

Research Article

Frequent Pattern Mining of Eye-Tracking Records Partitioned into Cognitive Chunks

Noriyuki Matsuda¹ and Haruhiko Takeuchi²

¹ Department of Social Systems & Management, University of Tsukuba, Tsukuba 305-8573, Japan

² National Institute of Advanced Industrial Science & Technology (AIST), Tsukuba 305-8566, Japan

Correspondence should be addressed to Haruhiko Takeuchi; takeuchi.h@aist.go.jp

Received 23 July 2014; Accepted 27 October 2014; Published 23 November 2014

Academic Editor: Yongqing Yang

Copyright © 2014 N. Matsuda and H. Takeuchi. This is an open access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Assuming that scenes would be visually scanned by chunking information, we partitioned fixation sequences of web page viewers into chunks using isolate gaze point(s) as the delimiter. Fixations were coded in terms of the segments in a 5×5 mesh imposed on the screen. The identified chunks were mostly short, consisting of one or two fixations. These were analyzed with respect to the within- and between-chunk distances in the overall records and the patterns (i.e., subsequences) frequently shared among the records. Although the two types of distances were both dominated by zero- and one-block shifts, the primacy of the modal shifts was less prominent between chunks than within them. The lower primacy was compensated by the longer shifts. The patterns frequently extracted at three threshold levels were mostly simple, consisting of one or two chunks. The patterns revealed interesting properties as to segment differentiation and the directionality of the attentional shifts.

1. Introduction

Eyes seldom stay completely still. They continually move even when one tries to fixate one's gaze on an object because of the tremors, drifts, and microsaccades that occur on a small scale [1]. Hence, researchers need to infer a fixation from consecutive gaze points clustered in space [2]. We may regard such a cluster of gaze points as a perceptual chunk, a familiar term in psychology after Miller [3] in referring to a practically meaningful unit of information processing.

During fixation, people closely scan a limited part of the scene they are interested in. They then quickly move their eyes to the next fixation area by saccade, which momentarily disrupts vision. However, it normally goes unnoticed thanks to our vision system that produces continuous transsaccades perception [4–6]. It means that successive fixations constitute a higher order chunking over and above the primary chunking of gaze points. Put metaphorically, the (gaze, fixation, fixation-chunking) relationship is analogous to the (letter, word, phrase) relationship. For the sake of brevity, a chunk of fixations will be referred to as a chunk.

In viewing natural scenes or displays, a chunk continues to grow until interrupted by one or more isolate gaze points resulting from drifting attention or by accident. These do not participate in any fixation. Whatever causes the interruption, we believe that such isolate points serve as chunk delimiters, like the pauses in speech. As a pause can be either short or long, interruptions by isolate points can vary in length. Figure 1 illustrates two levels of chunking: (a) chunking of gaze points into fixations and (b) chunking of consecutive fixations with and without interruption.

Granting our conjecture, one may still wonder what particular merits will accrue from the analysis of chunks in lieu of ordinary plain fixation sequences. The expected merits are twofold: separation of between- and within-chunk patterns and extraction of common patterns across records. Neither of these is attainable when dealing with multiple records by heat maps of fixations accumulated with no regard to sequential connections [7], by network analysis of the adjacent transitions accumulated within and between records [8–10], or by scan paths that would be too complicated [11] unless reduced to frequently shared subpaths. The key to understanding this

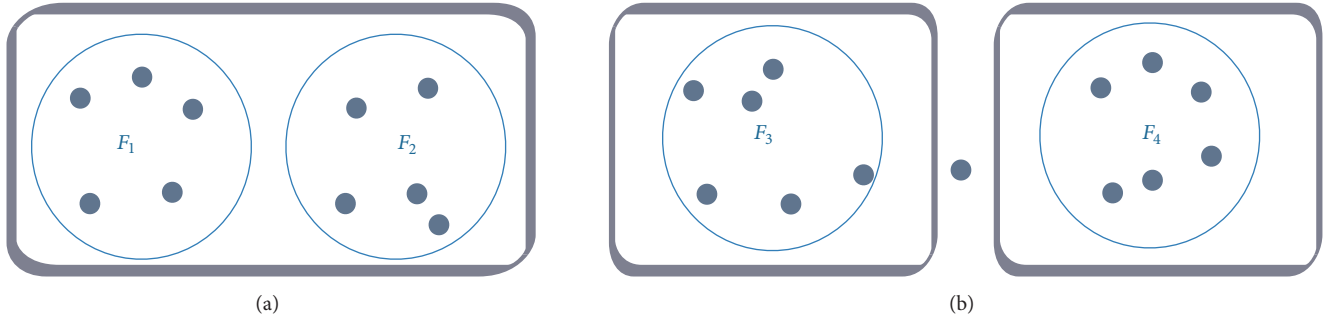


FIGURE 1: Two fixations in one chunk (a) and in separate chunks (b).

point lies in the structure of fixation sequences as explained below.

1.1. Structure of Fixation Sequences. Equation (1) presents two types of fixation sequences, one plain and the other partitioned, both arranged in time sequence. The former serves as a basis for heat maps, scan paths, and network analysis. The latter incorporates chunks delimited by isolate gazes or any other appropriate criterion. The essential nature of the sequences remains the same when fixations are coded in areas of interest (AOI) or grid-like segments.

