

Twitter の実況書き込みを利用したスポーツ番組のイベント抽出 A Method of Event Extraction from Sports program using Tweets

丸尾将輝[†]
Masaki Maruo

大野将樹[†]
Masaki Oono

沼尾雅之[†]
Masayuki Numao

1. はじめに

近年、テレビ放送の多チャンネル化や録画機器の進化によって視聴できる番組の選択肢が大幅に広がった。それに伴い番組を録画したビデオが大量に溜まってしまいそれらを消化するのが困難になるという問題が起こっている。そこで未視聴ビデオを効率良く視聴するためにビデオの内容検索やシーン検索技術が求められている。ビデオのシーン検索のためにはビデオを時間ごとにアノテーションする必要がある。山本らの研究 [3] によるとビデオに対するアノテーションの方法には、音声認識や画像認識によって完全に自動で行う方法、専用のツールを用いて人間の手を介して半自動的に行う方法、ネットワーク上から必要な情報を収集し利用する方法がある。このうち自動でアノテーションを行う方法は後述する既存研究のように利用法を限定しなくては精度を高めることが困難であることから、本研究ではネットワーク上で人間が投稿した情報を利用する方法に注目する。

Twitter とは 140 文字以内の短文を投稿することができるマイクロブログであり、SNS (ソーシャル・ネットワーク・サービス) の一つである。Twitter にはそのコメント投稿の手軽さから、今まさに起こったことや見ているもの、感じたことがコメントとして投稿されるという特徴がある。また Twitter ではハッシュタグと呼ばれる文字列を書き込みに含めることで、自分の投稿にタグを付けることができる。このハッシュタグを利用することで、特定の話題についての書き込みを閲覧したり、知り合いではない人との意見交換を楽しむことが可能となっている。このような特徴を利用して、Twitter ユーザーの中でテレビ番組を見ながらリアルタイムにコメントや意見を Twitter に投稿し、ハッシュタグで共有するという文化が広まっている。このような書き込みは「実況書き込み」と呼ばれ、新たなテレビ放送の視聴方法として一般的になりつつある。そしてこの実況書き込みは視聴者から即時的に発信される貴重な情報資源として、その利用価値の高さが注目されている。Twitter では開発用に複数の API を提供している。[2]

本研究ではビデオのシーン検索を行うための技術として、テレビ番組の特にスポーツ番組の時間毎のイベント抽出を行うことを目的とする。スポーツ番組を対象とする理由は、スポーツ番組を録画したビデオの中でシュートシーンを集めたハイライトを作成したいといったようなシーン検索に対する需要が高いジャンルであることや、スポーツなのでルールによって定められた展開をするためにイベントの定義がしやすい為である。ここで「イベント」とはそのスポーツ番組内で起こっ

ている出来事を意味する。さらにイベントの内容を表す「ラベル」を付与することで抽出されたイベントの情報としての価値向上を目指す。これらの目的実現の為に、ネットワーク上の情報である Twitter の実況書き込みを利用して、イベント抽出とラベル付与を行う手法について提案する。

本稿では、第 2 章にて関連研究について、第 3 章で提案する手法について、第 4 章で行った実験について説明し、第 5 章で考察、第 6 章でまとめを行う。

2. 関連研究

スポーツ番組のイベントを自動的に抽出することを目指した研究として、岩井らの研究 [4] がある。

この研究ではサッカー中継を解析の対象とし、画像認識の技術によってイベントの抽出を行なっている。手法としてゴールポストやゴールラインといったような不動であるオブジェクトとカメラの位置の相対関係から特定映像での特徴量を抽出し、検出時に抽出した特徴量と比較することでカメラの状態を認識してイベントの抽出を行う。特徴量の抽出には画像のエッジ検出が可能な SUSAN オペレーター [5] を用いる。実験では特定映像をコーナーキック時の映像としてイベント抽出を行った。結果はボールが蹴り出された瞬間や選手がゴールから離れていく様子は検出できているが、その後のボールの行方や選手の詳細な動きは検出できていなかった。この他画像認識の技術によってイベント抽出を試みる研究として、Nguyen ら [6] の研究では野球中継に対して P. Chang らの隠れマルコフモデルを利用した手法 [7] を元にシーンの抽出を行なっているが、この研究では既にダイジェストとなっているデータを入力として与えており、一連の番組内でのシーン検索を実現していない。これらの研究結果から自動的にスポーツ番組のイベント抽出をする方法は特定のシーン限定であるなど制限付き条件下でないと抽出が難しいことがわかる。

