

Discovering NBA Game Stories from Twitter

Zhichao ZHANG[†], Hisashi KOGA[†], and Youhei OGYU[†]

[†] Graduate School of Information Systems, University of Electro-Communications 1-5-1, Chofugaoka,
Chofu-shi, Tokyo, 182-8585 Japan

E-mail: †{zhang,koga,ogyu}@sd.is.uec.ac.jp

Abstract This paper proposes a system which generates the game summary for basketball games by quoting the representative tweets during the game period for the first time, whereas the previous similar researches treated soccer and American football. Basketball is a more challenging subject than football or soccer for making a summary. This paper introduces a unique idea to put the discussions at the break times between quarters and at the end of the game into the game summary, since spectators state the game status at that moment and comments upon the activities of the players and the teams then. Without specifying keywords, our heuristics attempt to such discussions as the peaks of long duration in the tweet volume graph. Experimentally, our summary covers up to 87% of the items written in the NBA Official Game Summary.

Key words Twitter Mining, Sport Event, Summary Generation

1. Introduction

Twitter has become the most popular micro-blog which has more than 0.5 billion users. A lot of tweets are published every day in Twitter, containing a myriad of information about what the users are doing and watching, which can be seen as describing opinions about various events. Recently, mining information about such events from the twitter stream has been a significant research topic. This research approach is categorized into two kinds: the first kind attempts to discover the event occurrence without knowing event types [2], while the second one intends to obtain the detailed descriptions about a specific event like an earthquake, a typhoon etc. [1] from the twitter stream and to summarize them.

In the line of the second kind, some previous litterateurs dealt with sport games. Chakrabarti et al. [3] and Sport-Sense [4] studied American football, while [5] examined the soccer's world cup. All these methods identify the important moments in sport games by picking up the time instances when the tweet volume per time unit rapidly increases. For the sport games, a remarkable event causes a sudden increase of the tweet volume, because many Twitter users comment on it. Basically, important events in sport games can be discovered by checking if the tweet volume suddenly increases.

Our research purposes to generate the summary for basketball games from the tweets regarding to the NBA (National Basketball Association) games. Basketball has different properties from American football and soccer. For example, the goal frequency of basketball game is much higher

than that of soccer game. The score of a basketball game can be 110-100, while the score of a soccer game can be only 1-1. Therefore, one shoot/goal is less important for basketball than for soccer. This paper introduces a unique idea to generate the game summary effectively for basketball games, while paying attention to the increase of the tweet volume in the same way as the previous researches. After discovering the important moments from the tweet volume graph which records the number of tweets per time unit, our method yields the game story by choosing the representative tweets for each important moment. Notably, our method works in an almost unsupervised way without requiring the event keywords specific to basketball games such as dunk, 3-pointer and so on. Only the two team names associated with the game have to be given in order to extract the relevant set of tweets from the twitter stream.

The rest of the paper is organized as follows: In Sect. 2, we brief the related works which focus on the event detection from the twitter stream. Section 3 analyzes the characteristics of the tweets regarding to the NBA games. Section 4 describes our summary generation system for the NBA games. Section 5 reports the experimental results. Finally, Section 6 concludes this paper and discusses the future work.

2. Related Works

This section briefly refers to the previous researches which attempted to detect events for sport games from the twitter stream.

Zhao et al. [4] developed a system named Sportsense which displays the major events rates fans' excitement level in the

middle of the NFL American football games. They first build event templates by learning event examples which shows the change of the tweet volume when the events happen. Once the event templates are complete, the same type of events can be detected on-line for the ongoing game by matching to the event templates. However, this event detection method requires supervision, since the keywords related to events, for example "touchdown" or "TD" for short, must be specified so as to collect the event examples. The researchers must have the domain knowledge about the NFL games in order to predetermine the event keywords. Chakrabarti et al. [3] use Twitter to generate summaries of long running, structure rich events under the circumstances that multiple event instances share the same underlying structure. Specifically, they learned the structure and the vocabulary of events for American football with a modified Hidden Markov Model (HMM). Here the tweets for many games need to be prepared to compose the learning data. Moreover, the learning process for the HMM is time-consuming.

