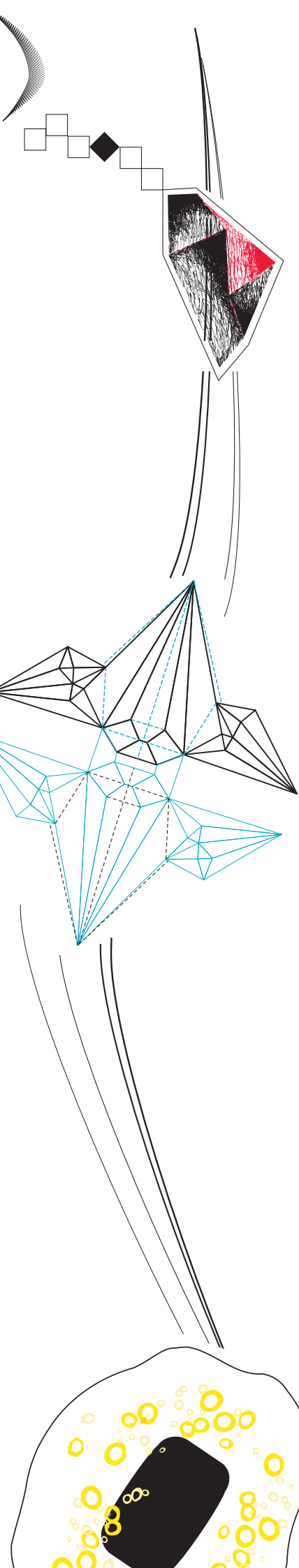


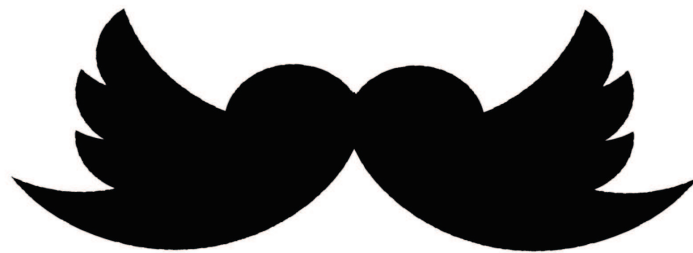
August 30th, 2016



Effects of Twitter Mentions on the Movember Campaign

Thesis Applied Mathematics
Stochastic Operations Research

Mike Visser



Supervised by:
Dr. N. Litvak
Dr. ir. T.A. van den Broek

Faculty of Electrical Engineering, Mathematics and Computer Science

UNIVERSITY OF TWENTE.

Effects of Twitter Mentions on the Movember Campaign

Mike Visser
Student number: s1221051

Supervised by:
Dr. N. Litvak
Dr. ir. T.A. van den Broek

August 30th, 2016

Table of Contents

1	Introduction	5
1.1	Background	5
1.2	Research goals	5
2	Data	7
2.1	Data sets	7
2.2	Data statistics and visualization	7
3	Centralities in the Twitter network of mentions	14
3.1	Definition of centralities	14
3.2	Comparison of centralities by using Kendall's weighted τ	15
4	Modeling Motivation	18
4.1	Notations for motivation models	18
4.2	Pure Growth Model: susceptibility	18
4.3	Pure Growth Model: susceptibility and reactions	19
4.4	Motivational decay	23
5	The Satiation-Deprivation Model: Construction and Analytical Results	25
5.1	Model requirements	25
5.2	Theoretical development of the SD Model	26
5.3	The SD Model as a stochastic model	27
6	The Satiation-Deprivation Model: Sequential Approach and Extensions	30
6.1	Notations for the sequential approach	30
6.2	The sequential approach for $ C' = 1$	30
6.3	The equivalence of the SD Recursions and the sequential approach for $ C' = 1$	31
6.4	The sequential approach for $ C' \geq 1$	32

6.5	Extensions	33
7	Donation potential model	35
8	Matching procedure	36
9	Numerical results for Twitter and Movember data	39
9.1	Modeling ϕ for Twitter users	39
9.2	Visual results for Satiation-Deprivation Model and donation potential	39
10	Statistical justification of Satiation-Deprivation Model	43
10.1	Goodness-of-Fit: The Kolmogorov-Smirnov Test	43
10.2	Co-occurrence of mentions and donations	44
10.3	Co-occurrence of high motivation levels and donations	45
10.4	The effect of the order of mentioners on donation occurrence	46
11	Discrete-Time Strategic Mentioning	48
11.1	State space, action set and value function	49
11.2	Brute-force algorithm for optimal strategies	52
11.3	Results	52
12	Continuous-Time Strategic Mentioning	55
12.1	Optimal mentioning for $M = 1$	56
12.2	Optimal mentioning for $M = 2$	57
12.3	Optimal mentioning for general M	59
12.3.1	The case $\phi = 1$	59
12.3.2	The case $\phi < 1$: deriving fixed-point equations	60
12.3.3	Numerical approach	65
12.3.4	Optimality of fixed-point equation	65
12.4	Effect of parameters on optimal average motivation values and optimal strategies .	66

12.5 Sensitivity of value to suboptimal strategies	68
13 Conclusions	70
14 Recommendations for Movember	71
15 Future research	72
16 Acknowledgments	73
A Kendall's weighted τ	75
B Critical values for the Two-sample Kolmogorov-Smirnov test	76
C Graphs	77

1 Introduction

1.1 Background

Movember The Movember Foundation concerns itself with four of the biggest worldwide issues concerning the health of men: prostate cancer, testicular cancer, poor mental health, and physical inactivity. Movember uses its money to create awareness and to fund research worldwide. It is named for its well-known activity in the month of November, when men around the globe let their moustaches grow to promote male health. By cultivating a moustache and finding sponsors, participants in the Movember campaign spread awareness for the cause and raise money for the foundation.