ネットワーク上の情報を用いてスポーツ番組のイベント抽出を目指す研究として、Twitter の実況書き込みを利用してイベントを抽出を行う、中澤らの研究 [8] がある。

手法としてはまず時間ごとの書き込み数の頻度から重要シーンの検出を行う。具体的には投稿された書き込みを一定時間ごとに数え、ニュートン法より時系列変化の極大値を求めることにより書き込みが急増するシーンの抽出を行う。この手法は David A. Shamma らの研究 [9] によるものである。次に検出された重要シーンの中で特徴的なキーワード、特に人物名に注目しイベントの主要人物を特定する。さらに、主要人物名が含まれる書き込みにおいて共起している名詞をイベント内容推定用のラベルとして付与する。実験は野球中継を

[†]電気通信大学情報理工学系研究科情報・通信工学専攻

対象とし、人手で付けたラベルとの一致率に関して評価を行なっている。評価結果として重要シーンの抽出が 94.6%、そのイベントでの主要人物の特定は 83.6%と高い一致率だったが、イベントの内容を示すラベル付けは 25.2%と低い結果になっていた。本研究では中澤らの研究と同じようにネットワークの情報として Twitter の実況書き込みを利用してイベント抽出を行うが、既存研究で実現ができていない書き込み頻度に依存しない、精度の高いラベル付与を伴ったイベントの抽出を目指す。

3. 提案手法

本研究ではイベントごとに出現しやすいキーワードを集めた辞書（イベント辞書）を作成し、解析に利用することによってイベントの抽出を行う手法を提案する。イベント検出をキーワードの出現頻度を指標として行うことで、書き込み頻度に依存しないイベントの抽出を目的とする。イベント辞書は番組のジャンルや内容ごとに一つ作成するものとする。例えば、サッカー中継のためのイベント辞書と野球中継のためのイベント辞書は別に用意する。イベント辞書は同じジャンルの複数の番組で取得した実況書き込みで学習を行い半自動的に作成する。作成したイベント辞書を利用して、最終的にラベル付けされていない実況書き込みデータから時間ごとに番組でどのようなイベントが起きたかを予測するシステムを作成する。本提案システムではイベント辞書の作成までをイベント辞書生成部、イベント辞書を使ったイベントの抽出までをイベント抽出部とする。

3.1. 提案システム

Twitter のストリーミング API によって得た実況書き込みをデータベースに保存し、提案システムの入力用実況書き込みはそのデータベース内から利用される。

提案するシステムはイベント辞書生成部とイベント抽出部に分かれ、図 1 に示すような構成になっている。本システムではイベント辞書生成部とイベント抽出部が個別で動く為に、もし仮に同じジャンルに対するイベント辞書が既に用意されている場合、見視聴番組の実況書き込みさえ用意すれば解析だけを行うことができる。既存のイベント辞書にさらに新規のイベントを追加していき、より詳細なイベント抽出を実現する新たなイベント辞書を作成することも可能とする。

