



Mapping Temporal Horizons

Adam Jatowt, Émilien Antoine, Yukiko Kawai, Toyokazu Akiyama

► To cite this version:

Adam Jatowt, Émilien Antoine, Yukiko Kawai, Toyokazu Akiyama. Mapping Temporal Horizons: Analysis of Collective Future and Past related Attention in Twitter. International World Wide Web Conference, IW3C2 International World Wide Web Conference Committee, May 2015, Florence, Italy. pp.11, 10.1145/2736277.2741632 . hal-01195687

HAL Id: hal-01195687

<https://inria.hal.science/hal-01195687>

Submitted on 12 Sep 2015

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Mapping Temporal Horizons

Analysis of Collective Future and Past related Attention in Twitter

Adam Jatowt
Kyoto University
adam@dl.kuis.kyoto-u.ac.jp

Yukiko Kawai
Kyoto Sangyo University
kawai@cc.kyoto-su.ac.jp

Émilien Antoine
Kyoto Sangyo University
first.last@cc.kyoto-su.ac.jp

Toyokazu Akiyama
Kyoto Sangyo University
akiyama@cc.kyoto-su.ac.jp

ABSTRACT

Microblogging platforms such as Twitter have recently received much attention as great sources for live web sensing, real-time event detection and opinion analysis. Previous works usually assumed that tweets mainly describe “what’s happening now”. However, a large portion of tweets contains *time expressions* that refer to time frames within the past or the future. Such messages often reflect expectations or memories of social media users. In this work we investigate how microblogging users collectively refer to time. In particular, we analyze half a year long portion of Japanese and four months long collection of US tweets and we quantify *collective temporal attention* of users as well as other related temporal characteristics. This kind of knowledge is helpful in the context of growing interest for detection and prediction of important events within social media. The exploratory analysis we perform is possible thanks to the development of visual analytics framework for robust overview and easy detection of various regularities in the past and future-oriented thinking of Twitter users. We believe that the visualizations we provide and the findings we outline can be also valuable for sociologists and computer scientists to test and refine their models about time in natural language.

Categories and Subject Descriptors

H.5.m [Information Interfaces and Presentation]: Miscellaneous

Keywords

Time Mention; Twitter; Social Network; Temporal Analysis; Visualization

1. INTRODUCTION

Memory of past actions or events as well as expectations or predictions about future occupy much of our thinking every

day. Rather than constantly focusing on the present, human thoughts tend to stray over different time perspectives within near or more distant time. Memory is useful for analyzing the past and for understanding the present, while future expectations and future-focused thinking serve an important role of arranging our actions and preparing for what may come.

As both remembering and expecting represent significant amount of human cognitive activities, not surprisingly, also in social media, users tend to leave traces of expressions related to the past and future. Such temporal expressions constitute a wealth of information about what users are going to do or what they expect to happen, as well as, which past events played important roles in their lives or what they recall from the past. Recently, there has been significant interest in building models for prediction of real-world outcomes using social media. For example, it has been shown that movie box-office revenues can be successfully forecasted using Twitter [4]. Most of the proposed approaches relied on detecting periodical patterns or trends of user activities and extrapolate to forecast the future. We advocate here another approach to directly analyze expectations expressed by users for judging their future actions. For example, when users tweet about a meeting scheduled in the evening, they leave a strong signal about a future activity at the end of the day. We believe that the traces of future-related thoughts are valuable to be mined and analyzed as they can be used for recommendation, targeted advertising, predictive user assistance and group behavior forecasting or even for tracking user whereabouts in case of sudden, massive disasters.

Similarly, we consider past-related expressions to carry high significance for assessing self-reported importance of events and activities users had in the past, or, in general, to characterize remembering and reminiscing patterns. In social sciences the concept of *collective memory* as a shared remembrance of the past has been known for quite a long time. However, it was usually measured in small scale or through controlled studies. It has become possible now to investigate the memory decay on the collective of social media users.

In this work, we advocate the concept of *memory and expectation sensing* as a complement to the well-known notion of *social sensing* in microblogging and social media. As we are aware that human memory and futuristic thinking are inherently complex matters, we thus first attempt to uncover any observable regularities and patterns in data through exploratory analysis of the way in which social media users

collectively refer to the future and the past. By this we hope to shed more light on the feasibility of implementing effective future prediction solutions through improved feature engineering and for understanding the importance of future and past events for users.

In particular, we focus on Twitter which is commonly used in the nascent field of *computational social science* [20]. First, we collect tweets with any explicit temporal references treating them as reflections of user thoughts on the future and the past. We make use of a half year long portion of Japanese tweets from mid-July 2013 to mid-January 2014; and a four month long collection of tweets in the USA from end of September 2013 to mid-January 2014. Then we map any detected temporal expressions to their corresponding time points. This allows us to determine the *time mention* (disambiguated time expression) of tweets. Next, we manipulate the two temporal signals: tweet timestamp (when the user sends the tweet) and *time mention* (to what time period the user refers) in tweets. Finally, we allow uncovering several findings through the visual analytics process.

Note that we choose two culturally different regions in order to perform comparative analysis of the way in which users refer to both the past and the future in Japan and in the USA. We find many commonalities in the shape of temporal attention but, at the same time, we also pinpoint certain differences.

To sum up, we make the following contributions in this paper. We are the first to analyze the way in which microblogging users refer to the past and to the future, and to propose the idea of *memory and expectation sensing* in social media. Our focus is on two distinct countries and on tweets written in different languages. For enabling visual analytics, we build an *interactive system* that collectively visualizes historical and future-oriented perspectives in microblogging. The system is available online [41, 42] for anyone interested in a quick overview or in detailed exploration of temporal aspects in shared messages. Lastly, we discuss the analysis results and outline several avenues for further research.

The reminder of this paper is structured as follows. In the next section we review the related work. Section 3 describes the data processing model, while Section 4 introduces our methodological approach for visualizing temporal patterns in tweets and overviews findings from this study. We formally describe the technical aspects in Section 5 and summarize our findings as well as their implications in the subsequent section. Finally, Section 7 concludes the paper and outlines our future work.

2. RELATED WORK

Numerous research works have shown the usefulness of microblogging for real-time information analysis in many domains, ranging from detecting natural disaster [36] and disease outbreaks [17]; via analyzing how the popularity of topics emerge and evolve on time [40] and space [3]; to opinion detection and analysis of the crowd about events [37]. An increasing number of users leave numerous traces of their daily life activities [20], especially, in social networks. Some works then exploit this new flow of personal data to predict upcoming activities in the near future [4, 38] or hot emerging topics [7].

Our work is also related to the subset of socio-psychological sciences about the study of memory and expectations. *Collective memory* is a term introduced in social sciences by

Halbwachs [11] to define the collective view of society on the past. Russell Jacoby [13] also coined the term *collective amnesia* for describing forcible repression of memories when whole groups or nations “forget” events from the past. In parallel to *collective memory*, Ebbinghaus [9, 10] investigated the characteristics of personal memory and the forgetting process. Both the social and individual-focused memory studies were limited to small samples of users and had subjective character. Thus, computer science approaches could already offer a significant complement to manual investigations exploiting large data and high coverage. Already, some recent works started investigating collective memory [8, 18] or estimating collective predictions in news [14], and on the web [15, 31, 32]. Unlike those works, in our research, we focus on social media and we perform large scale analysis on the way in which microblogging platform users refer to time.