A person who participates in the Movember campaign, a MoBro or a MoSista, thus acts as a *fundraiser*. He or she can create a profile, a MoSpace, on the Movember website. The MoSpace has the options for fundraisers to place a picture of themselves, state if they participate in a team with others, tell what kind of activities they do and give a description of why they participate. Fundraisers also have a Movember ID, a unique identification number. Friends and families of the fundraisers can sponsor the moustache or other activities that the fundraiser takes part in. It is important to keep in mind that donations associated with a fundraiser (or a Movember ID) therefore do not concern donations made by the fundraiser himself, but donations made by his or her sponsors.

Twitter Fundraisers can use any kind of social media in relation with the Movember campaign. Platforms like Twitter are excellent channels for promotion and participation in the activism. Fundraisers may use their profile to showcase their moustaches, spread information and find sponsors.

In 2014 the University of Twente was awarded a Twitter DataGrant, by writing a winning research proposal. They were among the six winners out of 1300 proposals. The data grant involves access to historical data from Twitter (in the form of Tweets) and is used by the University of Twente to research effectiveness of cancer-related campaigns.

Since the data grant was awarded, aspects of various campaigns have been researched. For the Movember campaign, the influence of prosocial norms and mention networks on donations worldwide was investigated, see [6]. Twitter data from 24 countries was analyzed. As most Tweets do not carry geolocation information, the data could be related to the country of origin only because of a country classifier built in [7]. We also use the country-classified data for this research.

For us the most interesting aspect of the Movember campaign on Twitter is when users are mentioning each other. A *mention* is made from USER1 to USER2 if a Tweet from USER1 contains the substring "@USER2". It can be regarded as a social 'boost' from USER1 to keep USER2 committed to the campaign, so as to USER2 to find himself sponsors, and consequently to receive more donations benefitting the Movember Foundation. We shall frequently use the terms *mentioner* and *mentione* in this report, where the mentioner is the writer of the Tweet containing a mention, and the mentionee is the receiver.

1.2 Research goals

Receiving a mention on Twitter can lead to an increase in the prosocial motivation of the mentionee. This kind of motivation has been widely researched in psychology and sociology and compared to other kinds such as intrinsic motivation, see [4]. Increases in prosocial motivation

can lead to a better performance which may occur in the form of task achievement, but also in the form of more and higher donations. The first research question of our research is therefore:

Research question 1: *How do mentions on Twitter affect donations to Movember?*

The idea is that receiving mentions leads to an increase in motivation and a higher motivation increases the activism of an individual. This can lead to more donations from his or her sponsors. A subgoal is therefore to model the height of the motivation by using mentions as an input. We aim to relate the height of motivation to donation occurrence.

Devising a measure for motivation is not a straightforward procedure. Even when a numerical value is associated with this psychological construct, it can be hard to determine the cause of increases and decreases in motivation. The causes for people to become motivated to behave or act in a certain way differ from culture to culture, from setting to setting, and from person to person.

In this research we shall not concern ourselves with capturing the exact nature of the relation between a mention and motivation. Instead, we introduce the notion of *motivation level* as a mathematical construction that is rooted in social exchange theory. As we shall only utilize two intuitive concepts from this theory, we shall not devote time to the research it, but for an overview we refer to [3]. After the construction of a motivation level, we shall use it as an intermediary construct to model the relationship between the observable mentions and donations. We shall show that there exists a positive relationship between height of motivation level and donation occurrence.

After establishing these relations, we aim to find mentioning strategies that maximize the mathematical motivation level. Under the assumption that the positive relationship is causal in nature, the optimal mentioning strategy then also maximizes donation occurrence. The second research question that we aim to answer is:

Research question 2: *What mentioning strategy should Twitter user i adopt to maximize the average motivation level of user j ?*

2 Data

2.1 Data sets

To answer the research question, we make use of two sets of data. The first consists of day-by-day country-classified networks of Twitter mentions, where a pair of users is included if one of them mentions the other in relation with the Movember campaign on that specific day. The model for motivation levels will be based on these mention networks. The other set consists of individual donations made to the Movember Foundation, which also contains time information.

Tweets From Twitter we received all cancer-related Tweets from May 2014 until January 2015. Then we regarded only Tweets that contained ‘Movember’, capitalized or not. By use of a naive Bayesian model, Tweets were country-identified and grouped as in [7]. For each country we obtain a set of Tweets that with high probability originate from this country. Among the country sets of Tweets were “The Netherlands” and “Sweden”; we shall extensively use the data from these networks for experiments in this research. We also received ‘Movember’-containing Tweets from 2013, which were country-classified for 24 countries. These latter Tweets are used to a smaller extent.