3.2. イベント辞書の作成

イベント辞書生成部でイベントに対する反応が見られる時間帯で手動で区分けした書き込みを入力し、キーワード抽出を行うことでイベント辞書を作成する。キーワードは形態素レベルで分解したものを 1 単語とする。形態素への分解には形態素解析器 MeCab[1] を用いる。分解した単語で学習時に定める不要語辞書と不要語ルールに該当しない単語だけキーワードの候補として採用する。キーワード抽出には TF・IDF 法を用いる。TF・IDF 値はイベントごとに纏めた書き込み群を 1 文書として以下の式で計算する。

$$TF(\text{単語 } i, \text{文書 } j) = \frac{\text{文書 } j \text{ での単語 } i \text{ の出現回数}}{\text{文書 } j \text{ 中の全単語の出現回数}} \quad (1)$$

$$IDF(\text{単語 } i) = \log \frac{\text{文書の数}}{\text{単語 } i \text{ が出現する文書の数}} \quad (2)$$

$$TF \cdot IDF = TF \times IDF \quad (3)$$

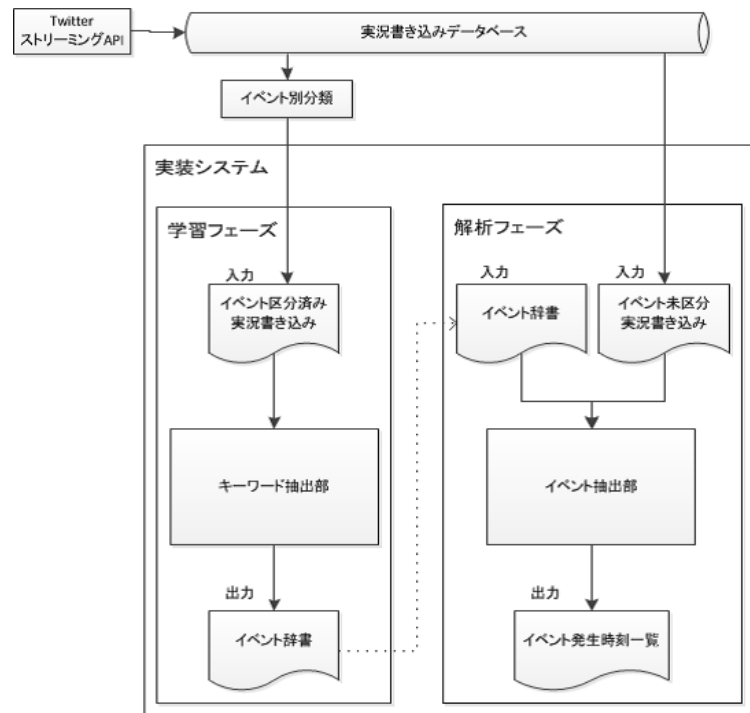


図 1: 提案システム概要図

イベント辞書にはイベントごとのキーワードと TF・IDF 値がセットになって保持される。学習した番組だけで特徴的なキーワードが登録されないように番組ごとに不要語辞書を持たせることでキーワードの制限を行うことを可能にする。

3.3. イベント抽出

イベント抽出部ではイベント辞書を利用して、イベントごとに定めた時間窓幅ごとにイベントの発生の有無を検出することで各イベントが発生した時刻一覧を作成する。

イベント抽出は以下のように行われる。

- (1) 書き込みを一定時間窓幅でフレームに分割。
- (2) フレームから書き込みを 1 つ取り出し単語に切り分ける。
- (3) 各単語とイベント辞書のキーワードの一致数をカウントする。
- (4) キーワード一致数の最も多いイベントをその書き込みの内容とする。
- (5) 2~4 を全フレームで行う。
- (6) フレーム j のイベント i に対する得点 P_{ij} を下式で算出する。

$$P_{ij} = \frac{\text{フレーム } j \text{ 内のイベント } i \text{ に対する書き込みの数}}{\text{フレーム } j \text{ 内の全イベントに対する書き込みの数}} \quad (4)$$

- (7) ($P_{ij} > \text{イベント検出閾値}$) の時フレーム j でのイベント i 発生を認定する。

なお利用するキーワードは TF・IDF 値からイベントごとに定めた閾値によって制限を与える。同時刻で複数のイベントが検出された場合、そこではそれらのイベントが同時もしくは連続して起きているとみなし