Memories and predictions rely on temporal analysis of text that models and leverage temporal expressions as well as their uncertainty [5, 14]. After prior detection and disambiguation, the temporal expressions can be included in the Temporal Information Retrieval (TIR) process [2, 5, 16, 19]. The objective of TIR is to improve the effectiveness of information retrieval methods by exploiting temporal information in documents and queries and by extending the traditional notion of *topical relevance* with *temporal relevance* [2, 6]. In fact, quite a significant fraction of Web queries have been found to contain explicit temporal expressions such as dates or names of calendar events [30, 34], while many other queries are known to have implicit temporal intent [26]. The importance of time comes from the fact that the value of information and its quality are intrinsically time-dependent, as shown in [25] that consider timeliness or currency as one of the five key aspects that determine a document’s quality; the others being relevance, accuracy, objectivity, and coverage. Studies in TIR have actually shown that the retrieval effectiveness of temporal queries can be significantly improved by modelling and taking into account publication timestamps and time mentions of documents [5, 16, 19]. Through this work we hope to provide better outlook on different ways in which users refer to diverse time periods and by this to offer clues for solving the challenge of designing effective TIR systems and for understanding society and humans.

Finally, regarding event detection within microblogging data, time analysis has been a major field of research as shown in [39, 34, 21]. Also, studies in [1, 23, 36] provide a suite of visualization capabilities to explore tweets from different dimensions relating to specific application like fire combat and earthquake detection. More generic platforms have been designed with visualization tools to portray time perspectives [24, 28, 29].

3. DATA MODEL

This section describes our model to organize the data to make it usable for visualization. The following entities are of our interest: *users*: people sharing messages; *time-referring messages*: atomic interactions of users with others that contain temporal expressions including expectations and memories about subjects; and *subjects*: objects mentioned in time-referring messages. Our field of study is a portion of Japanese and US Twitter data. Therefore, in our context the users are Twitter registered users, the messages are tweets written in Japanese and English and the subjects are topics of tweets. A tweet is a short message limited to 140 characters

used to communicate with others but also to react to certain events. We are particularly interested in time-referring tweet messages from which we can extract temporal expressions. The temporal expressions allow us to categorize tweets into those about future and those about past as well as to map them on timeline after prior disambiguation. We further describe the datasets in Section 4.1.

3.1 Data Structure

Each tweet is first represented as a tuple of raw data:

$\langle \text{user id}, \text{tweet content}, \text{timestamp} \rangle$

The *user id* is the Twitter identification number, the *tweet content* is the text of the tweet and the *timestamp* is the time when the tweet has been published on Twitter. Based on these attributes, the information related to time is extracted and is represented by two additional attributes:

$\langle \text{time mention}, \text{time diff} \rangle$

The *time mentions* are extracted and disambiguated from the *tweet content* by applying temporal entity-recognition techniques further detailed in Section 3.2. The *time diff* is then computed as the difference between the time mention and the timestamp. Intuitively, the quintuple of attributes should allow to answer simple but fundamental questions such as:

- *timestamp*: what people tend to say at the same time ?
- *time mention*: what people tend to say about the same time ?
- *time diff*: what people tend to say before or after the same time ?

The goal of our data model is to allow a combination of attributes to be contrasted leading us to infer information about the global and per user behavior of microblogging users.

3.2 Time Mention Extraction

As mentioned above we consider only tweets with time mentions. Note that while there are ready temporal taggers for English, to the best of our knowledge, none is available for extracting and anchoring temporal expressions from Japanese. Therefore, we use the Stanford CoreNLP tagger [22] for US dataset whereas we built our own tagging system for Japanese. We match tokens in the tweet content with a dictionary of the lexical expressions related to time in Japanese. For this we provide a set of common time expressions in Japanese accumulated from diverse resources such as the list of temporal expressions given by Japan Meteorological Agency¹ and others. The list we made includes different time-related phrases that refer either to near or distant time, both in future and in the past. We also asked 3 native speakers to ensure that no common temporal expressions have been missed and included their new suggestions. In addition, we defined a set of regular expressions to capture mentions of hours, days, including names of weekdays, months, seasons, national holidays and years. Eventually, the numbers of tweets containing time mentions detected in both datasets are comparable, as shown in Section 4.1, indicating that our dedicated system for Japanese language seems to be rather reliable.

We emphasize that differently to previous studies on analyzing or visualizing collective time references [14, 15], in this work, we use not only absolute temporal expressions (e.g., “January 2012” or “December 2, 2013”) but also we

extract and disambiguate relative temporal references (e.g., “next week” or “three months ago”). Intuitively, the latter are quite common in spoken and colloquial language and should be thus included in the analysis.

Lastly, we note that our model assumes the entire tweet content to be related to temporal expression that appears in the tweet. This assumption is reasonable given the short length of tweets.

Time Mention Disambiguation.

Each temporal expression is mapped to an absolute time period value in the past or in the future. The process is straightforward for absolute time expressions. Relative time expressions are resolved based on tweet timestamps.

Time Granularity Shifts.

It is known that humans tend to perceive time using scale similar to a logarithmic one [12]. When the time distance measured from the speaking time point increases, the granularity of temporal expressions used becomes coarser. For example, when someone talks about far future (e.g., several months or years ahead) he or she usually refrains from pinpointing exact hours, minutes or even days. On the other hand, when referring to near future, such as to actions in the same day or few days later, humans tend to use finer granularity time-references (hours or even minutes). This has been also confirmed by the analysis of time expressions extracted from large collections of news articles and then related to the article timestamps [14]. In our model we thus apply the logarithmic scale of time using three main time granularity levels inherent to human speech: minutes, hours and days depending on the distance to the start time of the disambiguated temporal expressions measured from the tweet timestamp (*time diff*). Below we show a sample of the mapping used in our model emphasizing heterogeneity of time granularity and the phenomenon of *granularity shifts* in natural language:

“Now”	→	timestamp (granularity of minutes)
“Tomorrow”	→	subsequent day of the timestamp (granularity of hours)
“Next month”	→	subsequent month of the timestamp (granularity of days)

Mapping Time Expressions.

Usually time mentions in human language do not express a specific point in time but rather a period. For example, for the time mention in the following tweet: “I will fly to Okinawa next week” the model should reflect the fact that the next week is a seven days long period, while the action (flight) may take place on any of them. To do this we represent time mentions by a uniform probability distribution over its time duration. For example, “tomorrow” is mapped into probability distribution over hours in the next day, while “next month” is converted to the uniform distribution spanning each day of the following month and so forth. For the Japanese dataset, we provide a set of this kind of mappings for each temporal expression based on the Japan Meteorological Agency guidelines². Figure 1 shows a sample of temporal expressions, translated from Japanese, referring to different time frames of a day (hour granularity) and their mappings to the absolute time values.