Mention graphs For each day we aggregated all Tweets and drew a graph where arc (u, v) was present if user u mentioned user v on this day. The graphs were stored in .graphml-files. We shall refer to them as mention graphs.

Donation data The Movember Foundation gave us access to an Excel sheet containing all donations of 2014 in The Netherlands and Sweden. Each donation carries a date, the Movember ID of the fundraiser and the donated amount of money in euros. Donation data has previously been used in [6] in an aggregated form.

Movember profiles For this research, the Movember Foundation provided us with the Movember IDs, names, team names and profile pictures (if present) from users in The Netherlands and Sweden. We used this data because many Movember profiles had been taken offline, so that a large quantity of information was not publicly available anymore. The data contains many hyperlinks to profile pictures that shall also come into use.

2.2 Data statistics and visualization

These data statistics rely on 2014 data where possible, but at times make use of Twitter data from 2013, which we had available earlier and for which we have more countries available. For 2014 we can look at The Netherlands, Sweden, the United Kingdom, the United States and Australia. Together with Djoerd Hiemstra and Han van der Veen we obtained the country-classified version of Tweets. The first two countries were requested by myself, the latter by Anna Priante who focuses on English-written Tweets.

Number of users For all countries we count the number of distinct Twitter mentioners and mentionees over the whole period in 2013, see Table 1. These are the users that obtain or give a

Country	AUS	BE	BR	CH	CZ	DE	DK	ES	FI	FR	HK	IE
# Users	16297	2575	8749	797	961	4135	1989	9989	3275	9335	832	10420
Country	IT	JP	MX	NL	NO	NZ	PT	S	SA	SI	UK	US
# Users	3168	1412	2955	12847	2016	2975	364	3082	11231	2074	392589	210321

Table 1: Number of Twitter users per country in 2013

mention with regard to Movember. We cannot do this for 2014 because the data has not all been country-classified. However, this table can give an impression of how activity around Movember is spread over the countries.

Number of mentions For four different countries we plot the number of mentions on a day in 2013 as a histogram, see Figures 1 through 4. We see that November 1st is a very popular mentioning date, as we would expect. What is less expected is the peak that we find at December 2nd in Sweden, the United Kingdom and the United States. It appears that this is a special day just after the campaign finished. One hypothesis is that on this date it has become clear how much money was raised during the Movember campaign, exciting users on Twitter to write to each other about it. In The Netherlands we can also see a slight rise in mentions just after the campaign finished.

For 2014, we have the mention graphs for the Netherlands and so we can make a histogram for this country, see Figure 5. Again a small rise in mentions can be seen in the beginning of December.

Donations For the Netherlands and Sweden we have donation data with timestamps and MovemberID of the fundraiser for the period from September 9th 2014 until May 6th 2015. We look at a few characteristics in the data set, specifically for The Netherlands. We can see in Figure 6 that the higher donations occur relatively rare compared to lower donations, and that people tend to give rounded amounts of money (like 1000, 1500, 2000).

We look at how often fundraisers receive donations. There is an exponential decay visible in the histogram of Figure 8, and a perhaps surprising amount of fundraisers obtaining many donations, some of them even around fifty. It seems that the number of times that a person has received a donation does not really influence the amount that he gets, as is evidenced by Figure 9. The graph of the median does not show large fluctuations as long as there twenty or more fundraisers. Of course for $n > 130$ the median is fully defined by the amount that the single fundraiser receives and it is thus not strange that here the median starts to fluctuate tremendously. Striking is the stability among the median heights of donations, that is, its insensitivity to the number of preceding donations.

Mentioner-mentonee structures For this part we have used Twitter data from The Netherlands in 2014.

It is good to get an overview of the mentioner/mentonee structure. We look at how many Twitter users are mentioning someone and how many mentonees are posting tweets on Movember themselves. We also see whether mentonees are more likely to mention anyone themselves. We include a figure displaying how often users get mentioned, see Figure 10.

We only consider users that post a tweet on Movember and found that 10.9% of all these Twitter users are mentioning someone else. The next natural question was: 'Do mentonees post tweets