認定する。イベント辞書に登録された全てのイベントの検出が終わった後に、登録されたイベントではないが注目度の高いイベントを書き込み数によって解析する。平均の 2 倍以上の書き込みが観測され、かつ登録されたどのイベントにも分類されない時「未定義イベント」として補助的にイベント認定される。全ての解析段階を終えると時刻ごとに抽出されたイベントが一覧となって表示される。

4. 実験

4.1. 実験データ

今回の実験では対象番組のジャンルをサッカー中継とした。特に国際大会のサッカー中継のデータを利用した。

- ・男子日本代表を応援するハッシュタグ #daihyo
- ・女子日本代表を応援するハッシュタグ #nadeshiko

それぞれ試合が中継されている時間に投稿された書き込みをハッシュタグでフィルターをかけて取得した。

スポーツ中継を対象とする場合、その書き込みがどちらのチームを応援しているかを判別する必要があるが、今回は日本を応援するハッシュタグを含む書き込みを対象として解析を行なっているので取得した書き込みデータは全て日本を応援している書き込みであるとした。

今後の実験データにおいて「味方」としているのは「日本」側を意味する。

表 1: 実験に利用したデータ一覧

id	書き込み数	味方チーム	敵チーム
a	5,077	日本女子	タイ女子
b	11,490	日本女子	韓国女子
c	6,564	日本女子	オーストラリア女子
d	29,652	日本	北朝鮮
e	38,887	日本	ウズベキスタン
f	24,348	日本	タジキスタン
g	5663	U-22 日本	U-22 クウェート
h	22,709	日本	韓国
i	16,734	日本	北朝鮮
j	9,600	U-22 日本	U-22 バーレーン
k	10,061	U-22 日本	U-22 シリア

抽出を試みるイベントとラベルを

- (1) 味方チームにレッドカード、イエローカードが出されたシーン (味方カード)
- (2) 味方チームがゴールを決めたシーン (味方ゴール)
- (3) 味方チームがシュートを外したシーン (味方シュート)
- (4) 味方チームの選手交代が行われたシーン (味方交代)
- (5) 味方チームにレッドカード、イエローカードが出されたシーン (敵カード)
- (6) 味方チームがゴールを決めたシーン (敵ゴール)
- (7) 味方チームがシュートを外したシーン (敵シュート)
- (8) 味方チームの選手交代が行われたシーン (敵交代)
- (9) 前半開始、前半終了、後半開始、後半終了が宣言されたシーン (時間)

以上の 9 つとした。

4.2. イベント辞書の作成

イベント辞書作成を行う実験を行った。イベント辞書作成の為に利用する不要語ルールと不要語辞書は以下のように定めた。品詞や単語の区別は MeCab によって行った。

以下の条件を満たす単語は不要語と判定し、イベント辞書のキーワードに採用しない。

- (1) 「動詞」「名詞」「形容詞」以外の品詞である。
- (2) 「非自立」「数」「人名」「地域」「組織」「代名詞」「サ変接続」のいずれかに該当する。
- (3) ひらがな 2 文字もしくはひらがな 1 文字である。
- (4) 「ー」で始まる単語である。

不要語辞書

不要語辞書には以下の情報を与えた。今回は JFA によって公式に発表されている情報のみを与えた。これらの情報と一致する単語はイベント辞書のキーワードに採用しない。

- (1) 敵味方のチーム名
- (2) 敵味方の選手名
- (3) 大会名

学習には表 1 で示した全 11 試合の書き込みデータを使用した。学習を行った結果は 11 試合で、選別を行った後の全書き込み数は 180539 であり、イベントごとの書き込み数は表 2 のようになった。最終的に作成されたイベント辞書のキーワードと TF・IDF 値の上位 10 単語は表 3 のようになった。

表 2: イベントごとの書き込み数

ラベル	書き込み数
味方カード	447
味方ゴール	15561
味方シュート	11043
味方交代	4769
敵カード	2338
敵ゴール	2748
敵シュート	5837
敵交代	492
時間	8331
other	128973

