

GIẢI THUẬT HIỆU NĂNG CAO KHAI THÁC TẬP SINH CỦA TẬP ĐÓNG PHỔ BIẾN

HIGH-PERFORMANCE ALGORITHM IN MINING GENERATORS OF CLOSED FREQUENT ITEMSETS

Phan Thành Huân

Đại học Quốc gia Tp. Hồ Chí Minh; huanphan@hcmussh.edu.vn

Tóm tắt - Trong khai thác dữ liệu, khai thác luật kết hợp là một trong những kỹ thuật quan trọng và được nghiên cứu nhiều. Đặc biệt là kỹ thuật khai thác luật kết hợp chính xác và không dư thừa, một số tác giả đã đề xuất khai thác luật kết hợp này từ tập sinh của tập đóng phổ biến trên dữ liệu giao dịch. Trong bài viết này, chúng tôi đề xuất giải thuật song song MCP-GCFI khai thác nhanh tập sinh của tập đóng phổ biến trên bộ xử lý đa nhân. Giải thuật đề xuất dễ dàng mở rộng trên nhiều hệ thống tính toán phân tán như Hadoop, Spark. Kết quả thực nghiệm trên bộ dữ liệu thực có mật độ dày của UCI và bộ dữ liệu giả lập có mật độ thưa của trung tâm nghiên cứu IBM Almaden, cho thấy giải thuật đề xuất hiệu quả.

Từ khóa - Bộ xử lý đa nhân; tập đóng phổ biến; tập sinh của tập đóng; giải thuật song song MCP-GCFI

1. Giới thiệu

Năm 1993, R. Agrawal đã đề xuất mô hình cơ bản khai thác luật kết hợp từ dữ liệu giao dịch dạng nhị phân theo hai giai đoạn [1, 2]: Phát sinh *tập phổ biến* (Frequent Itemset - FI) và sinh *luật kết hợp* từ tập phổ biến. Khi đó, số lượng itemset phổ biến được sinh ra là rất lớn và chiếm rất nhiều thời gian. Nhiều tác giả có đề xuất giải thuật sinh *tập đóng phổ biến* (Closed Frequent Itemset - CFI) [3-5], đây là tập không dư thừa và có kích thước nhỏ hơn FI rất nhiều lần ($CFI \subseteq FI$). Tuy nhiên, để sinh các *luật kết hợp chính xác* cũng như *không dư thừa*, nhiều tác giả có đề xuất khai thác tập luật kết hợp từ tập sinh của tập đóng phổ biến (Minimal/Generator Closed Frequent Itemset - mGCFI/GCFI) [6-10] như giải thuật GrGrowth [6], Zart [8] và DefMe [9, 10]. Các giải thuật trên tựa *Apriori*, tìm kiếm theo chiều sâu và dùng cấu trúc FP-Tree, IT-Tree – các giải thuật này tốn nhiều bộ nhớ/ đọc dữ liệu nhiều lần và không hiệu quả trên các loại dữ liệu giao dịch có mật độ dày hoặc thưa.

Trong nghiên cứu, tác giả đề xuất giải thuật khai thác song song *tập sinh*. Giải thuật đề xuất theo hướng tiếp cận song song dữ liệu và cả chức năng, dưới đây là các giải thuật liên quan trong bài viết:

- Phát sinh mảng IndexCOOC lưu trữ *items xuất hiện đồng thời* và *items xuất hiện trong ít nhất một giao dịch* với từng *item-hạt-nhân*;

- Dựa trên IndexCOOC phát sinh nLOOCTree;

- Giải thuật tuần tự SEQ-GCFI khai thác *tập sinh* của tập đóng phổ biến dựa trên cây nLOOCTree;

- Giải thuật song song MCP-GCFI khai thác nhanh *tập sinh* trên bộ xử lý đa nhân (BXLĐN).