Plain and Partitioned Fixation Sequences. Consider

$$\text{Plain: } [F_1 \ F_2 \ F_3 \ \cdots \ F_i \ F_{i+1} \ F_{i+2} \ \cdots]$$

$$\text{Partitioned: } [[F_1 \ F_2], [F_3], \dots, [F_i \ F_{i+1} \ F_{i+2} \ \cdots], \dots], \quad (1)$$

where F_i denotes the i th fixation.

Although it was not explicitly stated, McCarthy et al. [12] in effect extracted chunks from partitioned sequences in their work on the importance of web page objects and their locations. They grouped consecutive fixations within each AOI into a chunk called a glance to obtain plain sequences of glances coded in AOI. Their interest was to see how often areas of web pages would attract glances by varying the area locations and the types of tasks.

By focusing on the frequency of glances as an indication of importance, they disregarded the length of the chunks, that is, the number of fixations within glances. Also disregarded was the shift of glances, that is, between-chunk sequences. To us, both within- and between-chunk patterns seem to contain rich information worthy of investigation. The information can be extracted from partitioned sequences but not from plain ones. In addition, partitioned sequences will be of great value when some AOIs are nested into broader AOIs (see [16]), given the appropriate coding. The present study is extensible to such a hierarchical structure.

For the sake of simplicity, we will focus on the eye movements of web page viewers, and we will assume that the pages are divided into grid-like AOIs, that the fixations are coded in terms of the areas in which they fall, and that chunks are delimited by isolate gaze points.

1.2. Shifts of Interest within and between Chunks. The distance between two successive fixations indicates how far the interest

TABLE 1: Pattern extraction by prefix “a” at ms3.

Record	Initial sequences	Patterns prefixed by <i>a</i>
1	$[a][abc][ac][d][cf]$	$[abc][ac][d][cf]$
2	$[ad][c][bc][ae]$	$[_d][c][bc][ae]$
3	$[ef][ab][df][c][b]$	$[_b][df][c][b]$
4	$[e][g][af][c][b][c]$	$[_f][c][b][c]$
frequent code	$a/4, b/4/, c/4, d/4, e/3, f/3$	$b/4, c/4$

Note. The underscore $_$ means that the prefix was present in the chunk; for example, $_b$ implies ab .

shifted or did not shift in a looped transition that represents sustained interest in a given area. In our view, a chunk of fixations reflects continuous interest, and a new one begins after a momentary drift of the gaze. It seems natural to expect the distance distribution of the within-chunk shifts to differ, to some extent, from that of the between-chunk shifts.

The distance analysis explained above exploits information from the cumulative records across all viewers. Hence, it is possible that the results are influenced by some dominant patterns in particular records. If one is interested in sequential regularities often shared among records, frequent sequential pattern mining is useful, as explained below.

1.3. Frequent Sequential Pattern Mining. Among others, we will employ PrefixSpan, developed by Pei et al. [13, 14], because of its conceptual compatibility with the partitioned sequences of eye-tracking data. Their approach is briefly explained below using their example, seen in Table 1. (See the Appendix for a more formal explanation.) One can view the data as the eye-tracking records of four viewers in which fixations are alphabetically coded according to the areas of interest (AOI) they fall into: a, b, c, d, e, f , and g .

Codes a, b, c, d, e , and f are all frequent, shared by the majority, whereas code g is infrequent, appearing only once. For further scanning, any infrequent or rare code is to be removed from the records, since it will never appear in frequent patterns according to the *a priori* principle [15]. Let us set the level of being frequent at three for illustrative purposes. This level is called the minimum support threshold (abbreviated as ms).

TABLE 2: All of the frequent patterns extracted from Table 1 at ms3.

<u>[a]</u>	<u>[a][b]</u>	<u>[a][c]</u>	<u>[a][c][b]</u>	<u>[a][c][c]</u>	<u>[b]</u>	<u>[b][c]</u>
<u>[c]</u>	<u>[c][b]</u>	<u>[c][c]</u>	<u>[d]</u>	<u>[d][c]</u>	<u>[e]</u>	<u>[f]</u>

Note. Underscored patterns were extracted at ms4.

For every frequent code, one scans the reduced records, devoid of infrequent codes, for patterns prefixed by the given code. Those found for prefix “a” are listed in the second column of Table 1. These are subject to further scanning with respect to the frequent codes at this step, that is, *b* and *c*. This process recursively continues until no code is frequent or no patterns remain in the records. Note that prefixes grow in each step like “[a][b]”, “[a][c]” in the above example (see the Appendix for a more formal explanation).

Table 2 lists 14 frequent patterns extracted at ms3 from the initial record, including those found at ms4 as an embedded part. For instance, [a][c], found at ms4, is embedded in patterns [a][c] and [a][c][c], at ms3; that is,

$$[a][c] \subseteq \{[a][c], [a][c][c]\}. \quad (2)$$

Similarly, those found at ms3 are included in the patterns at ms2 reported by Pei et al. [13, 14]. Inclusive relations generally hold between different ms levels.

Ordinarily, one finds too few patterns at a high ms level and too many at a low level to make an interesting analysis. However, once one recognizes the inclusive relations, making use of multiple levels becomes a plausible solution for identifying strongly frequent patterns as opposed to mildly and weakly frequent ones. (See the Appendix for the relation networks among the patterns identified at ms2, ms3, and ms4.)

