalized Morrey spaces

Lemma 2 (Mizuhara T. 1991.) For $1 \leq D(\Phi) < 2^n, 1 \leq q < p < \infty$, then $M_q f(x) = \left(M(|f|^q)(x) \right)^{\frac{1}{q}}$ are bounded on $L^{p, \Phi}(R^n)$, where M are H-L maximum operators.

Lemma 3 (Ding, Y. 1997.) For $1 \leq D(\Phi) < 2^n, 1 < p < \infty$, then for $\forall f \in L^{p,\Phi}(R^n)$ we have

$$\|Mf\|_{L^{p,\Phi}} \leq C \|f^\#\|_{L^{p,\Phi}},$$

where $f^\#(x) = \sup_{x \in Q} \frac{1}{|Q|} \int_Q |f(y) - f_Q| dy$.

Lemma 4 If T are strongly singular integrals, $1 < r, s < \infty, b \in BMO(R^n)$, then for $a.e. x \in R^n, \forall f \in L^{p,\Phi}(R^n)$, $(1 < p < \infty)$ there exists a constant C independent of b, f so that the following inequations hold

$$(T_b f)^\#(x) \leq C \|b\|_* \left\{ \left(M|Tf|^{r'} \right)^{\frac{1}{r}}(x) + \left(M|f|^{r'} \right)^{\frac{1}{r}}(x) \right\}.$$

Proof: We only need to show for any $x_0 \in R^n$ the following inequations hold

$$(T_b f)^\#(x_0) \leq C \|b\|_* \left\{ \left(M|Tf|^{r'} \right)^{\frac{1}{r}}(x_0) + \left(M|f|^{r'} \right)^{\frac{1}{r}}(x_0) \right\}.$$

We choose a cube Q with side length δ , $s.t. x_0 \in Q$. Let $4\delta_0 = \delta_0^{\frac{1}{1+s'}}$ and we know δ_0 is a constant less than 1.

Case 1: $\delta < \delta_0$.

Suppose cube $\tilde{Q}$ with side length $\delta^{\frac{1}{1+s'}}$ has the same center as Q . We have

$$f = f\chi_{4Q} + f\chi_{\tilde{Q} \setminus 4Q} + f\chi_{R^n \setminus \tilde{Q}} \equiv f_1 + f_2 + f_3.$$

$$T_b f(x) = \int_{R^n} (b_Q - b(y))k(x-y)f(y)dy - (b_Q - b(x)) \int_{R^n} k(x-y)f(y)dy.$$

Let $C_Q = \frac{1}{|Q|} \int_Q \int_{R^n} (b_Q - b(y))k(z-y)f_3(y)dydz$. It is obvious that C_Q is a constant. So

$$\begin{aligned} & \frac{1}{|Q|} \int_Q |T_b(f)(x) - C_Q| dx \\ & \leq \frac{1}{|Q|} \int_Q |(b_Q - b(x))Tf(x)| dx + \frac{1}{|Q|} \int_Q \left| \int_{R^n} (b_Q - b(y))k(x-y)f_1(y)dy \right| dx \\ & + \frac{1}{|Q|} \int_Q \left| \int_{R^n} (b_Q - b(y))k(x-y)f_2(y)dy \right| dx \\ & + \frac{1}{|Q|} \int_Q \left(\frac{1}{|Q|} \int_Q \left| \int_{R^n} (b_Q - b(y))(k(x-y) - k(z-y))f_3(y)dy \right| dz \right) dx \\ & = A_1 + A_2 + A_3 + A_4. \end{aligned}$$

As for A_1 , we have

$$\begin{aligned} A_1 & \leq C \left(\frac{1}{|Q|} \int_Q |b_Q - b(x)|^{r'} dx \right)^{\frac{1}{r}} \left(\frac{1}{|Q|} \int_Q |Tf(x)|^r dx \right)^{\frac{1}{r}} \\ & \leq \|b\|_* \left(M(|Tf|)^r \right)^{\frac{1}{r}}(x_0). \end{aligned}$$