For the US dataset we consider as a *time expression*, an-

¹http://www.jma.go.jp/jma/kishou/known/yougo_hp/toki.html

²http://www.jma.go.jp/jma/kishou/known/yougo_hp/saibun.html

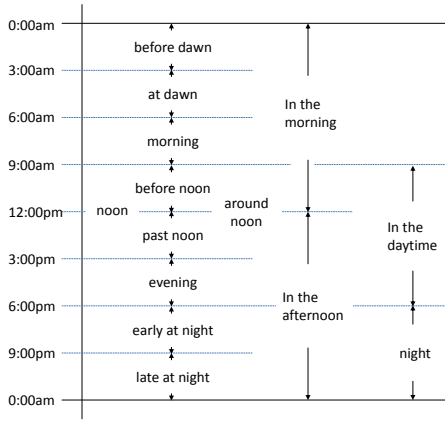


Figure 1: Sample of time expression mappings referring to different times of a day

notated expressions recognized as a *time* or *date* entity by the Stanford Name Entity Recognizer with some additional manually added ones to fit the same mapping style as the one for the Japanese dataset and to handle the granularity shift effect.

To sum up, in this work, the weight coming from a single time-referenced tweet on a particular time point is determined using the uniform probability distribution over the period defined by the temporal expression in the tweet and according to the selected granularity level for this expression. Following, the weight of the time point coming from several temporal expressions that cover this time point is calculated as the sum of their corresponding probabilities. The technical details of this probabilistic mapping are formally described in Section 5.1. Note that some previous works [14] adopted a more refined approaches to map activities occurring within a time period using either exponential, Gaussian or uniform distributions depending on temporal modifiers. However, deciding parameters of the probability distributions for different messages requires deeper NLP processing such as analyzing the character of events and activities, which is outside the scope of this work. Such processing would be also slow and require considerable computing resources for large datasets.

4. ANALYSIS

In this section we describe our visualization framework and detail findings we could obtain from the analysis. We defer the detailed description of the technical details until Section 5 to focus first on the conceptual way of portraying the data and on any observable patterns.

The following discussion will be mainly based on 2D plots called heat maps containing colored cells. A cell is the timespace on 2D time pane. The cell color represents the *intensity*, detailed in Section 5.1, with which tweets in the datasets refer to that cell. Below each graph, we display the color scale ranging from dark blue (the lowest *intensity*) to dark red (the highest *intensity*).

4.1 Distribution of Time References

First, we report the overall statistics of the data we use. The Japanese dataset has been build retrieving 31.6M tweets posted from Japan between July 21, 2013 and January 12, 2014; and the US dataset has been made by collecting 198M

tweets between September 25, 2013 and January 17, 2014. Unfortunately, due to technical obstacles and the lack of resources the data crawling was disrupted at certain times. This explains the blank sections in Figures 3 and 4.

We use the language detection method³, which is based on Naïve Bayesian filter and has nearly 99% precision, to select roughly 25M (millions) tweets written in Japanese in the Japanese dataset and 158M written in English from the US dataset. Among them we found 4M (16%) Japanese tweets and 30M (18%) English tweets that have temporal expressions. It means that a considerable fraction of messages contain some temporal clues. The slightly lower ratio of temporal expression usage in Japanese is more likely to be caused by our custom Japanese NLP parsing that may not be as efficient as the Stanford CoreNLP for English rather than by an actual difference within the datasets. The difference is rather small that indicates our custom Japanese *time expression* detection should work reasonably well.

In the table shown below, we report the overall strength of *temporal attention* of users towards the past, the present and the future as portrayed by the usage of temporal expressions towards each temporal class. Note that the present is determined by several words denoting the concept of “now” in Japanese and English.

	Japan	US
# time-referenced tweets	3.96M	29.92M
about past	38%	26%
about present	22%	16%
about future	40%	58%

These numbers seems to indicate that the overall *temporal attention* of Japanese Twitter users is equally distributed between the future and the past while the US Twitter users seem to be more oriented to future. This might be a language related difference since Japanese does not have an explicit future tense or difference related to the culture. However, we cannot rule out possible bias coming from different efficiency of parsers used for English and Japanese.

Next, we examine per user statistics. Figure 2 shows the distribution of users according to the frequency with which they use the *time mentions*. The horizontal axis is the ratio

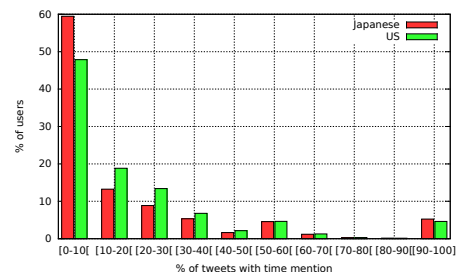


Figure 2: Rate of temporal expression usage in the tweets per user

between tweets with *time mentions* and the total number of tweet of a given user. The vertical axis denotes the percentage of users for each different usage rate of *time mentions*. We can notice that the ratio of users that at least once submitted a time-referenced tweet is 50% for Japanese and 65% for US. The value for the first bin [0 – 10%] shows that about 40% of Japanese users and over 50% of US users add time

³<http://code.google.com/p/language-detection>

mentions in more than 10% of their tweets. This indicates that temporal expressions are relatively commonly used in Twitter and can be used as a signal for measuring temporal orientation of user messages. The small bump (5.2% for US and 4.6% for Japan) in the last bin, which groups users who send more than 90% of tweets containing *time references*, is likely due to the presence of bots or automatically created accounts.

4.2 Scope of Temporal Attention

We expect the way in which users refer to the future or the past differs depending on whether the temporal perspective is near- or far-fetching. In particular, we expect that forecasting and remembering attention should drop over time. By this we mean that the amount of predictions and remembrances about near future and near past should be greater than the ones associated with more distanced events. Previous studies [14, 15] conducted on news article and web page collections have demonstrated the drastic drop in the count of references that point to time periods far away from document timestamps.

For testing this and other hypotheses we use the visualizations in Figure 3 to show the difference between the time mentioned in tweets and their timestamps. The top graph (Fig. 3(a)) displays the Japanese dataset and the bottom one (Fig. 3(b)) is for the US dataset. The timestamp values of tweets published on a particular day are given on the abscissa and their agglomerated time differences are shown in ordinate in logarithmic scale. Each graph in Figure 3 is composed of two parts, the top one showing the map of future-related expressions and the bottom one portraying the expressions about the past. For example, when looking at the future-related parts of the graphs, the cells falling into the segment of 1 - 7 days represent aggregate tweets that contain time references pointing to the period from 1 to 7 days later when counting from their timestamps. Similarly, the same segment on the past-related graph part represents the aggregate of tweets with time references referring to the span from 1 to 7 days ago from tweet timestamps. Since the horizontal line denotes the tweet timestamps, we should be able to observe fluctuations in the *temporal attention* over the whole time frame of our dataset. Note that as discussed in Section 4.1 data is missing at certain time points (indicated in the graphs in Figure 3 as blank columns).

To observe any potential calendar effect we indicate by black horizontal lines the start points of each week (dashed black lines) as well as the start points of each month (solid black lines). Below the graph we also display the curve of the total number of tweets crawled (in blue) and the percentage of the tweets with time mentions (in orange) to compare the visualization on particular day with the total amount of data used to visualize that day. On the right-hand side of the graph, the blue, vertical curve shows the amount of tweets mapped to every *time diff* interval (i.e. the sum of all the cell values on each lines).

We note that our system allows fine-grained investigation of tweets constituting any selected cell. Upon clicking on a given cell a new tab is opened in a browser with the list of top words, top temporal expressions and the table detailing all the attributes of tweets associated with the cell including their text. Similarly, selecting any gray cell on the right-hand side or on the top line will trigger pop-up window with ranked top terms for the cell. The computation of top words

will be described in Section 5.4.