表 3: 作成されたイベント辞書

味方カード	味方ゴール	味方シュート	味方交代	時間
キーワード TF・IDF 値	キーワード TF・IDF 値	キーワード TF・IDF 値	キーワード TF・IDF 値	キーワード TF・IDF 値
イエロー 0.072203	キキ 0.020931	クロスバー 0.013364	つかれる 0.013946	キックオフ 0.036531
イエローカード 0.055926	キター 0.008301	おしい 0.013258	乙 0.009796	ハジマタ 0.004676
面成敗 0.02222	キタ 0.006384	バー 0.012704	out 0.009647	オウタ 0.004378
カレー 0.021579	ゴールキタ 0.005969	惜しい 0.00957	in 0.008947	加好ー 0.004007
余計 0.015353	はいる 0.005297	しいる 0.007641	疲れ 0.008845	はじまる 0.004003
寸前 0.012236	ゴ 0.005093	ポスト 0.007343	IN 0.008341	つかれる 0.003848
カード 0.010101	ル 0.004923	宇宙 0.006445	野 0.007024	黒星 0.003803
もらう 0.009412	ッ 0.004464	正面 0.00511	代わる 0.006316	引き分け 0.003719
勿体 0.009319	けんこ 0.004226	嫌う 0.004383	アウト 0.005918	ハジマタ 0.003516
イエロ 0.008655	かざる 0.004016	いいい 0.004025	OUT 0.005088	おわる 0.003136
敵カード	敵ゴール	敵シュート	敵交代	
キーワード TF・IDF 値	キーワード TF・IDF 値	キーワード TF・IDF 値	キーワード TF・IDF 値	
レッド 0.053024	同点 0.01008	あぶる 0.036981	out 0.016265	
レッドカード 0.025422	orz 0.009661	あぶない 0.016601	下げる 0.015982	
キムチ 0.022923	滑る 0.007473	神 0.011228	世 0.010945	
イエロー 0.013838	あらあら 0.006139	危 0.009509	イグ 0.008697	
キンコン 0.011898	取り返す 0.005521	かわず 0.007241	ハゲ 0.008005	
レッドキタ 0.011831	追いつく 0.005303	スキー 0.006637	テセアウト 0.007926	
赤紙 0.010485	やりやりや 0.005266	こえる 0.006496	下ろす 0.007926	
赤 0.008955	あつけない 0.004876	アブ 0.00631	イグノ 0.006973	
イエローカード 0.008766	ああああ 0.004804	危ない 0.006213	イミフ 0.006955	
カレー 0.007092	あああ 0.004387	ナイス 0.005537	OUT 0.006638	

4.3. イベント抽出精度評価

評価は抽出されたイベントと正解データで F 値を算出して行った。5 試合の書き込みで学習し、5 試合の書き込みの解析を行った。試合ごとの平均 F 値は表 4、各イベントの適合率と再現率、F 値は表 5 のようになった。解析した 5 試合の中で最も平均の F 値が高かったデータ j についてのシステムによる出力は表 6 のようになった。比較の為に正解データを表 7 に示す。表 6 と表 7 ではシステムによる出力と正解データが一致した部分に色付けを行っている。

表 4: 試合ごとの F 値の平均

解析データ	F 値平均
g	0.216
h	0.259
i	0.371
j	0.560
k	0.552

表 5: イベントごとの F 値の平均

ラベル	適合率	再現率	F 値
味方カード	0.400	0.400	0.400
味方ゴール	0.650	0.867	0.707
味方シュート	0.403	0.46	0.423
味方交代	0.540	0.587	0.503
敵カード	0.310	0.300	0.212
敵ゴール	0.200	0.200	0.200
敵シュート	0.388	0.463	0.357
敵交代	0.110	0.267	0.138
時間	0.524	0.700	0.584