On the other hand, Nichols et al. [5] detected events and generated a journalistic summary from the tweets at a World Cup soccer game. They neither count on predetermined event keywords nor learn from the tweets for multiple games. Namely, their method works in an unsupervised way. We explain their method in details below, since we will extend it for basketball games in this paper.

After obtaining the tweets regarding to a given game using basic keyword filtering via the twitter API, they first draw a tweet volume graph whose x-axis presents the time in minute and the y-axis denotes the number of tweets per minute. From this graph, they extract spikes each of which is defined by the triple $\langle \text{Start Time}, \text{Peak Time}, \text{End Time} \rangle$ as shown in Fig. 1. The tweet volume starts increasing at the start time, reaches the peak at the peak time and stops decreasing at the end time. Among the derived spikes, only those whose slope between the start time and the peak time goes beyond a threshold are memorized as the important moments $M_1, M_2, \dots, M_m$, where m denotes the number of important moments. In [5], the threshold is empirically set to 3 times as large as the median of all the slopes for American football. Similar approaches are utilized by some other previous researches [6] [7] [8].

After the important moments in the game are identified, the game summary is constructed by selecting the N representative sentences from the set of tweets posted at each important moment M_j ($1 \leq j \leq m$). In this process, first, the longest sentence of each tweet included in M_j is abstracted. Let the set of such longest sentences for M_j be L_j . Then, the score of a sentence in L_j is computed by summing up the scores of its word tokens. Here, the score of a word token

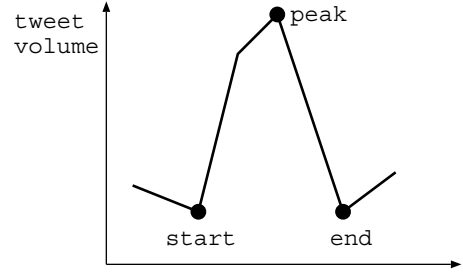


Figure 1 the Start, Peak, and End Times of a Spike

equals its word frequency. Finally, the top N sentences that do not share any non-stop word stemmed tokens are output as the game summary. For American football game, [5] sets N to 3, which means that only a few sentences are enough to cover the contents of an important moment.

3. Tweets for Basketball Games

This section first explains how to collect the tweet dataset regarding to the NBA games and then describe their features.

3.1 Collection of Tweets

We rely on the Twitter Streaming API (<https://dev.twitter.com/docs/streaming-apis>) to gather the tweets on the NBA games. This service allows developers to pull tweets in real-time which contain specified keywords. In our case, we collect the tweets for a certain NBA game between two teams by setting their team names as the keywords. For example, we can collect the tweets for the game between the Miami Heat and San Antonio Spurs, by specifying the set of hashtags: "#Heat", "#Spurs", "#MIA" and "#SAS" as the keywords. Here "MIA" and "SAS" are abbreviation of the two team names. We get the team names and their abbreviations from the section of "teams" from the NBA official website (www.nba.com). In order to get the tweets at the game period, we keep only the tweets with the timestamp between the begin time and the end time of the game, which we also learn from the official website. Our dataset consists of tweets for 30 games from the regular season, playoffs and the finals of NBA 2012-2013 season. We show two examples of the tweets below.

@TrappedInThe225 - Down by 6 still in the game. #Heat null Sat May 25 10:01:36 JST 2013

@dudeimspacely - #Pacers 28#Heat 22 End of 1st Quarter. #NBAPlayoffs null Sat May 25 10:01:36 JST 2013

3.2 Feature of the Tweets for NBA Games

Figure 2 illustrates the tweet volume graph which records the number of tweets per minute for the game between Lakers and Spurs in Apr. 27th, 2013. After examining several tweet volume graphs for different NBA games manually, we have noticed the next primary features of the tweet dataset for the NBA games.