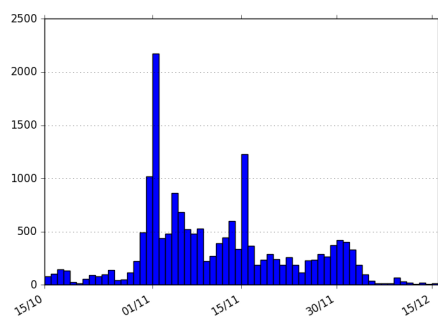


Figure 1: Number of mentions in The Netherlands, 2013

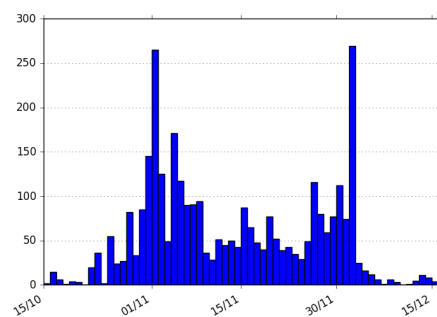


Figure 2: Number of mentions in Sweden, 2013

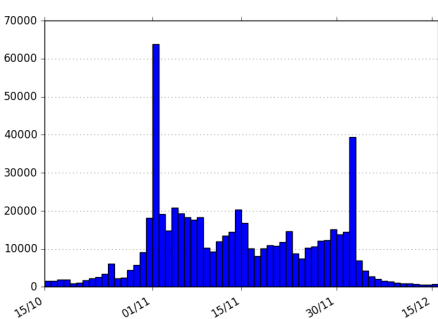


Figure 3: Number of mentions in United Kingdom, 2013

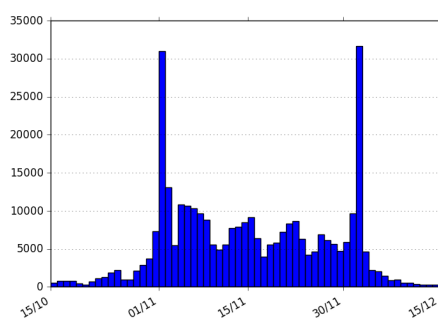


Figure 4: Number of mentions in United States, 2013

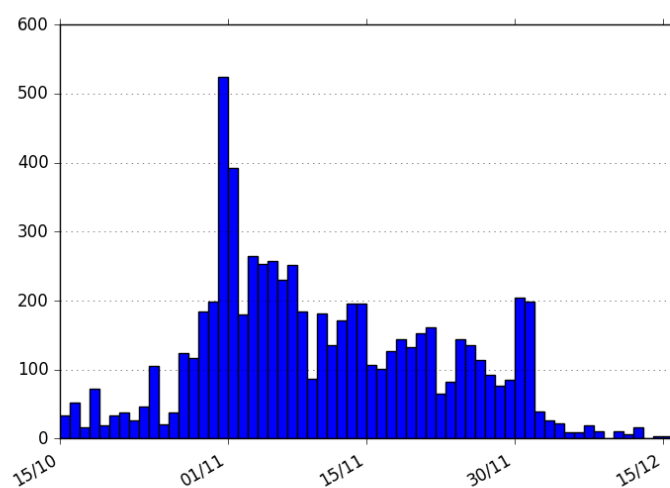


Figure 5: Number of mentions in The Netherlands, 2014

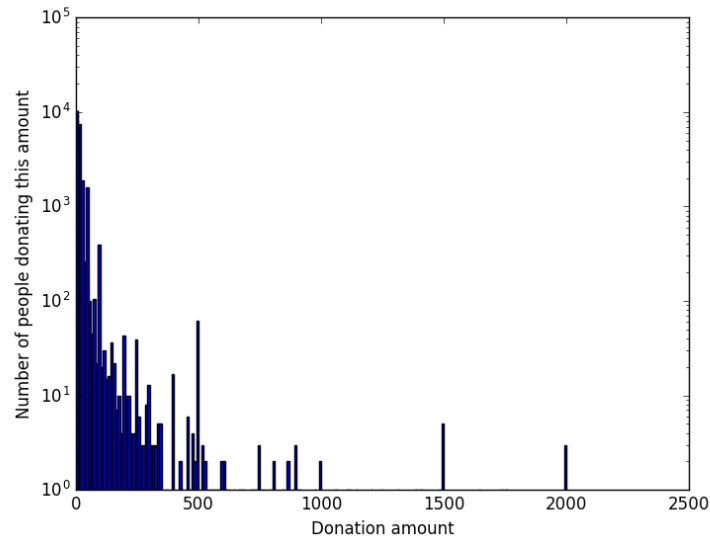


Figure 6: Histogram of donation height, Netherlands 2014

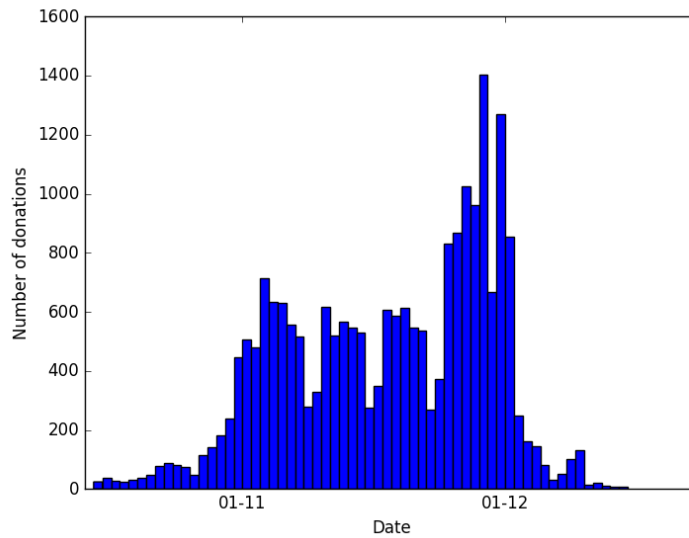


Figure 7: Number of donations over time, Netherlands 2014

about Movember themselves?’ We made a LIST1 of users that posted tweets about Movember, and a LIST2 of users that got mentioned (in relation with Movember). LIST1 contained 20 700 unique users and LIST2 contained 2 407 unique users. The overlap between the lists consists of 734 unique users. From this we gather that only $(734/2407=)30.5\%$ of mentioned Twitter users actually post something themselves concerning Movember. At the same time we can see that only $(734/20700=)3.5\%$ of users posting tweets on Movember was mentioned by someone in relation to the campaign.