表 6: システムによる出力

システム出力	システム出力
0:00:00 時間	0:44:00 敵カード
0:09:00 敵シュート	0:45:00 敵シュート
0:10:00 敵ゴール	0:46:30 敵シュート
0:12:00 味方シュート	0:50:00 味方ゴール
0:16:00 味方シュート	0:51:00 敵シュート
0:18:00 味方シュート	0:52:00 味方シュート
0:21:00 敵シュート	0:54:00 時間
0:22:00 敵カード	1:04:00 敵交代 味方シュート
0:24:00 敵カード	1:06:00 敵交代 味方交代 味方シュート
0:26:00 敵ゴール	1:08:00 時間
0:27:00 敵シュート	1:10:30 敵シュート
0:28:00 敵カード	1:15:00 敵シュート
0:28:30 敵シュート	1:16:00 味方交代
0:30:00 味方シュート	1:18:00 味方交代
0:37:30 敵シュート	1:20:00 味方シュート
0:38:00 敵シュート	1:22:30 敵シュート
0:40:00 味方シュート	1:24:00 味方シュート
0:42:00 味方シュート	1:25:30 敵シュート
	1:26:00 味方シュート 敵カード
	1:30:00 味方ゴール
	1:32:00 味方シュート
	1:34:00 味方交代
	1:38:00 味方シュート
	1:40:30 敵シュート
	1:42:00 敵交代
	1:44:00 敵カード
	1:46:00 敵カード
	1:48:00 味方交代 敵カード
	1:55:30 敵シュート
	1:56:00 味方シュート
	1:58:00 時間
	2:00:00 時間

表 7: 正解データ

正解データ
0:07:00 時間
0:26:00 味方シュート
0:30:00 味方シュート
0:36:00 味方シュート
0:40:00 敵シュート
0:42:00 味方シュート
0:45:00 味方カード
0:47:00 敵シュート
0:51:00 味方ゴール
0:54:00 時間
1:08:30 敵交代 時間
1:10:50 味方シュート
1:15:30 敵シュート
1:16:30 味方交代
1:20:30 味方シュート
1:22:30 敵交代
1:27:30 敵カード
1:30:30 味方シュート
1:31:30 味方ゴール
1:32:30 敵カード
1:35:30 味方交代
1:36:30 敵カード
1:42:40 敵交代
1:43:30 味方シュート
1:44:30 敵カード
1:48:30 味方交代
1:54:30 敵シュート
1:56:30 敵シュート
1:57:10 味方シュート
1:58:00 時間

5. 考察

5.1. イベント辞書について

抽出されたキーワードを見てみると、ゴールやシュートなど得点に関わるイベントでは感情的なキーワードが多く集まりやすくなっていた。顔文字を含んだ書き込みも多く見られたので、今回のシステムには実装していないが、顔文字の辞書を追加すれば特徴的なキーワードとして利用できると思われる。イベントごとのキーワード数や TF・IDF 値の差は学習に使用した書き込み数に差がある為に生じる。

学習するデータごとの書き込み数の差から特定の番組での特徴的なキーワードが強く出てしまうことがある。なるべく汎用的なイベント辞書を作成する為には、学習時にただ書き込みデータを追加していくのではなく、学習データの総書き込み数を考量して更新用ファイルとの結合を行う必要があると考えられる。

不要語辞書は、外国人選手の名前の除外に大きく寄与している。特にイベント「敵交代」でキーワードと TF・IDF に大きく影響を与えていることがわかる。しかし不要語ルールで「人名」を除外しているため元々日本人の名前はキーワードから除外できていた為、日本人の選手名情報は辞書作成の上ではあまり貢献しなかった。さらに不要語辞書作成の際には公式に公開されている本名を不要語辞書に登録しているので、愛称や外国人選手のカタカナ表記の揺れには対応できない為、いくつか人の呼称である単語がキーワードとして抽出されてしまっている。愛称の問題に関しては Wikipedia の愛称の項を利用して愛称を吸収する手法や、外間らの検索エンジンを利用して愛称を抽出する手法 [10] が利用できると考えられる。