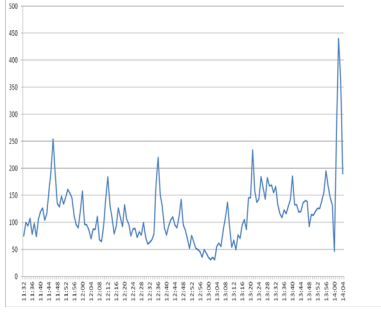


Figure 2 Tweet Volume Graph for a Basketball Game

(1) The graph fluctuates more often in the NBA games than in the games of NFL and soccer, which implies that it contains more spikes. This is probably because the basketball games are accompanied by more continuous actions than American football and soccer. We also noticed that spikes with sharp slopes tend to be related to real-time actions such as splendid slam dunks and turnovers.

(2) The contents of the tweets contained in one spike are very diverse for the NBA games, since multiple events such as slam dunk, assist, three pointer, game winner etc. can happen in just one minute.

Our system in the next section exploits the two above features to generate an attractive game story for an NBA game.

4. Our Event Detection System

The purpose of this system is discovering useful information as much as possible, and displaying them to the NBA fans as the games story. So the input of our system is the stream of tweets regarding to a certain NBA game, and the output is the summary, i.e. the story of this game. We develop this system by tailoring the method of Nichols et al. [5] for the basketball games. It operates the next three steps in order so as to output the final games story:

(1) Our system first determines the important moments by choosing some of the spikes in the tweet volume graph.

(2) Next, the tweets that can describe the contents of the important moments are chosen from the set of tweets published at the important moments.

(3) The final game summary is derived by excluding the similar tweets from the tweets chosen at the previous step. We explain the three steps from now on.

4.1 Decision of Important Moments

Like [5], the important moments are derived by searching peculiar spikes from the tweet volume graph. Although [5] seek spikes whose slopes are steeper than a certain threshold as the important moments for soccer, considering spikes with sharp slopes is insufficient for basketball, because they are related to the real-time actions such as beautiful slam dunks as pointed out in Sect. 3.2, Thus, the general game status

information is missed.

To contain such game status information, this paper uniquely pays attention to the discussions at the break times between quarters and at the end of the game, which always contains important information such as the game status and the comments on the previous quarter. For instance, during the halftime of the NBA games, spectators give their impression of the first half and, therefore, we should not neglect the tweets published then so as to obtain useful information. Note that the tweets in the break times are ignored or slighted in the previous researches for soccer and football. We show an example of the tweet issued at a break time below. This tweet was posted at the break time after the 3rd quarter of a game between the Indiana Pacers and the Miami Heat. It surely exhibits the game status at that moment and is significant.

At the end of the 3rd, the Pacers are ahead of the Heat by 13. Hibbert, George lead with 22 points each.

Unfortunately, the discussions at the break times cause only gentle excitement and result in the peaks with gentle slopes in the tweet volume graph. Then, how can we find such discussions from the graph without knowing specific keywords? In this paper, we propose a simple efficient heuristic approach to regard the spikes having long duration as the discussions between the break times. The rationale of this idea is as follows. At the beginning of a break time, people start talking about the previous quarter and a uphill slope is formed. Then, since they gradually leave from their PCs or smart phones, a downhill slope is formed. Interestingly, this downhill slope becomes a very long tail, as it is never interrupted until the next quarter starts.

We take both the slope and the duration of a spike into consideration. Particularly, we evaluate the value of a spike P according to Eq. (1). A spike is evaluated higher, as this formula becomes larger for the spike.

$$Score(P) = f_s \times \frac{Slope(P)}{MaxSlope} + f_a \times \frac{Area(P)}{MaxArea}. \quad (1)$$

f_s and f_a are weighting parameters to control the contribution of slope and area. We currently set both f_s and f_a to 0.5. MaxSlope denotes the biggest slope and MaxArea is the biggest area size of the peaks over the whole tweet volume graph. Though the area size of a spike appears in Eq.(1) instead of the spike duration, be aware that the area size of the spike is roughly proportional to its duration. In the same way as [5], we choose the peaks with the score higher than a threshold θ as the important moments. In the experiments at Sect. 5, θ is set to 0.1.