Trong nghiên cứu này tác giả trình bày: Các vấn đề cơ bản về tập phổ biến, tập đóng phổ biến và *tập sinh* của tập đóng phổ biến; Giải thuật xây dựng mảng chứa *itemset* đồng xuất hiện và xuất hiện với *item-hạt-nhân* ít nhất trong một giao dịch,

Abstract - Association-rule mining is one of the most important and well-researched techniques of Data Mining. Particularly, the techniques of mining are exact association rules and non-redundant, some authors have proposed mining these association rules from generators of closed frequent item sets. In this paper, we propose the parallel MCP-GCFI algorithm to fast mining generators of closed frequent item sets on Multi-Core processor, as well as to expand the algorithm on distributed computing systems such as Hadoop, Spark. The experimental results show that the proposed algorithms perform faster than other existing algorithms on both real-life datasets of UCI and synthetic datasets generated by IBM Almaden.

Key words - Multi-core processor; closed frequent itemset; generator itemsets; parallel MCP-GCFI algorithm.

giải thuật sinh cây nLOOCTree và giải thuật tuần tự SEQ-GCFI khai thác *tập sinh*; Trên cơ sở giải thuật tuần tự xây dựng giải thuật song song MCP-GCFI khai thác hiệu năng của bộ xử lý đa nhân; Và cuối cùng là kết quả thực nghiệm, cũng như hướng phát triển cho các nghiên cứu tiếp theo.

2. Vấn đề cơ bản về tập phổ biến

Cho $I = \{i_1, i_2, \dots, i_m\}$ là tập gồm m thuộc tính, mỗi thuộc tính được gọi là *item*. Tập các $itemX = \{i_1, i_2, \dots, i_k\}, \forall i_j \in I (1 \leq j \leq k)$ được gọi là *itemset*, itemset gồm có k item thì được ký hiệu là k -itemset. D là dữ liệu chứa các giao dịch, gồm có n bản ghi dữ liệu được gọi là tập chứa các giao dịch $T = \{t_1, t_2, \dots, t_n\}$, với giao dịch $t_k = \{i_{k1}, i_{k2}, \dots, i_{km}\}, i_{kj} \in I (1 \leq k_j \leq m)$.

Định nghĩa 1: Cho itemset $X \subseteq I$ số lượng các giao dịch trong D có chứa X được gọi là *độ phổ biến* (support) của itemset X . Ký hiệu, $sup(X)$.

Định nghĩa 2: Cho itemset $X \subseteq I$, nếu $sup(X) \geq minsup$ thì X được gọi là *itemset phổ biến* ($minsup$: ngưỡng phổ biến tối thiểu cho trước).

Định nghĩa 3: Cho $X \in FI$, X gọi là *itemset* phổ biến đóng nếu X không có tập cha cùng độ phổ biến. Ký hiệu, CFI là tập chứa các *itemset* phổ biến đóng.

Định nghĩa 4: Cho $X \in CFI$, tất cả các *itemset* con thực sự của X có *bằng độ phổ biến* với X gọi là *itemset* sinh của X . Ký hiệu, GCFI là tập chứa các *itemset* sinh.

Cho tập dữ liệu giao dịch D như trong Bảng 1.

Bảng 1. Dữ liệu D sử dụng cho Ví dụ

TID	Items							
	O	Q	R	S	U	V	W	Z
t1	1		1				1	
t2	1		1		1	1		
t3	1		1	1		1	1	

t_4					1			1
t_5					1			
t_6	1		1		1		1	
t_7	1		1	1				
t_8	1	1	1		1			
t_9	1		1		1	1	1	
t_{10}	1	1	1		1		1	

Dữ liệu ở Bảng 1, gồm 8 items $I = \{O; Q; R; S; U; V; W; Z\}$ và có 10 giao dịch $T = \{t_1; t_2; t_3; t_4; t_5; t_6; t_7; t_8; t_9; t_{10}\}$.

Bảng 2. Sinh tập CFI và GCFI trên D với $minsup = 3$

k-itemset	Tập CFI (#CFI=6)	Tập GCFI (#GCFI=13)
1	U	V, W, O, R
2	OR	VO, VR, WO, WR, WU, UO, UR
3	VOR, WOR, UOR	WUO, WUR
4	WUOR	

Bảng 2, liệt kê theo k -itemset của tập CFI và GCFI ở ngưỡng $minsup = 3$. Số phần tử trong các tập lần lượt $|CFI| = 6$ và $|GCFI| = 13$. Tỷ suất $|CFI|/|GCFI| = 6/13 \times 100\% \approx 46\%$.