Looking at both the future and the past related heat map graphs in both the datasets we observe the pronounced color gradation with some abrupt changes in color drawing parallel to the abscissa which indicate intensity levels that vary depending on the temporal scope of future and past references. This shows the tendency in human speech to use landmark expressions such as “today”, “this afternoon”, “next year” or “next week”. It appears that the strongest *temporal attention* seems to be mainly related to actions within the period up to the next 1 day in the future and down to 1 day in the past for both the languages (orange and yellow colors). Within these horizons, the activities up to the next hour and down to the past 10-15min seem to be even more often referenced, as especially visible in the US dataset. We then observe significant decline in the amount of temporal attention further above or below the lines of 1 day as marked by green or light yellow color areas. In the US dataset we can also observe the effect of Fridays and Saturdays during which users seem to make slightly more plans and forecasts up to 6-12h ahead (note the periodical, horizontal stripes in orange color) than on other days. In the Japanese dataset this is more blurred.

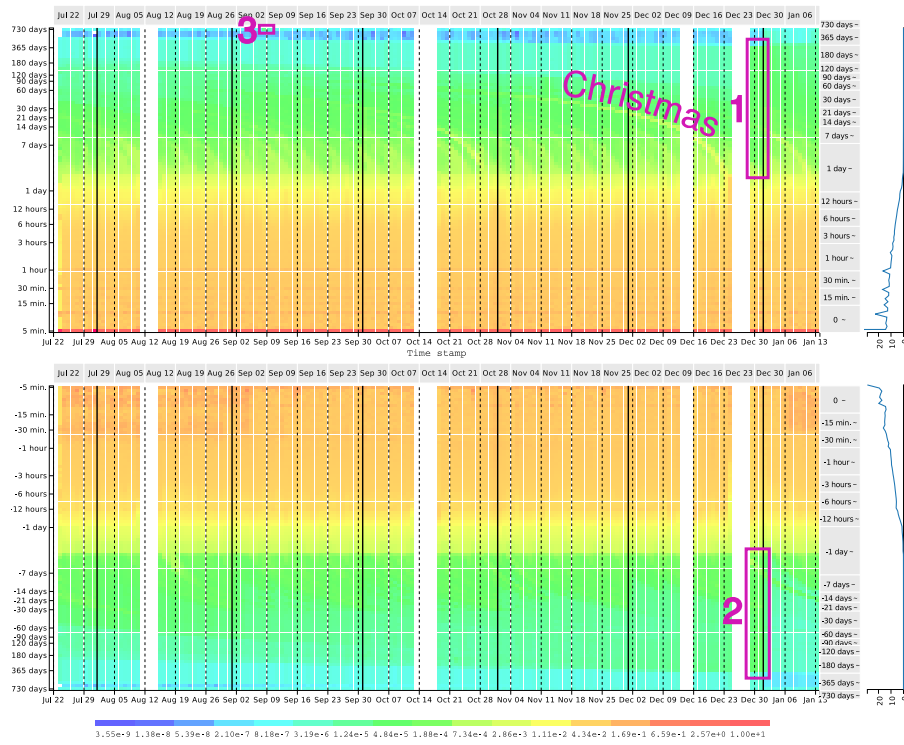
Lastly, we note that in general the far past and distant future are characterized by fewer references what confirms our initial hypothesis that the attention decreases the farther one looks into the past and the future.

4.3 Similarity of Past and Future

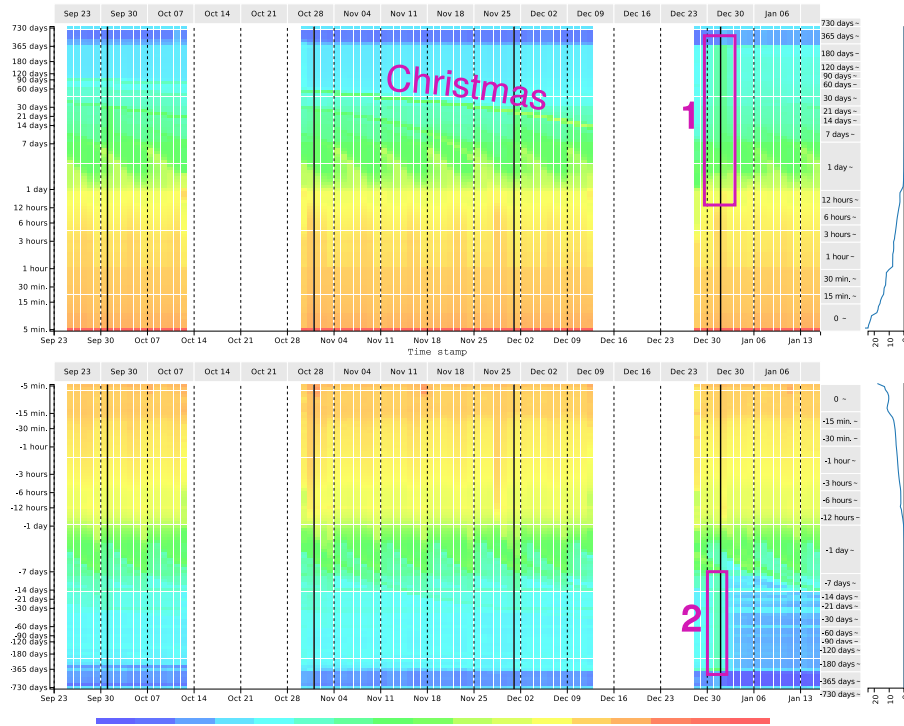
Perhaps, the most striking pattern one can observe is the relatively high similarity in the temporal attention related to the future and to the past. Thus, at this point, it is difficult to say that Twitter users have significantly different pattern of interest in the future than in the past or vice versa. Future were only slightly more common than the ones pointing to the near past. Of course, the symmetry is not perfect. There are certain differences noticeable such as curvy lines within the time frame further than several days from the timestamps (more visible in the US dataset). In general, the similarity between the past and future views is rather striking.

Next, we describe the periodical decreasing yellow or green curves in the future and the past-related parts of the graphs. After manual investigation of their top words, we found that they represent the tendency of people to use expressions referring to the end of a week. This is also confirmed by the fact that such lines last for a week when measured on abscissa and that they are vertically symmetrical based on dashed black lines indicating the change of the week. We note that in both the datasets the periodical curves are stronger in future (upper parts of the graphs) than in the past. When setting the different coloring scale (not shown in this paper) we managed also to observe similar situation for months. In addition, few curves longer than weekly or monthly curves can be noticed. As we observed they are due to tweets referring to significant temporal landmarks within our dataset such as Christmas (see the “Christmas” label in both the graphs).

Interestingly, we can notice another intriguing symmetry in the future and past-related areas of the graphs centered around the start of the New Year (see areas indicated with labels 1 and 2). In the past-related part (area with label 2 in both the graphs) we interpret it as the process of *re-calling things that happened throughout the last year*. This



(a) Snapshot of Japanese dataset available online [43]



(b) Snapshot of US dataset available online [44]

Figure 3: Heat maps of relative mentions of future and past along time (best viewed in color)

phenomenon starts to surface several days (and is especially evident just one day) before the end of the year⁴. We can observe the continuous green color of the cells on the column representing Dec 31, 2013. Note the rather drastic decrease in the past mentions referring to the time older than one day during the first days of 2014 what emphasizes that people tend to “forget” or, at least, to reminiscence less the events and things of the previous year, right after the new year has come. When looking on the part of the graphs representing the future (label 1 in both the graphs) we conclude that the process of *making plans for the next year* starts also few days before the end of the year but is the strongest on the first day of the next year (Jan 1, 2014) and then lasts for several subsequent days. Same as in the past part of the graphs on Dec 31, 2013, the column representing the first day of the new year in the future area (see area with label 1) has the pattern of green cells spanning almost the whole year (cells on the column of that day and also several next days). It is likely due to users discussing plans and wishes for their future activities to be done in various time periods of the new year.