5.2. イベント抽出について

試合ごとの抽出の精度は最高が F 値 0.560、最低で F 値 0.216、平均 F 値 0.391 という結果となった。この結果から試合によってイベントの抽出の精度に差が生じることがわかる。原因として同じイベントでも試合展開によって登場するキーワードに差がでることが挙げられる。特に今回はイベント辞書作成の学習を 5 試合でしか行なっていないために、様々な試合展開で各イベントに登場するキーワードが学習できず、この問題が生じていると考えられる。

また、評価方法をイベントの内容を考慮せずに一律で F 値によって決めているので、例えばゴールが 1 本もなかった試合でどこかの時刻にゴールイベントを 1 度誤検出してしまうだけで F 値が 1 から 0 になってしまうという、評価方法の問題もある。

イベントごとの抽出の精度は最高が F 値 0.707、最低で F 値 0.138、平均 F 値 0.391 という結果となった。この結果からイベントの内容によって抽出の精度に差が生じることがわかる。傾向としては反応が大きく書きこみ数も多いイベントの抽出精度は全体的に高くなっている。

今回は学習時に日本を応援する投稿のみを学習した為に敵チームのイベントに対する書き込みが全体的に学習できず、またイベント抽出時も対象となる書き込みが少なくなるためにイベントの抽出精度が低くなっている。スポーツ中継を解析対象とする際には両チームの応援書き込みを学習し、解析することで偏りのないイベント抽出ができると考えられる。

誤検出の一因として、ハーフタイムに交代が可能であるために「時間」と「交代」、ファールからフリーキックとなるために「カード」と「シュート」といったイベントが同時刻に起こることがあり、このために幾つかのイベントで学習時に余分なキーワードがイベント辞書に登録されている可能性が考えられる。

学習した試合数が少ない上に、さらに発生しにくいイベント「味方カード」や「敵ゴール」は学習した書き込みの量が少なく特徴的なキーワード抽出や、安定したパラメータの決定ができていないために抽出の精度が低くなった。最もイベント抽出が難しいのは「敵交代」でこのイベントは表 2 から分かる通りイベントに対する反応が低く、ほとんどの試合で書き込み数が増加しないイベントである。本提案システムではこのようなイベントを抽出することを目標としている。結果は F 値 0.138 と他のイベントと比べれば精度は低い再現実率 0.27 となっているので既存手法では埋もれてしまっていたイベントの抽出にある程度は成功している。しかしこのイベントの精度の低さはイベントの内容的に「味方交代」とキーワードが被ってしまう点、不要語辞書で完全回避が難しいキーワードが頻出するイベントである点も原因だと考えられ、単に書き込みの低さだけで精度が落ちているとは考えにくい。この点を考慮して特徴的なキーワードの取得が可能かつ、イベントに対する反応書き込み数が少ないイベントで再評価を行う必要がある。

最も F 値が高かった試合に注目してみると、システ

ムが検出したイベントは 54 個、そのうち正解であったのは 20 個で適合率 0.370 と適合率はあまり高いとは言えない結果であったが、正解データの個数が 30 個であることから再現率が 0.667 となり、半分以上のイベントが検出できていることがわかる。特にゴールシーンや味方交代シーンといったサッカーの試合で重要とされやすいイベントは逃していないことから、再現率は良い精度が出ていると言える。しかし実際にサービスとして本システム利用することを想定すると、適合率が低くシステムの出力に間違った情報が多く混ざってしまうと、快適なシーンサーチ機能を提供できない。よって適合率はさらに上げる必要がある。