As for A_2 , for $1 < q < t, u > 1$ so that $qu = t$, by Hölder inequation we have

$$A_2 \leq C \left\{ \frac{1}{|Q|} \int_Q |T((b_Q - b)f_1)(x)|^q dx \right\}^{\frac{1}{q}}$$

$$\begin{aligned}
&\leq C \left\{ \frac{1}{|Q|} \int_Q |(b_Q - b)f(x)|^q dx \right\}^{\frac{1}{q}} \\
&\leq C \left\{ \frac{1}{|Q|} \int_Q |(b_Q - b(x))^{qu'}|^{qu'} dx \right\}^{\frac{1}{qu'}} \left\{ \frac{1}{|Q|} \int_Q |f(x)|^{qu} dx \right\}^{\frac{1}{qu}} \\
&\leq C \|b\|_* \left(M|f|^{t'} \right)^{\frac{1}{t'}}(x_0).
\end{aligned}$$

As for A_4 , by noting $x \in Q, z \in Q, y \notin \tilde{Q}$, and mean value theorem, we have

$$\begin{aligned}
|k(x-y) - k(z-y)| &\leq \frac{C\delta}{|y-z|^{n+s'+1}}. \\
A_4 &\leq C \frac{1}{|Q|} \int_Q \left(\delta \int_{|y-x_0| > \frac{1}{2}\delta^{1+s'}} |b_Q - b(y)| |y-z|^{-(n+s'+1)} |f(y)| dy \right) dz \\
&\leq C \frac{1}{|Q|} \int_Q \left(\sum_{k=0}^{\infty} 2^{-k} (2^k \delta)^{\frac{-n}{1+s'}} \int_{|y-z| < 2(2^k \delta)^{\frac{1}{1+s'}}} |b_Q - b(y)| |f(y)| dy \right) dz \\
&\leq C \frac{1}{|Q|} \int_Q \sum_{k=0}^{\infty} 2^{-k} \left\{ (2^k \delta)^{\frac{-n}{1+s'}} \int_{|y-z| < 2(2^k \delta)^{\frac{1}{1+s'}}} |b_Q - b(y)|^{t'} dy \right\}^{\frac{1}{t'}} \left\{ (2^k \delta)^{\frac{-n}{1+s'}} \int_{|y-z| < 2(2^k \delta)^{\frac{1}{1+s'}}} |f(y)|^{t'} dy \right\}^{\frac{1}{t'}} dz \\
&\leq C \|b\|_* \left(M|f|^{t'} \right)^{\frac{1}{t'}}(x_0).
\end{aligned}$$

As for A_3 , we use the similar ways in Chanillo (Chanillo, S. 1984.). Let $l > 1, m > 1$ satisfying $2 + s' < l', lm = t$, then

$$\begin{aligned}
&\int_{\mathbb{R}^n} (b_Q - b(y)) k(x-y) f_2(y) dy \\
&= \int_{\mathbb{R}^n} \frac{e^{i|x-y|^{-s'}} v(x-y)}{|x-y|^{\frac{n(2+s')}{l'}}} \left(\frac{1}{|x-y|^{n(1-\frac{2+s'}{l'})}} - \frac{1}{|x_0-y|^{n(1-\frac{2+s'}{l'})}} \right) (b_Q - b(y)) f_2(y) dy \\
&+ \int_{\mathbb{R}^n} \frac{e^{i|x-y|^{-s'}} v(x-y)}{|x-y|^{\frac{n(2+s')}{l'}}} \frac{(b_Q - b(y)) f_2(y)}{|x_0-y|^{n(1-\frac{2+s'}{l'})}} dy \\
&= B_1 + B_2.
\end{aligned}$$

As for B_1 , by the mean value theorem, we get

$$\begin{aligned}
B_1 &\leq C \sum_{k=0}^{\infty} 2^{-k} (2^k \delta)^{-n} \int_{|y-x_0| < 2^{k+1} \delta} |b_Q - b(y)| |f(y)| dy \\
&\leq C \|b\|_* \left(M|f|^{t'} \right)^{\frac{1}{t'}}(x_0).
\end{aligned}$$

As for the estimate of B_2 , we need the following lemma.

Lemma 5 (Chanillo, S.) For $0 < s < 1, s' = \frac{s}{1-s}, t > 1, \frac{s'+2}{t} < 1$ $\tilde{K}_{s',t} \equiv \frac{e^{i|x|^{-s'}}}{|x|^{\frac{n(s'+2)}{t}}}$, we have $\forall f \in L^{t'}(\mathbb{R}^n)$,

$$\left\| \tilde{K}_{s',t} * f \right\|_{L^{t'}(\mathbb{R}^n)} \leq C \|f\|_{L^{t'}(\mathbb{R}^n)}.$$