4.2 Selection of Tweets for Important Moments

Next, for an important moment, we select the set of tweets which can describe its contents. In this process, a tweet is ranked according to the relevance of words that the tweet contains. Here, the relevance of a word is determined by its frequency in the group of tweets that belong to the important moment. In counting the word frequency, we exclude the English stop words by utilizing English stop-word dictionaries opened to the public on the web. In addition, the two team names which compete in the concerned NBA game are also discarded. Because the team names are used as the filtering keywords for the Twitter Streaming API, they are contained in almost all the tweets. Thus the team names cannot describe the contents of the important moments well.

After calculating the frequency of all the words over the tweets belonging to the important moment, we get the top K words which have the highest word frequency at the important moment. In our current implementation, $K = 10$.

Next, the score of a tweet is computed. We regard a tweet is more important as it contains more highly-ranked keywords. The score of a tweet t is denoted by $V(t)$ in Eq. (2).

$$V(t) = \sum_{i=1}^n \text{score}(i) \quad (2)$$

Here, n is the number of the top K words in t , and $\text{score}(i)$ presents the value of the top i -th word w_i . For $K = 10$, $\text{score}(i)$ is set to $20 - i$. Hence, the word with higher frequency has a higher score. In particular, the top keyword is assigned about twice as large score as the 10th keyword. Finally, the set of tweets t for which $V(t)$ becomes greater than some threshold τ are passed to subsequent processing discussed in Sect 4.3.

To confirm if we can successfully get the relevant tweets using the top 10 keywords, and if the diversity of these tweets is high, we preliminarily apply our method to several important moments. Here we report one case example for one important moment of the game between Houston Rockets and Oklahoma City Thunder on Apr. 30th, 2013. For this case, the top 10 keywords are as follows:

durant, kevin, left, dunk, seconds, pointer, lead, driving, harden, cuts.

The five tweets with the highest score for this important moment is shown below.

- (1) @okcthunder: Kevin Durant with a 3-pointer and a driving dunk in 29 seconds. #Thunder cuts #Rockets lead to 2. 105-103. 1:13 left in Gam
- (2) harden missed 3-pointer clutch shots in a row. watch durant pull up and win it with a 3 at the buzzer. #Thunder
- (3) James harden blowing the games for #ROCKETS

(4) Harden with the airball with 53 seconds left. Did I say soft? I meant flaccid. #Rockets

(5) James Harden picks up his 5th foul, he will sub out

This example shows that the diversity of the top 10 keywords is high. The top 10 keywords are divided into several types: (1) players' names, (2) event names and (3) general information words. Despite only the 10 words are considered in selecting tweets for this important moment, the contents of the chosen tweets have high diversity. We observe the similar tendency for other important moment examples. We guess the reason of this phenomenon as follows.

(1) On condition that one event is associated with one player, the top 10 keywords for one important moment includes several player's names for the most cases. Thus, multiple events are discovered with the top 10 keywords.

(2) One word out of the top 10 keywords may correspond to multiple events. For example, a player's name can be related to several different events.

(3) Even for the identical event, multiple tweets describing it can supply different information, while they contain the common keywords.

4.3 Removal of Similar Tweets

The similar tweets are excluded from the tweets chosen at the previous step. Removal of the similar tweets is necessary here, since a lot of spectators issue very similar tweets on the same event and displaying such similar tweets annoys the NBA fans.

This step first uses clustering to classifying the set of tweets into several clusters of similar tweets and then outputs one representative tweet per cluster. As a clustering algorithm, we use average linkage method, one of the well-known agglomerative hierarchical clustering algorithms. The agglomerative hierarchical clustering begins with one-point clusters and recursively merges the most similar pair of clusters, until the number of clusters finally reduces to one. In the agglomeration step, the clustering algorithm searches the closest pair of clusters and merges them into a new single cluster. The hierarchical clustering algorithm is advantageous in that the number of clusters do not have to be specified a priori. This nice feature is suitable for our case, since a single important moment contains multiple events for basketball games as stated in Sect. 3.2, so that it is impossible to grasp the proper number of tweet clusters beforehand.