The next question is how many mentionees also act as a mentioner in the networks. By making lists for the mentioners and the mentionees (as above) we found that 19.2% of mentionees mention

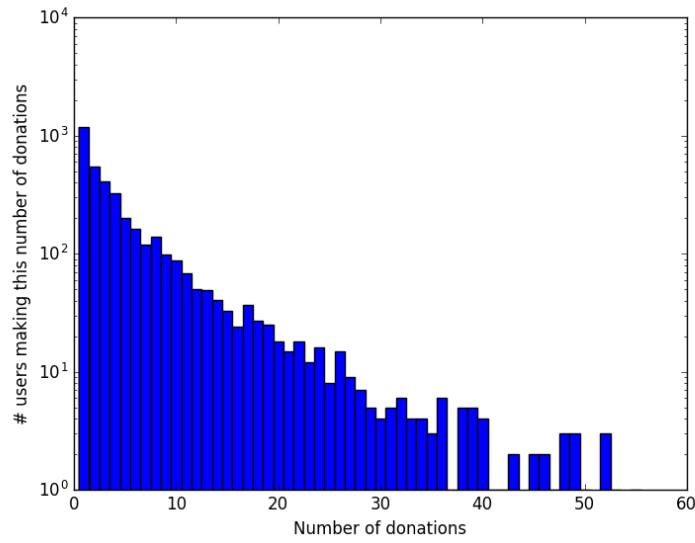


Figure 8: Number of users that make exactly n donations, Netherlands 2014

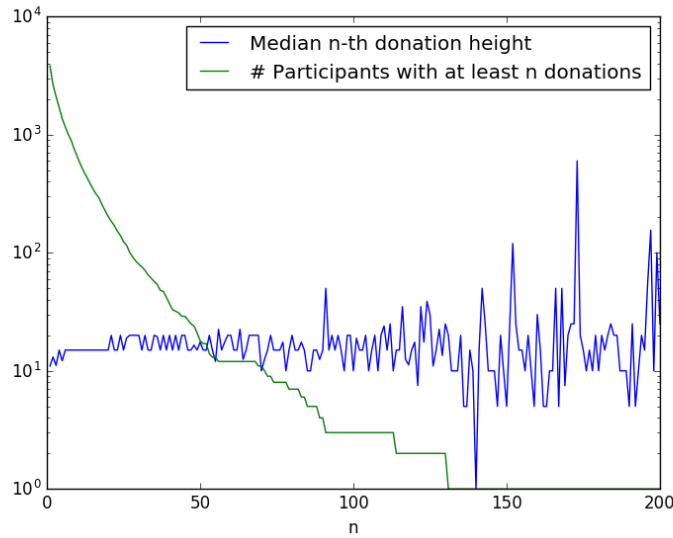


Figure 9: Median donation height of n -th donation, Netherlands 2014

someone else during the campaign. This implies that people who get mentioned are more likely to mention someone else than a random Twitter user. Conversely, 20.5% of the mentioners is being mentioned at some time during the campaign. This implies that mentioners are far more likely to be mentioned than a random Twitter user. A potential cause for this is reciprocity: a mentioner is being mentioned by his mentionee. The relation may also be explained by user activity as a cause: an active user posts more tweets on Movember, thus is more likely to mention someone at some point in time, and is at the same time more interesting to mention (because he has produced content that a user might want to refer to). Both hypotheses could be investigated with the current data. The next paragraph shows how we looked at reciprocity.

Reciprocity of mentions If user u mentions user v , we expect that the probability of v mentioning u is higher than if v is never mentioned by u . We refer to this as the reciprocity of the mentioning relation. First, we define reciprocity in two ways and then show the results for the current data set.

As a measure for reciprocity, we could look at whether users return mentions. Therefore we construct an accumulative mention graph, where all mentions are just accumulated over all time. To be more precise, let $\hat{w}(u, v)$ be the total number of times that u mentions v over all campaign days. Next, we define a graph $\hat{G}(U, E)$, where $(u, v) \in E$ if and only if $\hat{w}(u, v) > 0$. This accumulative mention graph $\hat{G}$, together with weight function $\hat{w}$ shall also be used in the next section, dealing with centralities. We can give the accumulative mention graph of The Netherlands in 2014 a reciprocity coefficient to indicate how often users in the network tend to mention people that mention them.

$$\text{Reciprocity} = \frac{\text{Number of arcs with inverse arc also present}}{\text{Number of arcs}}$$

The value is 1 if and only if every arc (u, v) has an inverse (v, u) , and 0 if and only if no arc is returned. More mathematically:

$$r_1(G) := \frac{|\{(u, v) \in U^2 | (u, v) \in E(G), (v, u) \in E(G)\}|}{|E(G)|} \quad (1)$$

Alternatively, we can count the number of reciprocal pairs of users and divide it by the pairs of users which have at least one arc running between them. If the users are labeled by natural numbers, we can write this as:

$$r_2(G) := \frac{|\{(u, v) \in U^2 | u < v \text{ and } (u, v) \in E(G), (v, u) \in E(G)\}|}{|\{(u, v) \in U^2 | u < v \text{ and } ((u, v) \in E(G) \text{ or } (v, u) \in E(G))\}|} \quad (2)$$

If we define reciprocity according to Equation (1) then we find reciprocity to be $r_1(G) = 0.146$ for the accumulative mention graph of The Netherlands in 2014. Using Equation (2), reciprocity for The Netherlands in 2014 equals $r_2(G) = 0.099$. Both results show that reciprocity is not a strong feature of the accumulative mention graph.