3. Giải thuật đề xuất

3.1. Tập chiếu và items xuất hiện ít nhất trên cùng một giao dịch với item-hạt-nhân có thứ tự [11]

Tập chiếu của $item_i$ trên dữ liệu giao dịch D : $\pi(i_k) = \{t \in D \mid i_k \in t\}$ là tập các giao dịch có chứa item i_k .

$$sup(i_k) = |\pi(i_k)| \quad (1)$$

Tập chiếu của itemset $X = \{i_1, i_2, \dots, i_k\}$, $\forall i_j \in I$ ($1 \leq j \leq k$): $\pi(X) = \{\pi(i_1) \cap \pi(i_2) \dots \pi(i_k)\}$.

$$sup(X) = |\pi(X)| \quad (2)$$

Để rút gọn không gian sinh, tác giả đưa ra Định nghĩa 5 và 6 ($P_k(X)$ – powerset của X có k item):

Định nghĩa 5: Cho $i_k \in I$ ($i_1 < i_2 < \dots < i_m$) thứ tự theo sup , ta gọi i_k là item-hạt-nhân. $X_{lexcooc} \subseteq I$ gọi xuất hiện đồng thời có thứ tự với i_k : $X_{lexcooc}$ là tập các item có thứ tự theo sup , xuất hiện đồng thời với i_k nghĩa là $\pi(i_k) \equiv \pi(i_k \cup X_{lexcooc})$, $\forall X_{lexcooc} \in P_{\geq 1}(X_{lexcooc})$, $\forall i_j \in X_{lexcooc}, i_k < i_j$. Ký hiệu, $lexcooc(i_k) = X_{lexcooc}$.

Định nghĩa 6: Cho $i_k \in I$ ($i_1 < i_2 < \dots < i_m$) thứ tự theo sup , ta gọi i_k là item-hạt-nhân. $X_{lexlooc} \subseteq I$ gồm items có thứ tự và xuất hiện cùng với i_k trong ít nhất một giao dịch, nhưng không xuất hiện đồng thời: $1 \leq |\pi(i_k \cup X_{lexlooc})| < |\pi(i_k)|$, $\forall X_{lexlooc} \in P_{\geq 1}(X_{lexlooc})$, $\forall i_j \in X_{lexlooc}, i_k < i_j$. Ký hiệu, $lexlooc(i_k) = X_{lexlooc}$.

Trong mục này, tác giả trình bày giải thuật sinh mảng IndexCOOC. Mỗi phần tử trong mảng IndexCOOC có 4 trường thông tin:

- **IndexCOOC[k].item**: lưu trữ item-hạt-nhân i_k ;
- **IndexCOOC[k].sup**: lưu trữ độ phổ biến của item-hạt-nhân i_k ;
- **IndexCOOC[k].cooc**: items xuất hiện đồng thời cùng với item-hạt-nhân i_k ;
- **IndexCOOC[k].looc**: items xuất hiện với item-hạt-nhân i_k trong ít nhất là một giao dịch.

Giải thuật 1. Tạo dựng mảng IndexCOOC

Đầu vào: Tập dữ liệu D .

Đầu ra: IndexCOOC, Ma trận BiM

1. Với từng phần tử k của IndexCOOC:
2. **IndexCOOC[k].item** = i_k
3. **IndexCOOC[k].sup** = 0
4. **IndexCOOC[k].cooc** = $2^m - 1$
5. **IndexCOOC[k].looc** = 0
6. Với từng giao dịch t_i thực hiện:
7. Lưu t_i vào **BiM**
8. Với từng item k có trong t_i thực hiện:
9. **IndexCOOC[k].cooc** = $vectorbit(t_i)$
10. **IndexCOOC[k].looc** = $vectorbit(t_i)$
11. **IndexCOOC[k].sup** ++
12. Mảng IndexCOOC được sắp xếp tăng theo sup
13. Với từng phần tử k của IndexCOOC:
14. **IndexCOOC[k].cooc** = $lexcooc(i_k)$
15. **IndexCOOC[k].looc** = $lexlooc(i_k)$
16. Trả về IndexCOOC, BiM

Minh họa giải thuật 1: thực hiện từ dòng 1 đến 11.