In the average linkage method, the distance between two clusters is defined as the average distance between any member (tweet in our case) of one cluster to any member of the other cluster. Here, the distance $D(t_i, t_j)$ between a pair of tweets t_i and t_j is defined as the Jaccard distance in Eq. (3).

$$D(t_i, t_j) = 1 - \frac{|S_i \cap S_j|}{|S_i \cup S_j|}, \quad (3)$$

where S_i and S_j symbolize the set of the words in t_i and t_j respectively. The Jaccard distance is derived by subtracting the Jaccard coefficient from 1. The Jaccard coefficient between two sets A and B is defined as $\frac{A \cap B}{A \cup B}$ and measures the extent of the overlap between them.

If we stop merging clusters before the cluster number decreases to 1, multiple clusters are extracted. Our implementation ceases merging clusters, when the distance between the two clusters to be merged exceeds a threshold value $D = 0.925$. Since we determine this value of D only empirically, the algorithm to derive an optimal value of D remains to be developed in future.

After having multiple clusters in the above way, we determine one representative tweet for every cluster with more than 3 members. Namely, small clusters are not adopted, since they are not admitted by many spectators. Consider a cluster C consisting of m tweets ($m > 3$). A tweet in C which is the most similar to other tweets in the same cluster is appointed to the representative of C . Here, the similarity of a tweet t_i in C to other tweets is measured by $\sum_{j=1, j \neq i}^m \frac{|S_i \cap S_j|}{|S_i \cup S_j|}$ which sums up the Jaccard coefficients between t_i and all the other tweets t_j in C .

Finally, the representative tweets of all the clusters for all the important moments constitute our game story.

5. Experiments

With the dataset in Sect. 3, we experimentally evaluate the performance of our system.

5.1 Accuracy of Detected Important Moments

One of the novel ideas in our method is that it takes both the slope and the area of the spikes into account to detect the important moments. Therefore, we evaluate our algorithm by comparing it with another policy which only considers the slope only. This policy is abbreviated as SL hereafter. SL is derived by setting $f_s = 1.0$ and $f_a = 0.0$ in Eq. (1) The accuracy of an algorithm is measured by how many percentages of the real important moments, i.e., the spikes with useful information are identified. We manually derive the real important moments for all the 30 games from their tweet volume graphs as the ground truth.

As the result, 14.86 genuine important moments are discovered manually per game on average. On average, our algorithm finds 13.23 important moments contained in the ground truth, while SL finds 10.56. Thus, the recall of our algorithm reaches 89% whereas that of the SL is 71%. Thus, our algorithm can detect important moments more accurately than SL.

To be more comprehensive, we give an explanation with respect to one game instance between the Miami Heat and

the Indiana Pacers on June 2nd, 2013. First, the ground truth important moments discovered manually are shown on the tweet volume graph for this game in Fig. 3. There, the spikes surrounded by the green rectangle correspond to the ground truth important moments. To see what kind of information the ground truth important moments include, Table 1 lists a typical tweet example chosen by us for the 14 important moments. Among them, IM1, IM5, IM9 and IM14 correspond to the discussions at the break times between quarters and at the end of the game.