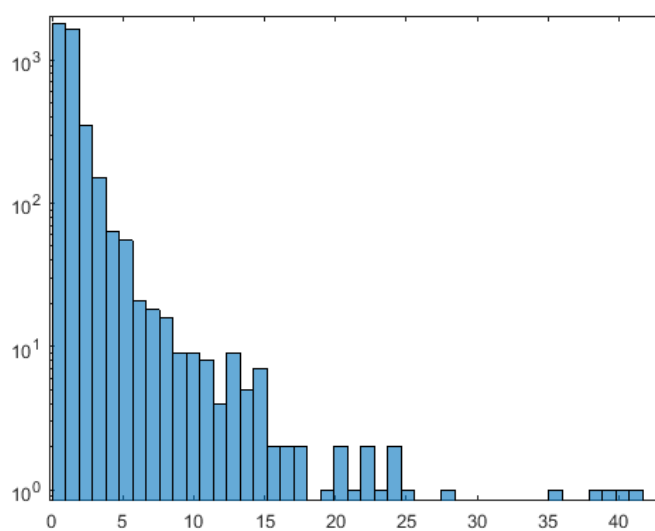


Figure 10: How often do individuals get mentioned? The number of mentions is on the horizontal axis; the number of people that are mentioned this amount of times is on the vertical axis. The Netherlands, 2014.

3 Centralities in the Twitter network of mentions

In this section we look at the centralities of individuals in the mention networks. We show that, although different measures of centrality exist in literature, a few of the most commonly used ones yield similar results. Many results in this section are based on the Twitter data set of 2013.

Introductory remarks We mainly use centralities defined on static graphs. Although dynamic centralities exist, they are not yet well-established. Therefore we use the accumulative graph, as constructed previously, where all mentions are just accumulated over all time. This accumulative graph $\hat{G}$, together with weight function $\hat{w}$ is used to calculate some centrality measures. The weight value may also be interpreted as the number of arcs from u and v , so that $\hat{G}$ is a multigraph.

3.1 Definition of centralities

We considered four types of centrality in this work:

- Weighted in-degree (Number of mentions)
- In-degree (Number of mentioners)
- PageRank centrality
- Harmonic centrality

For completeness we give a definition of PageRank centrality and harmonic centrality.

PageRank centrality This centrality is developed and used by Google. In a nutshell, if we regard a network of web pages referencing each other via hyperlinks (directed edges), the PageRank centrality of a web page equals the steady-state probability that a random surfer, arbitrarily following hyperlinks between pages, will be on this page. Formally, let $\bar{A}$ be the l_1 -normalized adjacency matrix of a graph, α a damping factor and v a preference vector. If U denotes the set of users, then the vector $p = (p_i)_{i \in U}$ of PageRanks is given by:

$$p = (1 - \alpha)v(1 - \alpha\bar{A})^{-1}.$$

For our purposes we chose $\alpha = 0.9$ and a uniform preference vector. The definition shown here was taken from [1].

Harmonic centrality Let $d(v, u)$ be the distance from user v to user u , which in our case equals the length of the shortest sequence of arcs to travel from v to u (weight does not play a role here). The harmonic centrality takes into account various axioms a centrality should rationally satisfy and is therefore included, see [1] for these axioms and the analysis. The formula is quite concise:

$$HC(u) = \sum_{v \neq u} \frac{1}{d(v, u)} \quad \forall u \in U$$

Also note that the case $d(v, u) = \infty$ is easily taken into account.

3.2 Comparison of centralities by using Kendall's weighted τ

To compare the centralities, we use Kendall's weighted τ , as introduced in [8]. The definition of Kendall's weighted τ is quite involved and therefore we include it in Appendix A. The choice for Kendall's weighted τ is made because we find differences in top centralities of higher importance than difference between persons with lower centralities (with respect to both centrality measures). So disagreement about which persons are very central is 'worse' than disagreement about centralities of peripheral persons. Kendall's weighted τ takes these ideas into account by a weighting procedure. If the Top-100's are very different for two centrality measures, we expect that Kendall's weighted τ has a low value, implying disagreement between the two measures. Following the reasoning in [8], we opt for a hyperbolic additive weighting scheme (also explained in Appendix A). Results for 2013 can be found in Table 2, results for 2014 (only The Netherlands) can be found in Table 3. The centrality scores for the number of mentioners agree most with all other scores.

Venn diagram Sweden 2013 We made a Venn diagram displaying how much overlap there is between the 'user top hundreds' that follow from the different centrality measures for Sweden in 2013, see Figure 11. Users falling outside the top hundred but with the same centrality as the user ranked at place 100 are also taken into account. We call this the set of top users for the centrality score under consideration. This means that for the score *number of mentions* we consider a set of 120 top users, for *PageRank* a set of 100 top users and for *harmonic centrality* a set of 126 top users. The resulting diagram shows that there is a large degree of agreement among the different centralities.