Khởi tạo đầu tiên cho mảng IndexCOOC: (cooc và looc được minh họa theo hexa) số item từ dữ liệu D đã cho ở Bảng 1 là $m = 8$.

item	O	Q	R	S	U	V	W	Z
sup	0	0	0	0	0	0	0	0
cooc	0xFF	0xFF	0xFF	0xFF	0xFF	0xFF	0xFF	0xFF
looc	0x00	0x00	0x00	0x00	0x00	0x00	0x00	0x00

Duyệt giao dịch đầu tiên t_1 : {O, R, V} ở dạng bit tương ứng là **10100010**(0xA2)

item	O	Q	R	S	U	V	W	Z
sup	1	0	1	0	1	1	0	0
cooc	0xA2	0xFF	0xA2	0xFF	0xA2	0xA2	0xFF	0xFF
looc	0xA2	0x00	0xA2	0x00	0xA2	0xA2	0x00	0x00

Tương tự, duyệt giao dịch cuối cùng t_{10} : {O, Q, R, U, W} ở dạng bit tương ứng là **11101010**(0xEA)

item	O	Q	R	S	U	V	W	Z
sup	8	2	8	2	7	3	5	1
cooc	0xA0	0xE8	0xA0	0xB0	0x08	0xA4	0xA2	0x09
looc	0xFE	0xEA	0xFE	0xB6	0xEF	0xBE	0xFE	0x09

Dòng 12, sắp xếp IndexCOOC tăng dần theo sup của từng item, ta có kết quả:

item	Z	Q	S	V	W	U	O	R
sup	1	2	2	3	5	7	8	8
cooc	U	O, R, U	O, R	O, R	O, R	∅	R	O
looc	∅	W	V, W	S, U, W	Q, S, U, V	O, Q, R, V, W, Z	Q, S, U, V, W	Q, S, U, V, W

Từ dòng 13 đến 15 – cho kết quả rút gọn ở Bảng 3.

Bảng 3. Mảng IndexCOOC có thứ tự tăng dần theo độ phổ biến của item, đồng thời cooc và looc cũng có thứ tự

item	Z	Q	S	V	W	U	O	R
sup	1	2	2	3	5	7	8	8
cooc	U	O, R, U	O, R	O, R	O, R	∅	R	∅
looc	∅	W	V, W	W, U	U	O, R	∅	∅

Xét itemW, $lexcooc(W) = \{O, R\}$ sinh tập $GCFI_{|W|} = \{(\underline{W}, 5), (\underline{WO}, 5), (\underline{WR}, 5)\}$ (dòng 3) xem cây $nLOOCTree(W)$: có đường đi đơn $\{W \rightarrow U\}$ với

$sup(\underline{WU}) = 3 \geq minsup$. Vì vậy, $GCFI_{[W]} = \cup\{(\underline{WU}, 3), (\underline{WUO}, 3), (\underline{WUR}, 3)\}$ (dòng 4 đến 10).

Xét $itemU$, $lexcooc(U) = \{\emptyset\}$ nên $GCFI_{[U]} = \{\emptyset\}$, xem cây $nLOOCtree(U)$: có một đường đi đơn $\{U \rightarrow O \rightarrow R\}$ với $sup(\underline{UOR}) = 5 \geq minsup$. Vì vậy, $GCFI_{[U]} = \{(\underline{UO}, 5), (\underline{UR}, 5)\}$.

Xét $itemO$, $lexcooc(O) = \{R\}$ sinh tập $GCFI_{[O]} = \{(\underline{O}, 8)\}$ (dòng 3) và $lexcooc(O) = \{\emptyset\}$, nên không sinh $GCFI$ từ $nLOOCtree(O)$ (dòng 4).

Xét $itemR$, đồng xuất hiện hoàn toàn với $itemO$: $GCFI_{[R]} = \{(\underline{R}, 8)\}$ (dòng 11).

Tập sinh $GCFI$ từ dữ liệu D ở Bảng 1 với $minsup = 3$:

Bảng 4. Tập sinh $GCFI$ trên D với $minsup = 3$

item	Tập sinh $GCFI$ (# $GCFI = 13$)		
V	(V, 3)	(VO, 3)	(VR, 3)
W	(W, 5)	(WO, 5)	(WR, 5)
	(WU, 3)	(WUO, 3)	(WUR, 3)
U	(UO, 5)	(UR, 5)	
O	(O, 8)		
R	(R, 8)		

4. Giải thuật song song MCP-GCFI

Ngày nay, nhiều máy tính cá nhân và máy trạm có trên hai nhân cho phép nhiều luồng xử lý được thực hiện đồng thời - điều này làm cho máy tính có được tốc độ xử lý nhanh và khả năng đa nhiệm tốt hơn. Để tận dụng hiệu năng của BXLĐN, cần phân phối xử lý đồng thời trên nhiều nhân cho nhiều pha/ bài toán để tiết kiệm thời gian và nâng cao hiệu năng.