4.4 Variation of Patterns in Temporal Mentions over Time

The next observation is the apparent stability of temporal attention both for the future and for the past, if one neglects the effect of major calendar events such as the Christmas and the New Year’s Eve. It means that the time horizons of users are essentially unchanging over the time period of our datasets.

When zooming in graphs by changing the upper limit of *intensity* used by the color scale, we were able to find some collective expectations of future events, such as announcement of Tokyo city as a host of Olympics 2020 that was made public on Sunday of Sep 8, 2013 (area with label 3 in the Japanese graph). In the snapshot in Figure 3(a), this is not so obvious due to cell normalization effect, described latter in Section 5.3. To observe this effect the cell normalization can be deactivated in the options. Another collective expectation can be found on top cell of Jan 13, 2014 which is due to the *Coming of Age Day* in Japan⁵. The celebrating students are expected to graduate from their schools after two years, hence many tweets on that day refer to the time two years later.

4.5 Topics of Expectations and Memories

We hypothesize that, on average, different expressions should be found when someone refers to events or actions in very near time (e.g., expecting breakfast or remembering lunch meeting at the end of the day) as opposed to happenings associated with more distant time (e.g., planning a vacation trip well ahead or reminiscing it some time later).

We examine then key terms related to different temporal scopes. Our system allows seeing the representative words used in tweets that refer to given time frames. These appear in a pop-up window upon pointing mouse pointer over gray buttons on the right-hand side of both future-related and past-related plots. The calculation of word scores is detailed in Section 5. Similarly, for exploration purposes it is possible to see representative words chosen from all the time-refereed

⁴Unfortunately, due to data loss this pattern cannot be seen in its entirety

⁵http://en.wikipedia.org/wiki/Coming_of_Age_Day

tweets on a given week by selecting the gray buttons on the top (bottom) of the future-related (past-related) plots.

Table 1 shows the top representative words for tweets mentioning selected time frames both in the past and the future in the case of the Japanese dataset (we skip the US dataset due to space constraints). We can observe that there are few expressions about the weather conditions for the near future (6-12h ahead), while the far future is dominated by the information about IOC decision to hold 2020 Olympics in Japan. The near past (6-12h ago) is about common things people have just done (ate lunch, studied, returned home, etc.), while the distant past expressions associate with tweets containing recollections and reminiscences.

future over 730 days	Olympics, marriage, live, good luck, Olympics (kanji), Japan, hold, words, decide, imagination
future 6-12 hours	weather, good morning, announcement, forecast, sleep, lunch, information, caution, exam, school
past 6-12 hours	weather, lunch, arrive, practice, sleep, study, announcement, good job (well done), leisure, tired
past over 730 days	nostalgic, change, recall, in those days, I, old, young, remember, know, period

Table 1: Top words in tweets related to near/distant future/past

Then, in Table 2 (future) and Table 3 (past) we show top words from aggregated predictions and memories for two different weeks, one in summer (Aug 19-25) and one at the change of the year (Dec 30-Jan 5). We see that for example, in the summer period, that the tweets about future often concern weather predictions, plans for summer vacations or about sports matches. For the past related words in the summer we observe memories of Obon vacation period (Aug 13-16) and some fireworks events. There are also some words about the earthquakes that happened at that time. Looking at the top words for the predictions and memories at the time of the year change we observe many customary expressions for greeting the new year, remembering the old year and visiting shrines to celebrate the new year and to thank for the good things in the old year.

Aug 19-25	high temperature, °C, caution, thunders, information, sports match, summer vacation, hot, beach, rain
Dec 30-Jan 5	start of the new year, celebrate, (the new year) begins, first visit of the year to a shrine (typical custom in Japan), new year, welcome the new year and forget the old, shrine, kind regards, new year (kanji), welcome

Table 2: Top words in future related tweets in Aug 19-25 and Dec 30-Jan 5 in the Japanese dataset

4.6 Comparison of Time Mentions and Timestamps

We investigate here to which days tweets written on a particular day refer to. Intuitively, the majority of tweets written on a given day should refer to the actions during that day. We test this hypothesis using the visualization shown in the graphs in Figure 4. Figure 4(a) refers to the Japanese and Figure 4(b) to the US dataset. The horizontal axis in the graphs denotes tweet timestamps, while the horizontal one denotes tweet mentions, both represented in days. Note

Aug 19-25	obon (holiday period in Japan on Aug 12-16), scary, fireworks, km, happen, deep, earthquake, hot, beach, breaking news
Dec 30-Jan 5	greetings, thank you (for everything this past year), farewell to the old year, new year, old year, (the new year) begins, the last year, congratulations, various things (of this year), year-crossing

Table 3: Top words in past related tweets in Aug 19-25 and Dec 30-Jan 5 in the Japanese dataset

that these matrix-like graphs have been computed based on all *time expressions* in our datasets, hence, not only on *time expressions* of daily granularity. We perform a Jaccard-style normalization, further described in Section 5.2, that is mix of column-wise and row-wise normalization, allowing us to avoid the effect of a non-uniform distribution of the tweet *timestamps* and *time mentions*. Below the matrix graphs of Figure 4 we also display for reference the cumulative “popularity” of days mentioned in tweets, and on the right-hand side we show the frequency of tweets.

When looking at the graphs, indeed, we observe the strong impact of timestamp on the time mentions in tweets in the form of pronounced diagonal (red and orange colors) of a width of about 3 cells. This means that many tweets refer to the same day as their timestamp or just 1 day before or after it. In addition, we observe light blue periodical rectangles which fit inside the week cells. A week cell is bounded by dashed line. Keeping in mind that all the cells above the diagonal represent tweets about future and below the diagonal tweets about the past, we conclude that these periodical rectangles portray expectations or memories within the same week as the one of the tweet timestamp. This again confirms the usage of *temporal landmarks* in natural language. Users tend to relate to the future and the past in a way which is bounded by the start and the end of the current week. Also, as seen in the graphs, there seems to be more attention put on the future than on the past by the users.

Lastly, the three light blue vertical lines in the Japanese and the US datasets correspond to the long awaited major calendar events: Obon holiday (Aug 13-16), Halloween, Christmas and Halloween, Thanksgiving, Christmas, respectively.

5. TECHNICAL ASPECTS

The system has been implemented in SCALA programming language on the server side and using D3 graphical library for enabling the web interface. The Japanese morphological analyzer Kuromoji⁶ was used for text processing tasks and Stanford CoreNLP [22] was used for text parsing and for detecting time expressions in English.

In the following of this section we present the definitions of the main aspects of the system. Formally, let T be the set of all the tweets in the dataset, $t \in T$ be a tweet characterized by a bag of words W_t . Visualizations are heat maps discretized in cells. Let also CEL be the set of all the cells in a heat map and $C \in CEL$ be one cell in the heat map. Finally, $C_{i,j}$ is a cell in row i and column j of a heat map with ROW and COL are respectively the number of lines and columns of a heat map.