The top users for the *number of mentioners* are not displayed, but this group consists of 101 users that are also present in the top users for the *number of mentions* and they are present in the top users of at least one of the other two (*PageRank* and *harmonic centrality*). So these users lie in the yellow, white and magenta patches. Better still, only four users of these regions do not belong to the top users of *number of mentioners*. We can conclude that the top users of the *number of mentioners* consist of users also present in the top hundred of at least two other centralities.

Venn diagram The Netherlands 2014 For The Netherlands in 2014 we made a similar Venn diagram, see Figure 12. In this case, for the number of mentions we consider a set of 117 top users, for *PageRank* a set of 100 top users and for *harmonic centrality* a set of 101 top users. We see that the agreement among the different rankings is rather low compared to the results for Sweden in 2013.

For the number of mentioners we consider a set of 103 top users. 88 of these users are also contained in the set of top users when number of mentions are considered. 95 of the 103 top users also figure in one of the other centralities.

Sweden 2013	# Mentions	# Mentioners	PageRank	Harmonic
# Mentions	1	0.985	0.892	0.954
# Mentioners	0.985	1	0.910	0.967
PageRank	0.892	0.910	1	0.875
Harmonic	0.954	0.967	0.875	1
The Netherlands 2013	# Mentions	# Mentioners	PageRank	Harmonic
# Mentions	1	0.987	0.897	0.904
# Mentioners	0.987	1	0.904	0.905
PageRank	0.897	0.904	1	0.818
Harmonic	0.904	0.905	0.818	1

Table 2: Weighted Kendall's τ for similarity between centrality rankings for Sweden and The Netherlands in 2013

The Netherlands 2014	# Mentions	# Mentioners	PageRank	Harmonic
# Mentions	1	0.977	0.895	0.896
# Mentioners	0.977	1	0.915	0.906
PageRank	0.895	0.915	1	0.843
Harmonic	0.896	0.906	0.843	1

Table 3: Weighted Kendall's τ for similarity between centrality rankings for The Netherlands in 2014