The present approach is expected to advance eye-tracking research along with conventional heat maps, scan paths, and network analysis recently developed by Matsuda and Takeuchi [8–10].

2. Method

2.1. Subjects (Ss). Twenty residents (seven males and 13 females) living near the AIST Research Institute in Japan were recruited for the experiments. They had normal or corrected vision, and their ages ranged from 19 to 48 years (average 30). Ten of the Ss were university students, five were housewives, and the rest were part-time workers. Eleven Ss were heavy Internet users, while the rest were light users, as judged from their reports about the number of hours they spent browsing online in a week.

2.2. Stimuli. The front (or top) pages of ten commercial web sites were selected from various business areas: airline companies, commerce and shopping, and banking. These were classifiable into three groups according to the layout types [8–10]. Due to space limits, we chose four pages with the same layout, the top and the principal layers. The principal layers were divided into the main area in the middle and subareas on both sides. The layers and the areas differed in size among pages.



FIGURE 2: Segment coding.

2.3. Apparatus and Procedure. The stimuli were presented with 1024×768 pixel resolution on a TFT 17" display in a Tobii 1750 eye-tracking system at a rate of 50 Hz. The web pages were randomly displayed to the Ss one at a time for 20 sec. The Ss were asked to browse each page at their own pace. The translated instructions are “Various web pages will be shown on the computer display in turn. Please look at each page as you usually do until the screen darkens and then, click the mouse button when you are ready to proceed.” The Ss were informed that the experiment would last for approximately five minutes.

2.4. Segment Coding. A 5×5 mesh was superposed on the effective part of each page, after the page was stripped of white margins that had no text or graphics. A uniform mesh was employed for ease of comparison among pages that varied in design beyond the basic layout. The distance of a shift between two segments was measured by the Euclidean distance, computed as the square root of $m^2 + n^2$, where *m* and *n* are the number of blocks (i.e., segments) moved along the horizontal and vertical axes.

The rows (and columns) of the mesh were alphabetically (and numerically) labeled in descending order: A through E (and 1 through 5). The segments were coded by combining these labels as seen in Figure 2: A1, A2, ..., A5 for the first row; B1, ..., B5 for the second; and so on through E1, ..., E5 for the fifth row.

2.5. Fixation Sequences. The raw tracking data for each subject consisted of time-stamped gaze points measured in *xy*-coordinates. The gaze points were grouped into a fixation point if they stayed within a radius of 30 pixels for 100 msec. Otherwise, they remained isolate.

Each fixation was then translated into code sequences according to the segments in which the fixation fell. Finally, each fixation sequence was partitioned into chunks using the isolate gaze points as delimiters.

2.6. Preprocessing the Codes for PrefixSpan. In accord with the algorithm, 25 segments were first recoded using letters *a* through *y*; then the codes in each chunk were alphabetically ordered with no duplication. In this process, we represented within-chunk loops by extra recoding. Consecutively repeated codes within a chunk were replaced by the corresponding capital letter, for example, [caaababaa] to [cAbabA]. After eliminating duplicates, we sorted the codes within each chunk, for example, [Aabc] from the original sequence. Consequently, we maintained the sequential order among chunks, but the within-chunk sequences could have been distorted. Due to this possibility, we were unable to identify between-chunk loops.

Frequent patterns were extracted at three levels of minimum support (denoted as ms12, ms14, and ms16) corresponding to 60, 70, and 80% of the subjects.

3. Results

The four pages used as stimuli will be referred to as P1, P2, P3, and P4.

3.1. Examination of the Chunks. The total number of chunks did not greatly differ among pages, ranging from 539 (P2) to 592 (P1). The pages agreed well on the lengths and proportions of primary, secondary, and tertiary chunks that contained one, two, and three fixations, respectively. Primary chunks accounted for 53.3 (P4) to 60.4% (P1) of the total chunks, and secondary chunks accounted for 21.9 (P1) and 25.1% (P4). Putting the primary and secondary chunks together, the vast majority of the chunks ($\geq 78.4\%$) were very short. The proportions of the tertiary chunks were much smaller, ranging from 6.9 (P3) to 12.2% (P4). The longer chunks accounted for 7.9 (P1) to 11.6% (P3).

The primary shifts of transitions within double-fixation chunks were loops (distance = 0) across pages. These accounted for 48.5 (P1) to 62.2% (P2). The pages agreed also on the secondary ($\sqrt{1}$) and tertiary ($\sqrt{2}$) distances, which involved adjacent segments connected laterally (or vertically) and diagonally, respectively. The proportion of the former ranged from 30.0 (P2) to 40.8% (P1). In contrast, that of the latter was much smaller ($\leq 8.8\%$). Put together, the overwhelming majority of the double-fixation chunks ($\geq 88.4\%$) were homogenous, that is, loops, or minimally heterogeneous ($\sqrt{1}$).

Loops and one-block shifts were also dominant among the chunks of length three or more. Loops accounted for 49.2 (P4) to 60.8% (P3) of the shifts, and one-block shifts accounted for 34.4 (P3) to 42.7% (P1) of them. Putting these together, the overwhelming majority ($\geq 91.0\%$) of the shifts within longer chunks were extremely short in distance.