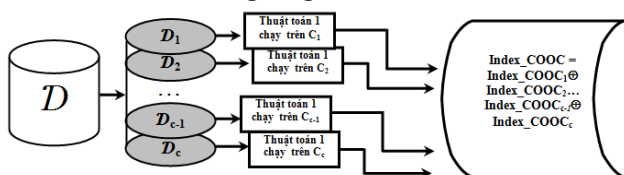
Tác giả xây dựng giải thuật song song MCP-GCFI (Multi Core Processor – Generators of Closed Frequent Itemsets) khai thác tập sinh trên BXLĐN dựa trên giải thuật tuần tự SEQ-GCFI.

Giải thuật tuần tự SEQ-GCFI, gồm 2 pha chính:

Pha 1: Xây dựng mảng IndexCOOC;

Pha 2: Giải thuật 3 khai thác tập sinh từ mảng IndexCOOC.

Bước thứ nhất, song song hóa Pha 1 theo sơ đồ Hình 2.



Hình 2. Sơ đồ song song hóa Pha 1

Hình 2, phân chia dữ liệu D thành c tập dữ liệu con $D_1, D_2, \dots, D_{c-1}, D_c$ ứng với từng nhân C_i thực hiện giải thuật 1 với đầu vào là dữ liệu D_i và đầu ra là mảng $IndexCOOC_i$ tương ứng. Để tính mảng $IndexCOOC$ cho D , thực hiện phép tính:

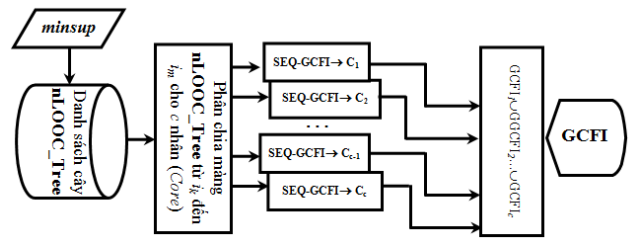
$$IndexCOOC = IndexCOOC_1 \oplus IndexCOOC_2 \oplus \dots \quad (3)$$

Sau đó, mảng $IndexCOOC$ được xếp tăng theo sup và chuẩn hóa theo dòng 13, 14 và 15 của giải thuật 1.

Thí dụ 2: Giả sử D được chia thành 2 tập - D_1 gồm có

$\{t1, t2, t3, t4, t5\}$ và D_2 là $\{t6, t7, t8, t9, t10\}$.

D_1 chạy giải thuật 1 trên nhân C_1 và trả về $IndexCOOC_1$:



Hình 3. Sơ đồ song song hóa Pha 2

Hình 3, phân chia mảng $IndexCOOC$ từ i_1 đến i_m thành c phần ứng với từng nhân C_j thực hiện giải thuật SEQ-GCFI với đầu vào là mảng $IndexCOOC$ từ phần tử thứ $[(j-1)*(m \div c) + 1]$ đến phần tử thứ $[j*(m \div c)]$ và đầu ra là tập sinh $GCFI_j$ tương ứng. Tập sinh cho tập dữ liệu giao dịch D , ta thực hiện theo phương trình:

$$GCFI_D = GCFI_1 \cup GCFI_2 \dots \cup GCFI_c \quad (4)$$

Thí dụ 3: Cho D được mô tả ở Bảng 1, $minsup = 3$. Kết thúc song song hóa Pha 1, có như Bảng 3.