5.1 Probability Mass Function

According to time mentions in a tweet t the cell $C_{i,j}$ is associated with the probability $P(C_{i,j}|t)$ such that $t \in C_{i,j}$ when the time expression is a point in time, $P(C_{i,j}|t) \rightarrow \{0, 1\}$. Otherwise, for a time period, $P(C_{i,j}|t) \rightarrow \{w \in \mathbb{R} : 0 \leq w \leq 1\}$ following a uniform distribution of the weight w of the tweet over the time period discretized according to the granularity described in Section 3.2. Intuitively, $P(C_{i,j}|t)$ is the amount of the probability mass that the tweet assigns to the cell. Hence, the sum of the weight of a tweet t in all the cells of a heat map, must be 1 (i.e. $\sum_{C \in CEL} P(t|C) = 1$).

Let $I(C_{i,j}) = \sum_{t \in T} P(C_{i,j}|t)$ be the *intensity* I of a cell $C_{i,j}$. This *intensity* value controls the color displayed in Figure 3 and 4.

5.2 Normalization

When looking at the blue curves at the bottom of graphs in Figure 3 and right of graphs in Figure 4, one can remark that the number of tweets issued each day is significantly varying. Hence, the system provides three possible normalizations to remove the effect of peak time. First, the row-wise normalization allows comparing the cells on the same row:

$$\text{Score}(r, c) = \frac{I(C_{r,c})}{\sum_{i \in ROW} I(C_{i,c})}$$

Second, the column-wise normalization allows comparing the cells on the same column: $\text{Score}(r, c) = \frac{I(C_{r,c})}{\sum_{j \in COL} I(C_{r,j})}$

The third one, a Jaccard coefficient normalization offers a way to “combine” both normalizations on the same visualization:

$$\text{Score}(r, c) = \frac{I(C_{r,c})}{\sum_{i \in ROW} I(C_{i,c}) + \sum_{j \in COL} I(C_{r,j}) - I(C_{r,c})}$$

The Jaccard coefficient is a ratio to measure the similarity between two sets. It is computed as the ratio between the intersection and the union of the two sets. In this application, for a given cell $C_{r,c}$ the two sets are, the set of all the tweets in the cells of row r , and the set of all the tweets in the cells of column c . Hence, it measures how much the cell is representative for the two values r and c on the axis x and y of the heat map.

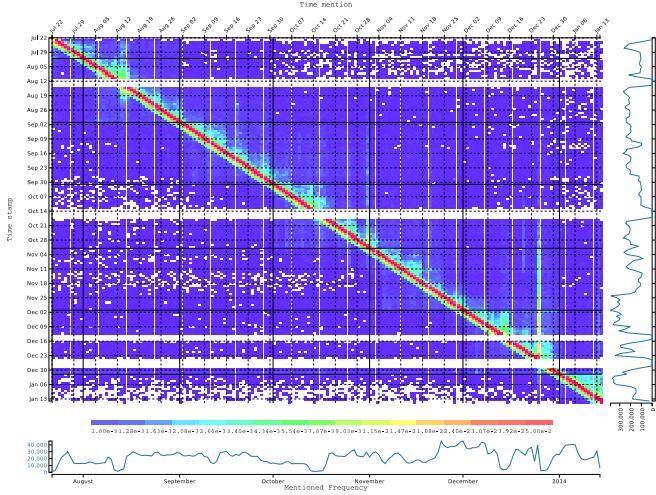
These normalizations are parameters, allowing the user to change the visualizations to spot different kind of patterns.

5.3 Normalization of Logarithmic Scale

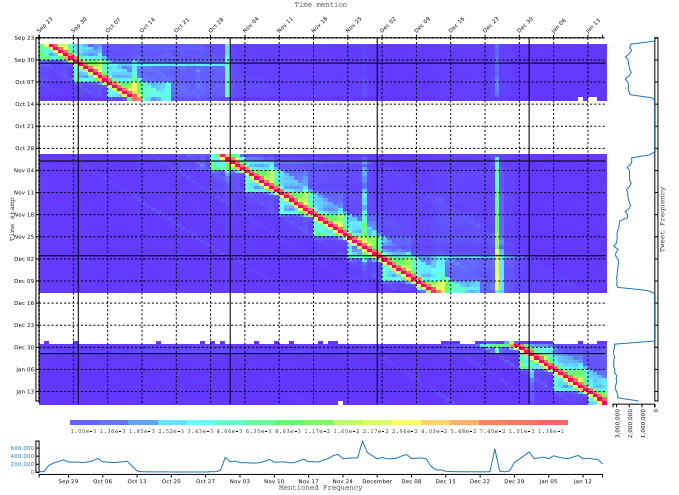
Here we address the problem of time visualization on a logarithmic scale. To illustrate the problem let’s consider an extreme case where all the tweets at a given timestamp ts have the time expression “this month”. In such case we would expect that the column ts would have cells spanning from beginning of this month until the end displayed in one uniform color and all the rest of cells in white. However, due to the ordinate axis drawn in a logarithmic scale the cells farther from the origin (more distant past or future) represent a larger time range than cells closer to the origin (near past or future). They will thus aggregate much more tweet’s weights and, consequently, will be represented with a much stronger color, giving the false impression than the far future and past are stronger referred to than they actually are. This is because the weight of the tweets will be uniformly distributed over time covered by the expression “this month”.

Let $I(C_{i,j}) = \sum_{t \in T} P(C_{i,j}|t)$ be the *intensity* I of a cell $C_{i,j}$. Intensity represents a raw measurement, as the sum of the probability of the tweets in the dataset to affect a cell.

⁶<http://www.atilika.org/>



(a) Snapshot of Japanese dataset available online [45]



(b) Snapshot of US dataset available online [46]

Figure 4: Heat map of timestamp vs. time mention matrix (best viewed in color)

Therefore, intensity should not be used directly to represent colors. Instead, we normalize it according to the scale of the axis. The color of a cell is then quantified as the *intensity* of that cell divided by time range of that cell:

$$color(C_{i,j}) = \frac{I(C_{i,j})}{size(C_{i,j})}$$

This normalization is called cell normalization and can be deactivated in the visualization options online.

5.4 Ranking Words

For each visualization, the system provides a ranking of the words in an arbitrarily defined area on the heat map called a region. The goal is to keep a top-k list of the most characteristic words for the region. The weight of each word is based on a TF-IDF weighting scheme to represent how much the word is characteristic for the given region. Intuitively, TF-IDF assigns high scores to terms that often appear in tweets associated with a given cell, while appearing infrequently in other cells. The novelty in our approach is in the count of term frequency using probability and in the definition of document based on the two levels of granularity defined for the regions: *cell* and *segment*. They are described in the subsequent paragraphs.

Cell.

In order to find the most relevant words appearing in tweets for each cell we measure the word relevance using a modified TF-IDF score where a term is a word and a document is a set of tweets corresponding to a given cell. The score to rank a word w in a cell $C_{i,j}$ using the whole set of tweets T is:

$$\text{Score}(w, C_{i,j}, T) = \frac{\sum_{t \in T} P(C_{i,j}|t) : w \in W_t}{\sum_{t \in T} P(C_{i,j}|t)} \times \log \frac{|CEL|}{|C \in CEL : \exists t \in C : w \in W_t|}$$

Segment.