Similarly, extremely short shifts ($\leq \sqrt{1}$) were modal among between-chunk transitions in reverse order and less prominent than within-chunk transitions. Primary one-block shifts accounted for 33.9 (P2) to 37.6% (P3) of the total between-chunk shifts, and loops accounted for 22.3 (P1) to 30.3% (P2). Their combined proportions ranged from 57.5 (P1) to 64.2% (P2).

TABLE 3: Number of patterns (*N*) by length (len) by ms.

	len	ms12		ms14		ms16	
		<i>N</i>	Loops	<i>N</i>	Loops	<i>N</i>	Loops
P1	1	18	B1/1	11		5	
	2	19		4		1	
P2	1	15	A1/1, B1/1, D1/1	9	B1/1	5	B1/1
	2	14	B1/2	4	B1/1	1	
	3	1					
P3	1	14	A1/1	12	A1/1	7	A1/1
	2	34	A1/8	11	A1/3	2	A1/1
	3	5	A1/2				
P4	1	18	A1/1	15	A1/1	6	
	2	20	A1/5	3		1	

Note. The length of a pattern (len) is the number of constituent chunks. Also listed are the identified within-chunk loops with the number of patterns in which they appeared.

The low prominence of the first two modal shifts was compensated by the relatively large proportions of the longer ones. Each of the two-block shifts ($\sqrt{2}$ and $\sqrt{4}$) exceeded 10% levels on all pages with the exception of 7.7% ($\sqrt{2}$) on P2. Compared to the paucity of shifts of three blocks ($\sqrt{5}$) within chunks ($\leq 1.7\%$), the corresponding distance between chunks, which ranged from 7.1 (P2) to 9.4% (P1), was noteworthy. Similarly noticeable was the size of the long-distance shifts ($\geq \sqrt{8}$), which ranged from 7.3 (P3) to 10.2% (P1), while such shifts were nonexistent or negligible (0.8% on P1) within chunks.

3.2. Examination of the Frequent Patterns. The frequent patterns extracted at three different ms levels (ms12, ms14, and ms16) are inclusive within each page in the sense that (a) subpatterns of a frequent pattern are also frequent at a given level and (b) the patterns extracted at a higher level are included in those at a lower level. For the sake of simplicity, the term “frequent” will be omitted below when obvious. Prior to mining, special coding was applied to the within-chunk loops as explained in Section 2.

As seen in Table 3, the patterns were generally short, consisting of one or two chunks across pages at all ms levels. The longer ones (one on P2 and five on P3), all of length three, were found only at ms12. The constituent chunks were simple in composition, being a single fixation or a single loop. The loops were limited to (A1A1), (B1B1), and (D1D1), all located in the first column of the mesh. (These will be denoted as (A1..), (B1..), and (D1..).) The (D1..) loop appeared only on P2 at ms12 by itself, unaccompanied by any other chunk. (B1..) appeared alone on P1 at ms12 and on P2 at all ms levels. Also, it was paired with B3 on P2 both as a prefix (ms12) and as a postfix (ms12 and ms14). (A1..) appeared by itself on P2 (ms12), P3 (ms12, ms14, and ms16), and P4 (ms12 and ms14) and also as a prefix to other chunk(s) on P3 (ms12, ms14, and ms16) and P4 (ms12). None of the corresponding segments were in the first column. The postfixes on P3 were A2 (ms12, ms14, and ms16); A3 and B3 (ms12 and ms14); and B2, B3, B4,

B5, C3, and D2 (ms12) in addition to A2A2 and B3B3 (ms12). Those on P4 were B4, C3, C4, D3, and D4 (ms12).

In the six patterns of length three found on P2 and P3 at ms12, the constituent codes were partially or totally homogenous. Five of them contained two repeated codes, either A2 or B3, including those prefixed by (A1..) as reported above. The remaining one, found solely on P3, contained A2. In the following examination of the double-chunk patterns, loops will be treated as single codes to reduce complexity.

The double-chunk patterns are listed in Table 4 by the direction of the sequences—upward, homogenous, horizontal, and downward. Superscripts L and R denote leftward and rightward sequences. Underscored patterns were extracted at ms14 and above. Those found only at ms16 are further emphasized in italicized bold face. The total number of patterns varied from 13 (on P2 and P4) to 34 (P3).

At ms16, the patterns were homogenous (B2B2 on P1; B3B3 on P2), horizontal (A1A2 on P3; C3C4 on P4), or downward (A2B3 on P3) sequences with the exception of down-rightward pattern A2B3 on P3. There was no leftward heterogeneous pattern.

The new patterns found at ms14 included an upward sequence (B2A3 on P3) and five downward sequences (B3C2 on P2; A1B3, A2B4, and A2C3 on P3; and A2B4 on P4) in addition to four homogenous sequences (B3B3 on P1; A2A2, B3B3 on P3; C3C3 on P4) and six horizontal sequences (A1A4 and B2B3 on P1; B3B1 on P2; and A1A3, A2A3, and B3B4 on P3). Among the 12 heterogeneous patterns, only two (B3B1 and B3C2 on P2) were leftward.