Song song hóa Pha 2: Nhân C_1 chạy giải thuật SEQ-GCFI từ $itemV$ đến U khai thác itemset sinh $GCFI_1$:

$GCFI_{1[V]}$	(V, 3)	(VO, 3)	(VR, 3)
$GCFI_{1[W]}$	(W, 5)	(WO, 5)	(WR, 5)
	(WU, 3)	(WUO, 3)	(WUR, 3)
$GCFI_{1[U]}$	(UO, 5)	(UR, 5)	

Nhân C_2 chạy giải thuật SEQ-GCFI từ $itemO$ đến R khai thác itemset sinh $GCFI_2$:

$GCFI_{2[O]}$	(O, 8)		
$GCFI_{2[R]}$	(R, 8)		

Tập sinh $GCFI$ trên D với $minsup = 3$ được tính $GCFI_D = GCFI_1 \cup GCFI_2$, ta có kết quả như Bảng 4.

5. Kết quả thực nghiệm

Giải thuật đề xuất được thực nghiệm cài đặt trên máy tính cấu hình: CPU Core Duo 2.0 GHz, bộ nhớ 4GB; ngôn ngữ lập trình C#(VS 2010).

Giải thuật được thực nghiệm trên 2 nhóm dữ liệu:

- Dữ liệu thu thập thực tế với mật độ dày: 2 tập Chess và Mushroom từ kho lưu trữ UCI.

- Dữ liệu chạy giả lập với mật độ thưa: 2 tập dữ liệu giả lập T10I4D100K và T40I10D100K từ trung tâm Almaden của IBM.

Bảng 5. Dữ liệu thực nghiệm

D	Số item	Số item nhỏ-nhất	Số item lớn-nhất	Số item trung-bình
Chess	75	37	37	37
Mushroom	119	23	23	23
T10I4D100K	870	1	29	10
T40I10D100K	942	4	77	40

Trong nghiên cứu, tác giả đề xuất giải thuật khai thác song song tập sinh. Đây là nghiên cứu đề xuất theo hướng

tiếp cận song song, nên chưa có giải thuật có cùng hướng tiếp cận nhằm đánh giá hiệu năng giải thuật. Vì vậy, tác giả đề xuất đánh giá hiệu năng giải thuật song song như sau: So sánh giải thuật tuần tự **SEQ-GCFI** với giải thuật **GrGrowth** [6] và giải thuật **Zart** [8]. Sau đó, so sánh hiệu suất của giải thuật song song **MCP-GCFI** với giải thuật tuần tự **SEQ-GCFI**. Trong thực nghiệm, tác giả tiến hành đánh giá theo từng ngưỡng tối thiểu *minsup* và tất cả 4 giải thuật trên đều cho kết quả như nhau.

Hiệu suất chạy giải thuật song song trên BXLĐN:

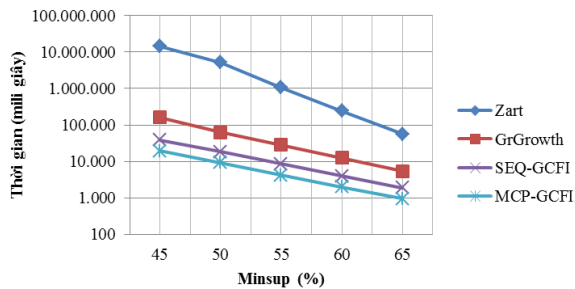
$$HS = 1 - \left(T_M - \frac{T_S}{c} \right) / \left(\frac{T_S}{c} \right) \quad (5)$$

Trong đó:

- T_S : thời gian thực hiện tuần tự;
- T_M : thời gian thực hiện song song trên BXLĐN;
- c : số lượng nhân của CPU (số core).

Phương trình (5): Đánh giá hiệu suất giải thuật song song **MCP-GCFI** so với giải thuật tuần tự **SEQ-GCFI**.

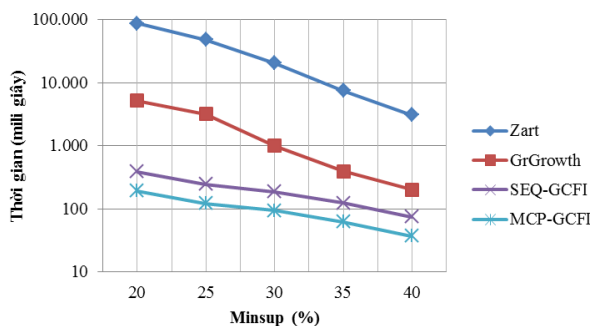
Chess



Hình 4. Thời gian thực hiện khai thác GCFI trên Chess

Hình 4 - kết quả được thực nghiệm từ tập dữ liệu thực **Chess** mật độ dày đặc (49,3%) cho thấy, giải thuật tuần tự **SEQ-GCFI** nhanh hơn giải thuật **GrGrowth**, **Zart** và giải thuật **MCP-GCFI** có thời gian chạy nhanh hơn giải thuật **SEQ-GCFI**. Hiệu suất trung bình của giải thuật **MCP-GCFI** là 76,1% và độ lệch chuẩn 1,8%.