また全体の精度の不安定さは選択されるシステム上のパラメータセットによって結果が大きく変わることが原因と考えられる。本システムではイベントをどのような時間窓幅で抽出するか、イベント辞書のどのキーワードを利用するか、イベントの検出を認める閾値をいくつにするかというようなパラメータの設定を行う必要がある。今回はパラメータの選択を 4 試合での評価結果で良い成績が得られたものが良いパラメータであるとして行ったが、試合内容によってイベント抽出精度が変わっていることから、この選択方法ではあまり良いパラメータの選択はできていないと考えられる。よって学習時に試合内容を吸収してイベント辞書と共に、試合内容ごとに最適なパラメータセットも出力するようなシステムにするか、もしくはパラメータに大きく依存しない辞書の作成とイベント抽出のアルゴリズムを考える必要がある。

6. おわりに

本研究では Twitter の実況書き込みを利用したスポーツ番組のイベント抽出をイベントごとに出現するキーワードを集めた辞書を利用することで、書き込みの頻度だけによらないイベントの抽出と一意なラベル付与を可能とする手法を提案し、実装した。評価実験では番組内でイベントを 9 つ定義して抽出を行い、イベント抽出の精度を F 値によって評価した。5 番組でイベント抽出を行った結果、抽出の精度は最高が F 値 0.560、最低で F 値 0.216、平均 F 値 0.391 という結果となった。イベントごとの精度は最高が F 値 0.707、最低で F 値 0.138 という結果となった。このように同じジャンルの番組でも番組の内容やイベントの内容によって精度に大きな差が生じた。また全体的に適合率より再現率の結果が良く、必要以上にイベントを抽出してしまっていることがわかった。

今回の実験ではイベントの抽出に利用するパラメータを数試合の解析を行った後に評価の平均が高くなったパラメータを採用することで決定した。しかし、この方法では試合展開によって精度が高くなるパラメータの選択が行えないことがある。そこで今後の課題として、試合展開を入力することで汎用的なパラメータが出力されるようなパラメータ決定用の学習システムを作成するか、あるいはイベント抽出時に解析対象の書き込み数や使用するイベント辞書によって自動的に適切なパラメータを選択するシステムを作成する必要

がある。また適合率を上げるために、イベントごとのキーワードの出現頻度の余韻の違いに注目することで、イベントの終了時刻を検出することができれば精度がより向上すると考えられる。

参考文献

- [1] <http://mecab.sourceforge.net/>
- [2] 辻村浩 (2010) 『Twitter API プログラミング』ワークスコーポレーション 335p.
- [3] 山本大介, 長尾 確: 閲覧者によるオンラインビデオコンテンツへのアノテーションとその応用, 人工知能学会論文誌, Vol. 20, No. 1, pp. 67-75, 2005.
- [4] 岩井, 丸尾, 谷内田 他: 「サッカー映像からの特定映像イベントの抽出」電子情報通信学会技術研究報告. MVE, マルチメディア・仮想環境基礎 99(183), 31-38, (1999-07-15)
- [5] S.M Smith, J.M Brady : “ SUSAN . A New Approach to Low Level Image Processing ”, International Journal of Computer Vision, 23(1):45-78, 1997
- [6] Nguyen Huu Bach, 篠田浩一, 古井貞熙. 隠れマルコフモデルを用いた野球放送の自動的インデキシング. 電子情報通信学会技術研究報告 PRMU2004-107, pp. 13-19, 2004.
- [7] P. Chang, M. Han, and Y. Gong, “ Extract highlights from baseball game video with hidden Markov models, ” Proc. of the Int. Conf. on Image Processing (ICIP '02), vol. 1, pp. 1.609-612, 2002.
- [8] 中澤, 帆足, 小野, KDDI 研究所: 「Twitter を用いたテレビ番組からのイベント検出及びラベル付与手法」一般社団法人情報処理学会 全国大会講演論文集 2011(1), 517-519, (2011-03-02)
- [9] D A Shamma, L Kennedy, E F Churchill “ Tweet the debates ”, Proc WSM'09, 2009.
- [10] 外間智子, 北川博之: “ Web データを用いた人物の呼称抽出 ”, DBSJ Letters, Vol.5, No.2, pp.49-52 (2006).