The patterns extracted at ms14 and above had no segments in rows D and E and no segments in the fifth column. None of the seven upward and downward sequences were strictly vertical, involving adjacent or nonadjacent columns in the ratio of 4 to 3. These vertical patterns mostly involved adjacent rows (6 out of 7).

Some of the constituent segments of the sequences at ms14 and above appeared solely as prefixes (A1 on P1 and P3; A2 on P4) or as postfixes (B3 on P1; B1 and C2 on P2; A3, B3, B4, and C3 on P3; B4 and C4 on P4).

The new double-chunk patterns found at ms12 had (a) segments in row D and in column 5, (b) notable positions of the new segments, (c) increased heterogeneous patterns, (d) increased sequences between nonadjacent rows, (e) strictly vertical sequences, and (f) bilateral sequence pairs. The segments in row D appeared only as postfixes in the downward sequences (D2 and D3 on P1 and P2; D2, D3, and D5 on P3; and D3 on P4). Similarly, the new segments found in row C were postfixes (C3 and C4 on P1; C3 on P2; C1 on P3; and C2 on P4) with a single exception (C2 on P3). The new segments in row B were mostly postfixes: B1, B4, and B5 on P1, B5 on P3, and B4 and B5 on P4. B2 and B3 on P4 were prefixes. An interesting case was B2 on P2 which was special, being a prefix to itself (B2B2). Dual roles were more notable than unary ones among the new segments in row A (A2 and A3 on P1, A1 on P2, and A3 on P4).

A total of seven new upward sequences were found, three on P1 and two on both P2 and P3, but still none on P4. These were prefixed by B2 (on P1 and P3), B3 (P2), or C3 (P3) and postfixes by the segments in row A—A1, A2, A3, or A4. Only

TABLE 4: Double-chunk patterns by direction.

Page	Direction	Pattern
P1	↑	B2A2 B2A3 ^R B2A4 ^R
	==	<u>B2B2</u> B3B3
	↔	<u>A1A4^R</u> <u>B2B3^R</u>
		A2A4 ^R A3A4 ^R B2B1 ^L B2B5 ^R B3B2 ^L B3B4 ^R
P2	↓	A1D3 ^R B2C3 ^R B2C4 ^R B2D2 B2D3 ^R B3D3
	↑	B3A1 ^L B3A3
	==	<u>B3B3</u> A1A1 B1B1 B2B2
	↔	<u>B3B1^L</u> B1B3 ^R B3B2 ^L
P3	↓	<u>B3C2^L</u> B3C3 B3D2 ^L B3D3
	↑	<u>B2A3^R</u> B2A2 C3A2 ^L
	==	<u>A2A2</u> <u>B3B3</u> B2B2 C3C3
	↔	<u>A1A2^R</u> <u>A1A3^R</u> <u>A2A3^R</u> <u>B3B4^R</u> A3A2 ^L B2B3 ^R B3B2 ^L B3B5 ^R C2C1 ^L
P4	↓	<u>A1B3^R</u> <u>A2B3^R</u> <u>A2B4^R</u> <u>A2C3^R</u> A1B2 ^R A1B4 ^R A1B5 ^R A1C3 ^R A1D2 ^R A2B2 A2B5 ^R A2D2 A2D3 ^R A2D5 ^R B2C3 ^R B3D2 ^L B3D3 C3D2 ^L
	↑	(none)
	==	<u>C3C3</u>
	↔	<u>C3C4^R</u> A2A3 ^R B3B4 ^R B3B5 ^R C3C2 ^L
P4	↓	<u>A2B4</u> A2C4 ^R A3B4 ^R B2C3 ^R B3C3 B3C4 ^R C3D3

Note. The sequence directions are upward (↑), homogenous (==), horizontal (↔), and downward (↓). Underscored patterns were extracted at ms14. Those extracted at 16 are also emphasized in italicized bold face. Leftward and rightward sequences are marked by superscripts ^L and ^R, respectively.

C3A2 involved nonadjacent rows. A strictly vertical sequence was present on each of P1, P2, and P3—B2A2, B3A3, and B2A2. The rest were rightward (B2A3 and B2A4 on P1) or leftward on P2 and P3 (B3A1 on P2; C3A2 on P3).

A total of five new homogenous sequences were found on P2 and P3, one in row A (A1A1 on P2), three in row B (B1B1 on P2 and B2B2 on P2 and P3), and one in row C (C3C3 on P3). Like those at ms14 and above, none of the constituents were in columns 4 or 5.

A total of 17 new horizontal sequences were found on P1 (two in row A and four in row B), P2 (two in B), P3 (one in A, three in B, and one in C), and P4 (one in A, two in B, and one in C). A2 and A3 appeared as a prefix or as a postfix, while A4 appeared only as a postfix. The same held for B1, B2, and B3, while B4 and B5 appeared only as postfixes. C2 assumed dual positions in C2C1 on P3 and C3C2 on P4, both of which were leftward. The ratio of leftward to rightward sequences was 2 : 4, 1 : 1, 3 : 2, and 1 : 3 in the order of P1, P2, P3, and P4.