In the visualization in Figure 3, the gray cells on the top or bottom line and on the right-hand side are regarded as virtual documents that represent *segments*. For instance, the top left gray cell in Figure 3 with label “Jul 22” gathers all the cells between column July 22 to 28 on every row. The *segments* are either row-wise or column-wise; they are sets of cells that span, respectively, all columns covered by the width of the selected rows, or all rows covered by the width of the selected columns.

A second level TF-IDF score is computed to rank top words in *segments*. The content of a virtual document (segment) is the set of words W_{S_i} built from all the words of the cells in that *segment*.

Let SEG be the set of all the horizontal segments in a heat map and $S \in SEG$ be one horizontal segment in the heat map that is a set of all the tweets affecting the cells of that segment. Finally, S_i is a horizontal segment in position i in a heat map. The weight of a word w in a segment is the sum of the weights of the tweets in which the word w appears.

$$\begin{aligned} \text{Score}(w, S_i, T) = & \frac{\sum_{t \in T} P(C|t) : C \in S_i : w \in W_C}{\sum_{t \in T} P(C|t) : C \in S_i} \\ & \times \log \frac{|SEG|}{|S \in SEG : \exists t \in S : w \in W_t|} \end{aligned}$$

6. DISCUSSION

We will first list the key observations we noticed using our system. Then we enumerate several directions for future research.

6.1 Key Findings

Amount of Temporal Expressions. We have found quite many temporal expressions in tweets. They can be used for anchoring tweet content on the timeline useful for many tasks including information retrieval, summarization,

trend detection and so on.

Short Time Horizons and Supremacy of the Present. As expected, tweets exhibit clear temporal constraints. The attention to distant time periods as measured by the amount of extracted temporal references is significantly lower than the one to the near time. The present when more generally defined by yesterday, today and tomorrow, dominates the time-span of the temporal attention. This observation confirms conclusions from some related studies on news articles or on the open web [14, 15].

Similarity of Remembering and Expecting. We have noticed strong symmetry in the visual maps of temporal references to the past and to the future. The time horizons and the patterns of temporal attention are quite similar with the future taking however more attention of the users than the past.

Temporal Landmarks. As observed users tend to use temporal landmarks such as the start and end of weeks or months. The strongest landmark is the end and the start of the year.

Relative Stability of Temporal Horizons. We observed that the patterns of attention are more or less stable over time apart from few exceptions (e.g., the time periods just before and after the New Year’s Eve). That is, for different periods of collected data there is relatively little variation in the distribution of temporal attention and the lengths of temporal horizons.

Diversity of Discussed Topics over Time. We observe certain variations in words used when referring to diverse time distances. Different time segments are characterized by differing sets of activities. The day life activities dominate the present, as shown in Table 1 but some popular events can be noticed for far future and past.

Cultural Differences. Some cultural variations can be observed by comparing views of the two datasets. The most obvious ones relates to different calendar events (e.g., Thanksgiving day in the USA). Nevertheless, on average more similarities can be observed.

Lastly, we note that our study is limited to only tweets containing temporal expressions. We also analyze only half a year long time period for the Japanese dataset and 4 months long period in case of the US dataset. It would be advantageous to analyze more implicit ways in which people can refer to future or past and to compare the results over longer datasets such as ones spanning a year or even few years. We plan to work on these issues in the future.

6.2 Applications and Future Extensions

Combined spatio-temporal analysis. The addition of location mentions and location stamps to the analysis should help to more precisely discover commonly remembered or expected events. We could also investigate time-space orthogonality through the comparative analysis of the temporal and spatial attention.

Future Prediction. As mentioned before, the expressed expectations can be utilized for predicting user activities to support any pro-active, reactive or adaptive systems. By studying the visualization graphs the feature engineering for classifiers can be improved for any activity prediction systems. For example, to predict user future mobility (e.g., holiday travel) we would expect the references to the travel to appear

long before the departure date. Thus, any evidence data hinting travel plans to be used for creating classifier features should be searched within certain time ahead of the start of the travel with corresponding granularity. Also, the typical variations in temporal attention that we observe in the graphs of Figure 3 suggest particular division of time into units from where features can be extracted. Instead of uniform time division one could use the one adopted to the significant changes occurring in temporal attention as portrayed by our visualizations. Obviously, different time divisions will result in different classification features collected and should impact prediction accuracy.

Recommendation. Based on captured user memories or expectations future systems could better suggest particular activities, products, services or locations to visit.

Sentiment analysis. Humans tend to retain positive memories more than negative ones (*rosy retrospection* [27]) and favor positive future views than negative ones (*valence effect* [35]). It is then interesting to perform sentiment analysis to confirm any potential sentiment biases.

Verifying expectations and predictions. It should be possible to evaluate the predictability of future expressions and credibility of future plans. For example, for a user tweeting about travelling to Hawaii in the next month, we might search for any tweets by this user with GPS location within Hawaii to confirm the fact of visiting the islands one month later.

Social computing. As future work we plan to extend our system to accept any datasets for visualization. This should be useful for social scientists who lack efficient means for making sense of their data. Also, a more precise name entity recognizer [33] designed to deal with the noisy nature of twitter will replace Stanford CoreNLP.

7. CONCLUSIONS

In this paper we studied the dynamics of expecting and remembering processes characteristic to social media users based on data collected over nearly half a year long time-span for the Japanese and 4 months for the US datasets. Our focus is a large population of microblogging users who continuously share online information about their daily lives and things that matter to them. We have demonstrated visual framework to portray particular ways in which users refer to the past and to the future. Our methodological approach allows observing the scope of temporal attention of users, temporal horizons of their perspectives, the typical expressions and topics associated with references to given time frames (e.g., near, distant past or near, distant future) and other related features. In the future we plan to work on the previously listed directions.

8. ACKNOWLEDGMENTS

This work was supported by the Strategic Information and Communications R&D Promotion Programme (SCOPE) of the Ministry of Internal Affairs and Communications of Japan, the JSPS KAKENHI Grant Number 26280042 and the JST research promotion program Sakigake: “Analyzing Collective Memory and Developing Methods for Knowledge Extraction from Historical Documents”.