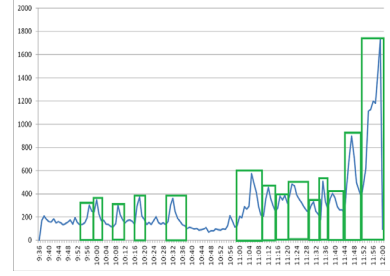


Figure 3 Ground Truth Detected Manually

Table 1 Examples of Tweets for each Important Moment

IM1	Paul George's dunk on bosh #nasty #heat vs pacers
IM2	At the end of the first quarter, the #Pacers trail the Heat 23-21
IM3	Ridiculous 3 point FGyo early in this game. #Heat are shooting lightsout at 85% (6-7) and #Pacers are shooting 50% (1-2). #Heat lead 25-24
IM4	2 missed dunks #pacers ouch
IM5	D-Wade steal, ends up with a #LeBron JAMI #Heat lead at the half 40-39 Over the #pacers
IM6	#pacers taking an 11 point lead against the #heats on the 3rd quarter #awesome #enjoyingthegame #keep-itgoing
IM7	The pacers are playing so good defensively. Let's go Heat!!! Let's go Heat!!! Letis go Heat!! #eastern finals#heat #game6
IM8	#Pacers are dominating the boards. They've got 43 rebs (13 Off 30 def) vs 28 rebs (10 Off 18 def) for #Heat. Lance Stephenson leads w/ 11
IM9	At the end of the 3rd, the #Pacers are ahead of the Heat by 13. Hibbert, George lead with 22 points each
IM10	Three pointer! lets go miller #heat
IM11	Intense #heat and #pacers!! win or go home game Let's go #Heat#NBA #HEATNATION
IM12	LeBron James just got a technical for throwing a tantrum and likely closing the door on a comeback. #heat #pacers
IM13	#Pacers go on a 9-0 to restore order and now lead 81-6 w/ 3:55 left. Roy Hibbert 24pts, 9rebs, Paul George 25pts 8 rebs. #heatvpacers
IM14	It's going to game 7 #Pacers

Our algorithm detects 13 ground truth important moments. It misses one true important moment, that is, IM10 that describes an important three pointer by the Miami Heat’s player Mike Miller. Note that all the discussions in the break times are found by our algorithm. On the other hand, SL discovers 10 important moments of which 9 match to the ground truth. It misses IM1, IM8, IM9 IM10 and IM12. Importantly, IM9 which describes the summary of the third quarter and corresponds to the break time is not recognized by SL considering the slope only. In this way, SL leaks more meaningful information than our algorithm.

5.2 Quality of Our Summary

The design purpose of our system is to describe the NBA game story by exploiting the useful information contained in the representative tweets. To examine if our system fulfills this goal, we evaluate the representative tweets outputted by our system. In particular, we compare our summary that consists of the chosen tweets with the NBA official Game Story and examine how our summary agrees with it. The NBA official Game Story is a detailed game summary written by the editors of the NBA official web site which contains the meaningful information about the games, such as game status, representative events, good move, bad move, quotations from player interviews telling the state of the game and so on.

As an example for one game, Table 2 shows the 13 items described in the official Game story and the matched tweets in our summary side by side. This game was played between the Miami Heat and the Indiana Pacers on June 2nd. Empty entries in the table mean that our summary misses the item corresponding to the row. Since there exist two empty entries, our summary covers the 11 items. Therefore the coverage rate of our summary is $\frac{11}{13} = 84.6\%$. Whereas our summary cannot cover all the items described in the NBA official Game Story, it succeeds in acquiring useful information not written there. Table 3 displays the examples of the tweets with useful information that is not stated in the NBA Official Game Story, but contained in our summary. These tweets are either more detailed descriptions about the game than the NBA official Game Story or comments on the teams and the players.

Figure 4 summarizes the coverage rate for multiple games, i.e., the 7 games of the NBA eastern semi-final between the Heat and the Pacers. Here, the x-axis denotes the game ID and the y-axis shows the coverage rate against the NBA Official Games Story. The mean coverage rate reaches up to 87%. Since our summary mines a lot of useful information not contained in the NBA Official Game Story, we consider that this coverage rate is acceptable.

Table 2 Sentences in the Official Game Story and the Matched Tweets in our Summary