A total of 29 new downward sequences were found, six on P1, three on P2, 14 on P3, and six on P4. The prefixes concentrated in rows A and B with two exceptions (C3D2 on P3 and C3D3 on P4). In contrast, the postfixes concentrated in rows C and D with exceptions of five patterns on P3 and one on P4. Half or more of the downward patterns on P1, P2, and

TABLE 5: Isolate primitives by ms level.

	msl2	msl4	msl6
P1	A5 C2 C5 E3	A2 A3 B1 B5 C3 C4 D3	A1 A2 A4 B3
P2	A2 B4 C1 D1	A1 A3 B2 C3 D3	A1 A3
P3	(none)	B5 C1 C2 D2	A3 B2 C3
P4	A5 B5 C1 C5 D2	B1 B2 B3 B5 C1 C2 C5 D3 D4	A2 B3 B4 C5

Note. Primitives in bold face were persistent at two or three ms levels.

P3 involved nonadjacent rows (A-D/1 and B-D/3 on P1; B-D/2 on P2; and A-C/1, A-D/5, and B-D/2 on P3, where n denotes the number of cases), whereas only A2C4 out of six patterns did so on P4. The strictly vertical patterns were limited to columns 2 and 3 (B-D/2 on P1; B-C/1 and B-D/1 on P2; A-B/1, A-D/1, and B-D/1 on P3; and B-C/1 and C-D/ on P4). The rest were rightward on P1 and P4, leftward on P2, or mixed on P3.

Among all of the patterns in Table 4, the heterogeneous sequences were mostly unilateral in that the symmetric pairs were limited in number (B2B3-B3B2 on P1; B1B3-B3B1 on P2; A2A3-A3A2, A2B2-B2A2, A2C3-C3A2, and B2B3-B3B2 on P3; and none on P4). Four of these were horizontal sequences. The constituents were limited to a subset consisting of the first three rows and columns, that is, {A2, B1, B2, B3, and C3}.

The individual constituents of the multichunk patterns were frequent by themselves as primitive patterns at a given ms level, but not vice versa. Table 5 lists the isolate primitive patterns not participating in any multichunk pattern at a given ms level. While the number of total primitive patterns monotonically decreased from msl2 to msl6, the ratio of the isolate primitive patterns to the total primitive patterns monotonically increased on all pages almost perfectly. The ratios at {msl2, 14, 16} were {4/17, 7/11, 4/5}, {4/13, 5/8, 2/4}, {0/13, 4/11, 3/6}, and {5/17, 9/14, 4/6}, in the order of P1, P2, P3, and P4. The sole exception was the second and the third ratios on P2. There were no isolates on P3 at msl2.

Generally, an isolate primitive at a given ms level would become a member of sequence(s) at a lower level and would not be present at a higher level. Exceptionally, C5, located in the rightmost column, persisted on P4 as an isolate at all ms levels. Partial persistence was observed between msl4 and msl6 on P1 (A2), P2 (A1, A3), and P4 (B3) as well as between msl2 and msl4 on P4 (B5, C1). No persistence was observed on P3. The persistent ones on P1 and P4 were limited to the first three columns of the top row, {A1, A2, A3}, whereas those on P4 spread over rows B and C in columns 1, 3, and 5, that is, {B3, B5, C1, C5}.

Finally, E3 on P1 at msl2 was the sole frequent segment in the bottom row E where segments were generally infrequent across pages at all ms levels.

4. Discussion

Eye-tracking researchers have inferred a fixation from gaze points closely clustered in space and time, treating it as a

meaningful unit of information processing, that is, a *chunk*, a familiar concept in psychology. Chunking of lower-level chunks into a higher one is not uncommon as seen in the relationships (letter, word, phrase, sentence, paragraph, ...). The present paper examined the patterns of second-order chunks, that is, chunks of fixations, using isolate gaze point(s) not participating in any fixation as the delimiter. The delimiter was assumed to play an auxiliary role in chunking, like a pause in speech.

Most of the identified chunks were short, consisting of one or two fixations. Also, the transitions within multifixation chunks and between chunks were mostly short in distance, either loops or one-block shifts to adjacent segments. These seem to be attributable to the minimum criterion of the delimiter we employed—at least one isolate gaze point. Hence, even an accidental dislocation of one's gaze resulted in chunking. It would be ideal if we could separate cognitively meaningful chunking from accidental chunking. Until an effective method is established, the best we can do is to be cautious in interpreting the results.

Actually, setting an appropriate criterion is a difficult task due to the possible individual and situational variations. Perhaps individuated criteria will be appropriate instead of a uniform criterion. Further investigation of the distributions of gaze points participating in fixations and those that are isolated is necessary.

As reported earlier, within- and between-chunk transitions were similar in that the first two modal distances were zero (i.e., loops) and one block. However, these differed in order and in magnitude. Loops were primary among within-chunk transitions but secondary among between-chunk transitions. The opposite was true for the one-block shifts. Next, the proportions of the primary and secondary distances of the within-chunk transitions exceeded the respective proportions pertaining to the between-chunk transitions. Similarly, there were more long-distance shifts between chunks than within them.

These results seem to suggest that the attention of our subjects was most likely shifted, after a pause, to an adjacent segment one block away or within the same segment. The medium or long-distance shifts were also separated by pauses, though their proportions were smaller than the short ones. Shifts without a pause, that is, within-chunk shifts, were short, chiefly occurring in the same segment or between adjacent segments one block away.