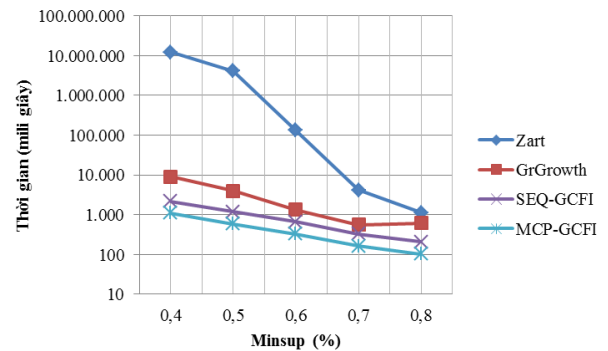
Mushroom



Hình 5. Thời gian thực hiện khai thác GCFI trên Mushroom

Hình 5 - kết quả được thực nghiệm từ tập dữ liệu thực **Mushroom** mật độ dày (19,3%) cho thấy, giải thuật tuần tự **SEQ-GCFI** nhanh hơn giải thuật **GrGrowth**, **Zart** và giải thuật **MCP-GCFI** có thời gian chạy nhanh hơn giải thuật **SEQ-GCFI**. Hiệu suất trung bình của giải thuật **MCP-GCFI** là 74,3% và độ lệch 1,5%.

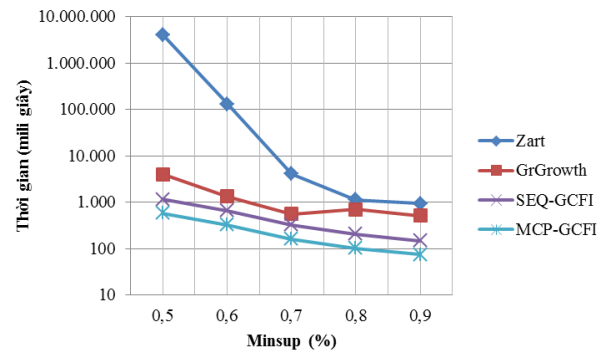
T10I4D100K



Hình 6. Thời gian thực hiện khai thác GCFI trên T10I4D100K

Hình 6 - kết quả được thực nghiệm từ tập dữ liệu giả lập **T10I4D100K** mật độ rất thưa (1,1%) cho thấy, giải thuật tuần tự **SEQ-GCFI** nhanh hơn giải thuật **GrGrowth**, **Zart** và giải thuật **MCP-GCFI** có thời gian chạy nhanh hơn giải thuật tuần tự **SEQ-GCFI**. Hiệu suất trung bình của **MCP-GCFI** là 80,1% và độ lệch 1,4%.

T40I10D100K



Hình 7. Thời gian thực hiện khai thác GCFI trên T40I10D100K

Hình 7 - kết quả được thực nghiệm từ tập dữ liệu giả lập **T40I10D100K** mật độ thưa (4,2%) cho thấy, giải thuật tuần tự **SEQ-GCFI** nhanh hơn giải thuật **GrGrowth**, **Zart** và giải thuật chạy trên BXLĐN là **MCP-GCFI** có thời gian chạy nhanh hơn giải thuật tuần tự **SEQ-GCFI**. Hiệu suất trung bình của giải thuật **MCP-GCFI** là 81,5% và độ lệch chuẩn 1,6%.

Qua thực nghiệm cho thấy, giải thuật khai thác song song tập sinh **MCP-GCFI** trên BXLĐN tốt hơn rất nhiều so với giải thuật **GrGrowth**, **Zart**. Giải thuật **MCP-GCFI** cần được thực nghiệm mở rộng trên các dữ liệu giao dịch có kích cỡ lớn và so sánh thêm với các giải thuật chạy trên hệ thống phân tán như Hadoop, Spark.