References

- [1] Fabian Abel, Claudia Hauff, Geert-Jan Houben, Richard Stronkman, and Ke Tao. “Twitcident: Fighting Fire with Information from Social Web Streams”. In: *WWW Companion*. 2012.
- [2] Omar Alonso, Ricardo Baeza-yates, Jannik Strötgen, and Michael Gertz. “Temporal Information Retrieval: Challenges and Opportunities”. In: *TempWeb*. 2011.
- [3] Sebastien Ardon, Amitabha Bagchi, Anirban Mahanti, Amit Ruhela, Aaditeshwar Seth, Rudra Mohan Tripathy, and Sipat Triukose. “Spatio-temporal and Events Based Analysis of Topic Popularity in Twitter”. In: *CIKM*. 2013.
- [4] Sitaram Asur and Bernardo A. Huberman. “Predicting the Future with Social Media”. In: *WI-IAT*. Vol. 1. Aug. 2010.
- [5] Klaus Berberich, Srikanta Bedathur, Omar Alonso, and Gerhard Weikum. “A Language Modeling Approach for Temporal Information Needs”. In: *ECIR*. 2010.
- [6] Ricardo Campos, Gaël Dias, Alípio M. Jorge, and Adam Jatowt. “Survey of Temporal Information Retrieval and Related Applications”. In: *ACM Comput. Surv.* 47.2 (2014), 15:1–15:41.
- [7] Yan Chen, Hadi Amiri, Zhoujun Li, and Tat-Seng Chua. “Emerging Topic Detection for Organizations from Microblogs”. In: *SIGIR*. 2013.
- [8] Au Yeung Ching-man and Adam Jatowt. “Studying How the Past is Remembered: Towards Computational History Through Large Scale Text Mining”. In: *CIKM*. 2011.
- [9] Hermann Ebbinghaus. “Über das Gedchtnis. Untersuchungen zur experimentellen Psychologie.” In: (1885).
- [10] Hermann Ebbinghaus. *Memory: A Contribution to Experimental Psychology*. Columbia University. Teachers College. Educational reprints. no. 3. 1913.
- [11] Maurice Halbwachs. *La mémoire collective*. Les Presses universitaires de France. 1950.
- [12] Christoph Hoerl and Teresa McCormack. *Time and memory : issues in philosophy and psychology*. No.1. 2001.
- [13] R. Jacoby. *Social Amnesia: A Critique of Contemporary Psychology*. 1997.
- [14] Adam Jatowt and Ching-man Au Yeung. “Extracting collective expectations about the future from large text collections”. In: *CIKM*. 2011.
- [15] Adam Jatowt, Hideki Kawai, Kensuke Kanazawa, Katsumi Tanaka, Kazuo Kunieda, and Keiji Yamada. “Analyzing Collective View of Future, Time-referenced Events on the Web”. In: *WWW*. 2010.
- [16] Nattiya Kanhabua, Klaus Berberich, and Kjetil Nørvg. “Learning to Select a Time-aware Retrieval Model”. In: *SIGIR*. 2012.
- [17] Nattiya Kanhabua and Wolfgang Nejdl. “Understanding the Diversity of Tweets in the Time of Outbreaks”. In: *WWW*. 2013.
- [18] Nattiya Kanhabua, Tu Ngoc Nguyen, and Claudia Niederee. “What triggers human remembering of events? A large-scale analysis of catalysts for collective memory in Wikipedia”. In: *JCDL*. Sept. 2014.
- [19] Nattiya Kanhabua and Kjetil Nørvg. “Determining Time of Queries for Re-ranking Search Results”. In: *ECDL*. 2010.
- [20] David Lazer, Alex (Sandy) Pentland, Lada Adamic, Sinan Aral, Albert Laszlo Barabasi, Devon Brewer, et al. “Life in the network: the coming age of computational social science”. In: *Science* 323.5915 (Feb. 6, 2009), pp. 721–723.
- [21] Chenliang Li, Aixin Sun, and Anwitaman Datta. “Twevent: Segment-based Event Detection from Tweets”. In: *CIKM*. 2012.
- [22] Christopher D. Manning, Mihai Surdeanu, John Bauer, Jenny Finkel, Steven J. Bethard, and David McClosky. “The Stanford CoreNLP Natural Language Processing Toolkit”. In: *ACL*. 2014.
- [23] Adam Marcus, Michael S. Bernstein, Osama Badar, David R. Karger, Samuel Madden, and Robert C. Miller. “Twitinfo: Aggregating and Visualizing Microblogs for Event Exploration”. In: *CHI*. 2011.
- [24] Andrew J. McMinn, Daniel Tsvetkov, Tsvetan Yordanov, Andrew Patterson, Rrobi Szk, Jesus A. Rodriguez Perez, and Joemon M. Jose. “An Interactive Interface for Visualizing Events on Twitter”. In: *SIGIR*. 2014.
- [25] M. J. Metzger. “Making sense of credibility on the Web: Models for evaluating online information and recommendations for future research”. In: *JASIST* 58.13 (2007), pp. 2078–2091.
- [26] Donald Metzler, Rosie Jones, Fuchun Peng, and Ruiqiang Zhang. “Improving Search Relevance for Implicitly Temporal Queries”. In: *SIGIR '09*. 2009.
- [27] Terence R. Mitchell, Leigh Thompson, Erika Peterson, and Randy Cronk. “Temporal Adjustments in the Evaluation of Events: The “Rosy View””. In: *Journal of Experimental Social Psychology* 33.4 (1997), pp. 421–448.
- [28] Fred Morstatter, Shamanth Kumar, Huan Liu, and Ross Maciejewski. “Understanding Twitter Data with TweetXplorer”. In: *KDD*. 2013.
- [29] Mashaal Musleh. “Spatio-temporal Visual Analysis for Event-specific Tweets”. In: *SIGMOD*. 2014.
- [30] Sérgio Nunes, Cristina Ribeiro, and Gabriel David. “Use of Temporal Expressions in Web Search”. In: *ECIR*. 2008.
- [31] Kira Radinsky, Sagie Davidovich, and Shaul Markovitch. “Learning Causality for News Events Prediction”. In: *WWW*. 2012.
- [32] Kira Radinsky and Eric Horvitz. “Mining the Web to Predict Future Events”. In: *WSDM*. 2013.
- [33] Alan Ritter, Sam Clark, Mausam, and Oren Etzioni. “Named Entity Recognition in Tweets: An Experimental Study”. In: *EMNLP*. 2011.
- [34] Alan Ritter, Mausam, Oren Etzioni, and Sam Clark. “Open Domain Event Extraction from Twitter”. In: *KDD*. 2012.
- [35] David Rosenhan and Samuel Messick. “Affect and expectation”. In: *Journal of Personality and Social Psychology* 3.1 (1966), pp. 38–44.
- [36] Takeshi Sakaki, Makoto Okazaki, and Yutaka Matsuo. “Earthquake shakes Twitter users: real-time event detection by social sensors”. In: *WWW*. 2010.
- [37] George Valkanias and Dimitrios Gunopulos. “How the Live Web Feels About Events”. In: *CIKM*. 2013.
- [38] W. Weerkamp and M. de Rijke. “Activity Prediction: A Twitter-based Exploration”. In: *TAIA*. 2012.
- [39] Jianshu Weng and Bu-Sung Lee. “Event Detection in Twitter.” In: *ICWSM*. 2011.
- [40] Jaewon Yang and Jure Leskovec. “Patterns of Temporal Variation in Online Media”. In: *WSDM*. 2011.

Online Resources

- [41] *Portal to visualizations of the Japanese dataset*. URL: <http://scope13.cse.kyoto-su.ac.jp/JP4/chartlist.html>.
- [42] *Portal to visualizations of the USA dataset*. URL: <http://scope13.cse.kyoto-su.ac.jp/US4/chartlist.html>.
- [43] *Time difference for Japanese dataset*. URL: <http://scope13.cse.kyoto-su.ac.jp/JP4/timediff.html>.
- [44] *Time difference for US dataset*. URL: <http://scope13.cse.kyoto-su.ac.jp/US4/timediff.html>.
- [45] *Time transition matrix for Japanese dataset*. URL: <http://scope13.cse.kyoto-su.ac.jp/JP4/timematrix.html>.
- [46] *Time transition matrix for US dataset*. URL: <http://scope13.cse.kyoto-su.ac.jp/US4/timematrix.html>.