Now we turn to a discussion of the frequent patterns (i.e., subsequences) extracted by PrefixSpan. The patterns were simple in structure, mostly consisting of single or double chunks. Furthermore, the chunks themselves contained single fixations or single loops as expected from the chunk properties discussed above. More complex structures might have resulted if we had employed less stringent criteria for the delimiter. Even so, beneath the structural simplicity, interesting properties emerged as to the segment differentiation and the directional unevenness in attentional shifts.

First, the within-chunk loops were limited to (A1..), (B1..), and (D1..), all of which were in the leftmost column. While the presence of (D1..) was quite limited, the leading roles of (A1..) and (B1..) as prefixes in the multichunk sequences are

noteworthy. These roles might be attributable to menu items placed in the segments. Second, the multichunk sequences chiefly consisted of the segments in rows A, B, and C. In particular, the leading role of A1 on P1 and P3 was noteworthy, like the loop (A1..), though its dual role as pre- and postfix was observed on P2. In contrast, A4, B4, and C4 were consistently positioned as postfixes. The same held for the segments in row D, which appeared only at the lowest ms level. The segments in row E were totally absent in multichunk sequences.

Third, the sequences at ms14 and ms16 were more likely to be horizontal, including homogenous codes, than downward and, to much less extent, than the upward sequence, which remained least likely among the additional patterns found at ms12. The order between horizontal and downward sequences varied across pages at ms12.

By chunking eye-tracking records into smaller units, we discovered interesting properties of the eye movement of web page viewers. However, further studies seem necessary to enhance the present approach, for example, by setting up nested AOI's to reflect the hierarchical structure of the web objects [16] and by adjusting the chunk delimiters to accommodate individual and task variations. Besides these refinements, we are planning an application of mined frequent patterns to simultaneous clustering [17] of subjects and the properties of their eye movement and other relevant indices.

Appendix

We briefly explain frequent sequential pattern mining by PrefixSpan (prefix-projected sequential pattern mining) developed by Pei et al. [13, 14]. Interested readers should consult the original articles for formal descriptions and evaluations in comparison with other competing algorithms.

Let us use Table 1 as the DB (database) to be scanned. It consists of four sequences whose elements are nonempty subsets of items $\{a, b, c, d, e, f, g\}$. An element is composed of a set of items: a, b, c, d, e, f , and g . PrefixSpan assumes that items in an element are alphabetically ordered with no duplication, for example, $[a]$, $[ac]$, and $[df]$.

The goal of PrefixSpan is to find subsequences frequently shared among the records in DB. A subsequence is defined as the list of nonempty subsets of the elements of a given sequence, where the sequential order of elements is preserved. For example, $[bc][c][d][f]$ is a subsequence of $[a][abc][ac][d][cf]$. The threshold of frequent occurrence is called the minimum support (abbreviated as ms in this paper). Its value is to be specified by the user.

Subsequences of special importance are a prefix and the associated suffix. For instance, a frequent item a , with $ms = 3$, can serve as a prefix of the ensuing pattern (i.e., the suffix) to be scanned next. The patterns listed in the second column of Table 1 are the suffix sequences constituting the $\langle a \rangle$ -projected database. Similar databases are to be constructed for every frequent item. With ms2, ab and $\underline{b}$ together will be considered frequent, where the underline $\underline{}$ implies a . Hence, ab will serve as a prefix, yielding only the two suffixes $[\underline{c}][ac][d][cf]$ and $[df][c][b]$.

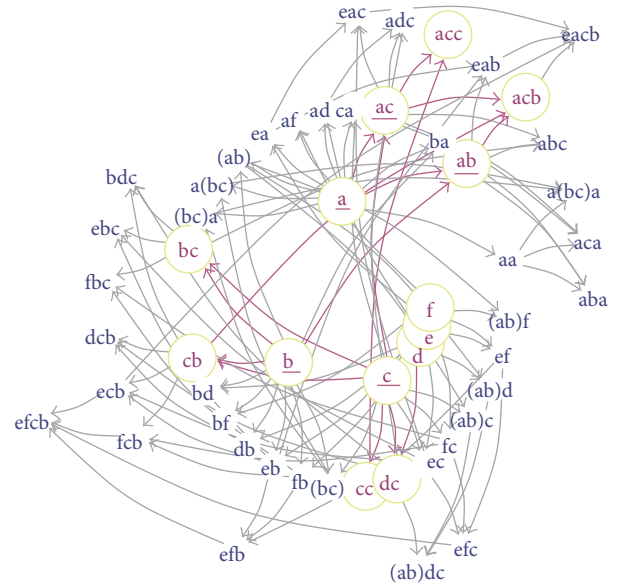


FIGURE 3: Network of the frequent patterns extracted at ms2 (small letters in dark blue), ms3 (large letters in dark red), and ms4 (underlined). Note: $[]$ is omitted for the single-code chunks and is replaced by $()$ for the multiple-code chunks for the sake of simplicity. See the first column of Table 1 for the initial sequences.

The network of the frequent patterns extracted at $ms = 2, 3$, and 4 is illustrated in Figure 3 to help grasp the inclusive relations among them in two senses: (a) nn element of a frequent pattern is also frequent; and (b) a frequent pattern at a given ms level is also frequent at a lower level.