	Official Game Story	Our Summary
1	The pacers defeated the Miami Heat 91-77 in Game 6	#Pacers win # Heat by 91-77
2	The Pacers limited Miami to 36-percent shooting and dominated things inside	
3	The Pacers won the rebounding battle 53-33 and outscoring the Heat 44-22 in the paint	Rebound 53-33, score in the paint 44-22, this the reason why the #Pacers won this game
4	LeBron James scored 29 with 7 rebounds and 6 assists	LBJ got 29 pts, 7 rebounds and 6 assists in game 6
5	In quarter 3, the Pacers outscored Miami 29-15, including a 12-0 burst early in the quarter , took a 68-55 lead into the fourth quarter.	At the end of the 3rd, the #Pacers are ahead of the Heat by 13.IND 68, MIA 55
6	Miller made his only two shots, both 3-pointers, for the Heat.	Mike miller is a game changed #heat
7	Had a putback dunk and then Miami unraveled completely. Called for an offensive foul	LeBron James just got a technical for throwing a tantrum
8	George Hil pushing Indiana’s lead to 81-68 with 3:55 remaining.	4 mins left and #pacers are up 81-68.
9	Dwayne Wade and Chris Bosh combined for 15 points on 4-for-16 shooting.	D.Wade and C.Bosh totally got 15 points
10	West missed his first seven shots and finished 5-for-14 for Indiana.	Wake up, Mr. West!!!
11	Joel Anthony came back into the rotation	
12	Paul George (28 points) and Roy Hibbert (24)	George 28 pts, 8 rebound and 5 assists, Roy Hibbert 20 pts, 11 rebounds.
13	Miami was without forward Chris ”Birdman” Andersen	This game just shows how Miami would have lost the last game if Birdman got ejected.

Finally, we show that our approach to compile the representative tweets for the clusters produced by the hierarchical clustering algorithm is effective to make the summary diverse. The tweets with respect to the NBA games are classified into the following 6 types:

(1) General status: showing the score of the game. For example, the scores of the two teams at half time and at the

Table 3 Tweets with Useful Information that is not Stated in the NBA Official Game Story

1	That dunk by Paul George over bosh!!! #pacers
2	Led by @PaulGeorge24's 9pts/4rebs, the #Pacers trail the Heat by 2 at the end of the 1st quarter.?MIA 23, IND 21
3	#Pacers only down 2 thru the 1st. Despite #Heat going 6-7 from downtown. #GoPacers #HEATvsPACERS #NBAPlayoffs
4	At least 15 points left on the board due to missed dunks layups #Pacers
5	8-1 free throw and 10-3 foul advantage for the #pacers in the first half. Where are all the people claiming the nba wants the #Heat to win
6	Hibbert with an easy layup, #pacers lead the #Heat 51-42 midway thru the 3rd quarter
7	A tale of two 3rd quarters in game 5 Miami won the third quarter 30-13 in game 6 Pacers won the quarter 29-15 #ECF #heat #pacers #games6
8	Omg the heat coming back 2 fast. C'mon #pacers
9	Paul George Wt 3 t put #pacers 7 up
10	Bosh is so bad its 3 min left in the game and dude is sitting on the bench. How is this guy getting paid 15m??? #Heat #NBAPlayoffs
11	This is a good #nbaplayoffs game #heatvspacers! #Pacers came to play.
12	This game is really intense.#Heat #Pacers

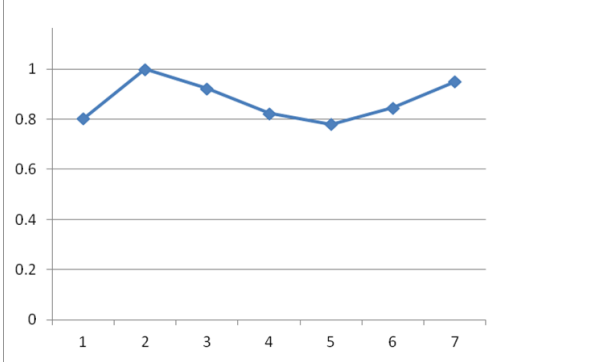


Figure 4 Coverage Rate of our Summary against the NBA Official Games Story

end of the game.

(2) Detailed information: showing the detailed information of the game, e.g., the number of rebounds, assists, and the points that a specific player achieves.

(3) General events: These are common events which are observed in many games, e.g., brilliant block, slam dunk, three pointer and game winner.

(4) Special events: These are rare events that may not happen in every game, e.g., breaking the record of three point and long consecutive win in the NBA history.

(5) Comments on the performance of the teams.

(6) Comments on the performance of players.

As for IM5 in Table 1 which corresponds to the discussion at the halftime, our summary covered all the 6 information types. Table 4 exhibits the tweet examples in our summary for each information type. On the other hand, when we chose randomly the same number of tweets from the whole set of tweets associated with IM5, only 2 information types were covered, because too personal tweets were put into the summary. The above fact provides some support for the claim that the agglomerative hierarchical clustering helps augment the diversity of our summary.