6. Kết luận

Bài viết đã trình bày giải thuật tuần tự **SEQ-GCFI** khai thác tập sinh của tập đóng phổ biến gồm ba bước: Đầu tiên là phát sinh nhanh cấu trúc mảng **IndexCOOC** có chứa items xuất hiện đồng thời với item-hạt-nhân và items xuất hiện ít nhất với item-hạt-nhân trong một giao dịch; Bước 2: Xây dựng **nLOOCTree** dựa vào mảng **IndexCOOC**; Bước 3: Khai thác nhanh tập sinh của tập đóng phổ biến dựa trên cây **nLOOCTree**. Từ giải thuật tuần tự **SEQ-**

GCFI, tác giả mở rộng và song song hóa thực hiện trên BXLĐN gọi là giải thuật **MCP-GCFI**. Hiệu suất trung bình khi song song hóa trên 4 tập dữ liệu thực nghiệm là 78% và độ lệch chuẩn 3,3%.

Trong các nghiên cứu tiếp theo, tác giả hướng đến việc mở rộng giải thuật **MCP-GCFI** để khai thác nhanh *tập sinh* trên hệ thống điện thoại thông minh đa lõi có tài nguyên hạn chế, cũng như thực nghiệm mở rộng trên hệ thống phân tán phổ biến hiện nay như Hadoop, Spark.

TÀI LIỆU THAM KHẢO

- [1] R. Agrawal, T. Imilinski and A. Swami, *Mining association rules between sets of large databases*, Proc. of the ACM SIGMOD Int Conf on Management of Data., 1993, pp. 207-216.
- [2] P. Huan and L. Bac, *A Novel Algorithm for Frequent Itemsets Mining in Transactional Databases*, In: Trends and Applications in Knowledge Discovery and Data Mining. PAKDD 2018. LNCS, 11154, Springer Cham, 2018, pp. 243–255.
- [3] Y. Bastide, N. Pasquier, R. Taouil, G. Stumme, and L. Lakhal, *Mining Minimal Non-Redundant Association Rules using Closed Frequent Itemsets*, In: 1st International Conference on Computational Logic, 2000, pp.972 – 986.
- [4] M. J. Zaki, “Mining Non-Redundant Association Rules”, *Data Mining and Knowledge Discovery*, 9, Springer, 2004, pp. 223–248.
- [5] P. Huan, *NOV-CFI: A Novel Algorithm for Closed Frequent Itemsets Mining in Transactional Databases*, ICNCC 2018, ACM, NY, USA, 2018, pp. 58-63.
- [6] G. Liu, J. Li and L. Wong, “A new concise representation of frequent itemsets using generators and a positive border”, *Knowl. Inf. Syst.*, 17(1), Springer, 2008, pp. 35-56.
- [7] T. Hamrouni, “Key roles of closed sets and minimal generators in concise representations of frequent patterns”, *Intell. Data Anal.*, 16(4), 2012, pp.581-631.
- [8] L. Szathmary, A. Napoli and S. O. Kuznetsov, *ZART: A Multifunctional Itemset Mining Algorithm*, The 6th Intl Conf on Concept Lattices and Their Applications, 2008, pp. 47–58.
- [9] A. Soulet and F. Rioult, *Efficiently Depth-First Minimal Pattern Mining*, In: Tseng V.S., Ho T.B., Zhou ZH., Chen A.L.P., Kao HY. (eds) *Advances in Knowledge Discovery and Data Mining. PAKDD 2014*. LNCS. 8443, Springer Cham, 2014, pp. 28-39.
- [10] A. Soulet and F. Rioult, *Exact and Approximate Minimal Pattern Mining*, In: Guillet F., Pinaud B., Venturini G. (eds) *Advances in Knowledge Discovery and Management. SCI. 665*, Springer Cham, 2017, pp. 61-81.
- [11] Phan Thành Huấn, Lê Hoài bắc, *Thuật toán hiệu quả khai thác tập hiếm tối thiểu*, FAIR – Nghiên cứu cơ bản và ứng dụng, 2018, pp. 497-505.

(BBT nhận bài: 21/10/2019, hoàn tất thủ tục phân biên: 08/5/2020)