More formally, a sequence β of length m is a prefix of another sequence α of length n ($m \leq n$) consisting of frequent elements in the database if and only if the first $m - 1$ elements are identical; the last element of β is a subset of the m th element of α .

The suffix of α with regard to β is a sequence, the first element of which is the difference between the m th elements of α and β . The remaining elements of the suffix are identical with the $(m + 1)$ th to the last element of α ; that is,

$$\text{element}_{1|\text{suffix}} = \text{element}_{m|\alpha} - \text{element}_{m|\beta}; \quad (\text{A.1})$$

if $m < n$

$$\begin{aligned} \text{element}_{j|\text{suffix}} &= \text{element}_{m+j-1|\alpha}, \\ j &= 2, \dots, n - m + 1. \end{aligned} \quad (\text{A.2})$$

Scanning with respect to the prefix stops when the suffix becomes nil ($m = n$) or no frequent item exists in the projected database. This process is executed in a depth-first manner for every code initially identified as frequent.

It must be noted that some of the extracted patterns may be hard to identify in the original sequences, due to the intermittent removal of infrequent items from the projected database during the process, for example, the extracted pattern $[a][c][b]$ in Table 2 and the sequence $[ef][ab][df][c][b]$ in Table 1. This point should be clear to those who are familiar

with masking (or wildcard) characters, such as an asterisk “*” in string matching. One can find original patterns by attaching a masking character to the extracted patterns.

Conflict of Interests

The authors declare that there is no conflict of interests regarding the publication of this paper.

References

- [1] S. Martinez-Conde, S. L. Macknik, and D. H. Hubel, “The role of fixational eye movements in visual perception,” *Nature Reviews Neuroscience*, vol. 5, no. 3, pp. 229–240, 2004.
- [2] D. D. Salvucci and J. H. Goldberg, “Identifying fixations and saccades in eye-tracking protocols,” in *Proceedings of the Eye Tracking Research and Applications Symposium*, pp. 71–78, November 2000.
- [3] G. A. Miller, “The magical number seven, plus or minus two: some limits on our capacity for processing information,” *Psychological Review*, vol. 63, no. 2, pp. 81–97, 1956.
- [4] D. Melcher, “Dynamic, object-based remapping of visual features in trans-saccadic perception,” *Journal of Vision*, vol. 8, no. 14, article 2, 2008.
- [5] D. Melcher, “Selective attention and the active remapping of object features in trans-saccadic perception,” *Vision Research*, vol. 49, no. 10, pp. 1249–1255, 2009.
- [6] J. Ross, M. C. Morrone, M. E. Goldberg, and D. C. Burr, “Changes in visual perception at the time of saccades,” *Trends in Neurosciences*, vol. 24, no. 2, pp. 113–121, 2001.
- [7] E. Cutrell and Z. Guan, “What are you looking for?: an eye-tracking study of information usage in Web search,” in *Proceedings of the 25th SIGCHI Conference on Human Factors in Computing Systems (CHI '07)*, pp. 407–416, May 2007.
- [8] N. Matsuda and H. Takeuchi, “Networks emerging from shifts of interest in eye-tracking records,” *eMinds*, vol. 2, no. 7, pp. 3–16, 2011.
- [9] N. Matsuda and H. Takeuchi, “Joint analysis of static and dynamic importance in the eye-tracking records of web page readers,” *Journal of Eye Movement Research*, vol. 4, no. 1, article 5, 12 pages, 2011.
- [10] N. Matsuda and H. Takeuchi, “Do heavy and light users differ in the Web-page viewing patterns? Analysis of their eye-tracking records by heat maps and networks of transitions,” *International Journal of Computer Information Systems and Industrial Management Applications*, vol. 4, pp. 109–120, 2012.
- [11] J. H. Goldberg and X. P. Kotval, “Computer interface evaluation using eye movements: methods and constructs,” *International Journal of Industrial Ergonomics*, vol. 24, no. 6, pp. 631–645, 1999.
- [12] J. D. McCarthy, M. A. Sasse, and J. Riegelsberger, “The geometry of web search,” in *People and Computers XVIII—Design for Life*, pp. 249–262, Springer, London, UK, 2004.
- [13] J. Pei, J. Han, B. Mortazavi-Asl et al., “PrefixSpan: mining sequential patterns efficiently by prefix-projected pattern growth,” in *Proceedings of the 17th International Conference on Data Engineering*, pp. 215–224, April 2001.
- [14] J. Pei, J. Han, B. Mortazavi-Asl et al., “Mining sequential patterns by pattern-growth: The PrefixSpan approach,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 16, no. 11, pp. 1424–1440, 2004.
- [15] R. Agrawal and R. Srikant, “Fast algorithms for mining association rules,” in *Proceedings of the International Conference on Very Large Data Bases (VLDB '94)*, pp. 487–499, 1994.
- [16] D. I. Brooks, I. P. Rasmussen, and A. Hollingworth, “The nesting of search contexts within natural scenes: evidence from contextual cuing,” *Journal of Experimental Psychology: Human Perception and Performance*, vol. 36, no. 6, pp. 1406–1418, 2010.
- [17] A. Prelić, S. Bleuler, P. Zimmermann et al., “A systematic comparison and evaluation of biclustering methods for gene expression data,” *Bioinformatics*, vol. 22, no. 9, pp. 1122–1129, 2006.