Table 4 Types of the Representative Tweets for IM5

Type of information	Instance of detected information
General status	Score of half time: 40-39. Pacers down by 1 point
Detailed information	1 8-1 free throw and 10-3 foul advantage for the pacers in the first half 2 Wade has only 1 point and Bosh with 3
General events	1 D-Wade steal, L-James dunk 2 Sam Young hits the Pacers' first 3-pointer of the night
Special events	broke the record for the number of missed dunks in a half
Comments on teams	1 Bad sign for Pacers. Down by 1 at the half. Should be up by at least 8. 2 Pacers end yet another quarter terribly 3 Both teams are lucky as hell.
Comments on players	1 Wake up Mr West 2 NOTHING from DavidWest

6. Conclusion

This paper proposes a system which generates the game summary for NBA basketball games by quoting the representative tweets issued during the games for the first time, while the previous researches dealt with soccer and American football. It is more challenging to make a game summary for basketball than for football or soccer, since one shoot/goal is less important for basketball. To enrich the summary, we actively put not only the real-time actions such as beautiful slum dunks which are described in the steep spikes in the tweet volume graph, but also the discussions at the break times between quarters which usually contain both the game status and the comments on the previous quarter. We propose a heuristic approach to focus on the spikes of long duration, even if they have rather gentle slopes in order to gather

such discussions. Our algorithm successfully finds the ground truth important moments in the game more accurately than the one which considers the slope of the spikes only. Furthermore, selecting the representative tweets per cluster which is derived by the hierarchical clustering algorithm applied to the tweets associated with an important moment increases the diversity of our summary. As the result, our summary covers up to 87% of the items in the NBA Official Game Story. Although our summary can not cover all the items stated in the NBA Official Game Story, it also contains a lot of useful information not contained there.

There remains problems to be solved in future: First, we need to make the evaluation method more reliable, since the current evaluation method depends on the manual works by ourselves much and is subjective to some extent. Next, our current system presumes an offline environment. We will extend it, so that meaningful tweets are detected real-time from the tweet stream and displayed to the NBA fans. Increasing the readability of our summary should be also pursued, since our current system displays the representative tweets as they are.

Acknowledgment

This work is supported by the Ministry of Education, Culture, Sports, Science and Technology, Grant-in-Aid for Scientific Research (C) 24500111, 2013.

References

- [1] T. Sakaki, M. Okazaki and Y. Matsuo, “Earthquake Shakes Twitter users: Real-time Event Detection by Social Sensors”, in Proc. of WWW’10, pp. 851–860, 2010.
- [2] S. Petrovic, M. Osborne and V. Lavrenko, “Streaming first story detection with application to Twitter”, in Human Language Technologies: The 2010 Annual Conference of the North American Chapter of the Association for Computational Linguistics (NAACL-HLT), pp.181–189, 2010.
- [3] D. Chakrabarti and K. Punera, “Event Summarization using Tweets,” in Proc. of ICWSM, 2011.
- [4] S. Zhao, L. Zhong, J. Wickramasuriya and V. Vasudevan, “SportSense: Real-Time Detection of NFL Game Events from Twitter,” CoRR abs/1205.3212 (2012).
- [5] J. Nichols, J. Mahmud, and C. Drews, “Summarizing Sporting Events Using Twitter,” in Proc. of IUI’12, pp. 189–198, 2012.
- [6] D.A. Shamma, L. Kennedy and E.F. Churchill, “Tweet the Debates: Understanding Community Annotation of Uncollected Sources,” in Proc. of the first SIGMM workshop on Social media, 2009.
- [7] D.A. Shamma, L. Kennedy and E.F. Churchill, “Peaks and Persistence: Modeling the Shape of Microblog Conversations,” in Proc. of CSCW’11, pp.355-358, 2011.
- [8] J. Weng and F. Lee, “Event Detection in Twitter,” in Proc. of ICWSM, 2011.