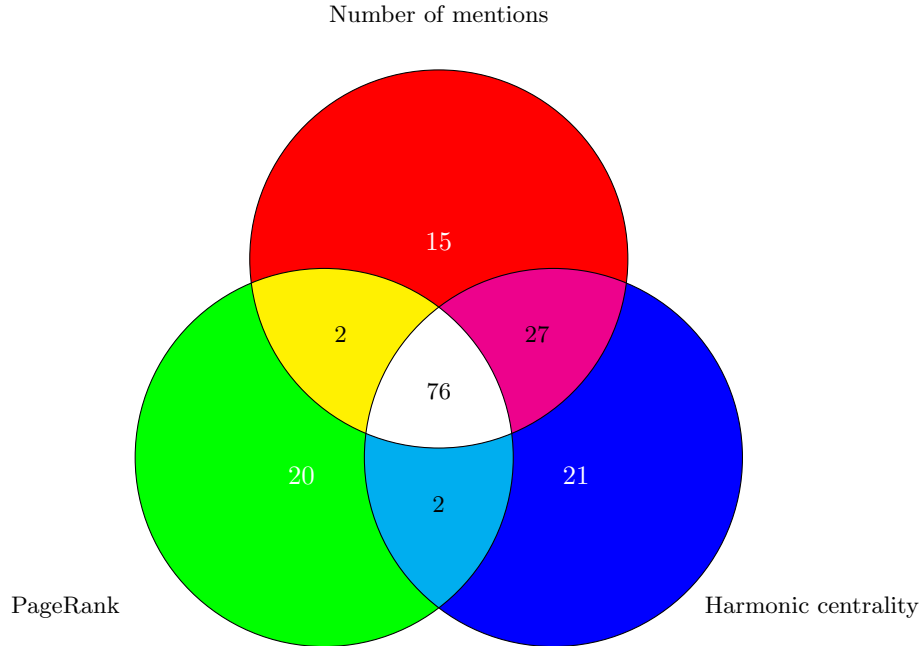


Figure 11: Swedish top users with respect to different centrality scores in 2013

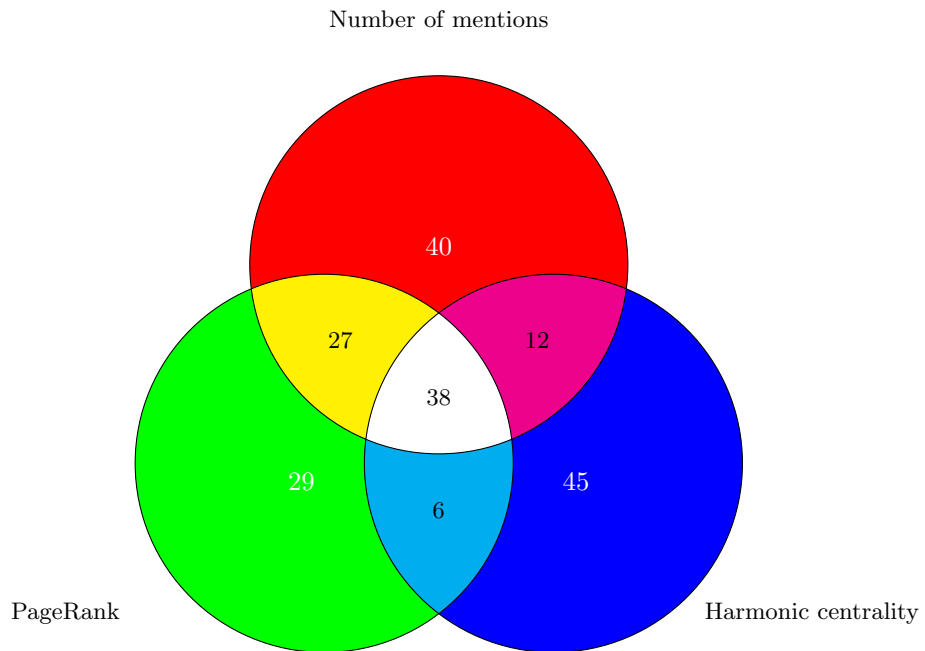


Figure 12: Dutch top users with respect to different centrality scores in 2014

4 Modeling Motivation

In this section we construct mathematical models to translate mentions into motivation levels. Some of the definitions and notations presented here will also come into use later when we introduce the Satiation-Deprivation Model.

4.1 Notations for motivation models

Let U be a set of Twitter users. In practice, the user set under consideration will be from a specific country in a specific year, so U may denote the Twitter users in The Netherlands 2014 that are included in the mention graphs. Individuals are frequently denoted by i and j , but also at times by u and v .

There is a discrete time axis denoted by $\mathcal{T} := \{0, 1, 2, \dots, T\}$. These numbers usually refer to days in the Movember campaign. The *motivation level* of an individual i at time t is denoted by $L_i^{(t)}$.

Individuals are assumed to start with an *initial motivation level* $L_i^{(1)} = \xi_i$.

The mention graphs imply an interaction between users. We define $Q^{(t)}$ for $t \geq 1$ to be a binary matrix called a *mention matrix*, where $Q_{ij}^{(t)} = 1$ if and only if user i mentions user j at time $t \in \mathcal{T} \setminus \{0\}$. Furthermore, $m_{i,j}^{(t)}$ is defined to be the amount by which user i is motivated by user j during time step t , for $i \neq j$. The method of calculating $m_{i,j}^{(t)}$ will depend on the model. We can then also define the total amount of received motivation of i during time step $t \geq 1$ as:

$$m_i^{(t)} := \sum_{j \in U \setminus i} m_{i,j}^{(t)}.$$

Translating the received amount of motivation $m_i^{(t)}$ into a motivation level $L_i^{(t)}$ will also depend on the model under consideration.

4.2 Pure Growth Model: susceptibility

In this section we start our first mathematical modeling with the Pure Growth Model. For this model, in addition to the above definitions, we introduce $m_i^{(0)} = \xi_i$ as the initial motivation level of i , obtained before any mentions take place. The total motivation level of user i at time $t \geq 1$ is then defined as:

$$L_i^{(t)} := \sum_{\tau=0}^{t-1} m_i^{(\tau)} = m_i^{(0)} + \sum_{\tau=1}^{t-1} \sum_{j \in U \setminus \{i\}} m_{i,j}^{(\tau)}. \quad (3)$$

Alternatively, $L_i^{(t)}$ can be recursively written as:

$$L_i^{(t)} = L_i^{(t-1)} + m_i^{(t-1)} \quad (4)$$

Note that the definition shows that $L_i^{(t)}$ is the motivation level of i after the motivation influences of time $t-1$ but before the motivation influences that will occur at time t . It can be considered an initial motivation level for time t .

In this model, the Pure Growth Model, we associate with every user i a susceptibility parameter $s_i \in [0, 1]$, which indicates how strongly an individual is affected by j 's motivation. Motivation only flows from j to i on day t if i mentions j on that day, given by $Q_{ij}^{(t)}$. We assume that i is then influenced by j 's motivation level just before the influences of time t occur. The fact that i mentions j is thus seen as a consequence of j being a motivating agent for i .

Additionally, for computational reasons, we define $m_{i,i}^{(t)} := L_i^t$ and $Q_{ii}^{(t)} = \frac{1}{s_i}$ for all t in the Pure Growth Model. This translates to saying that during time step t user i is motivated by himself by an amount equal to his motivation level after time $t - 1$. This subtlety is used to capture the (accumulative) motivation level in the current time step.

We now use the above to construct a recursive equation:

$$m_{i,j}^{(t)} = Q_{ij}^{(t)} s_i L_j^{(t)}, \quad \text{where } L_j^{(t)} = \sum_k m_{j,k}^{(t-1)}. \quad (5)$$

The formula says that, if i mentions j , the amount by which user i is motivated by user j is proportional to his own susceptibility and the motivation level of j . Note how in the formula for $L_j^{(t)}$ the term $m_{j,j}^{(t-1)}$ concisely captures the previous level $L_j^{(t-1)}$. Furthermore, for $j = i$, we have

$$\begin{aligned} m_{i,i}^{(t)} &= \frac{1}{s_i} s_i L_i^{(t)} \\ &= L_i^{(t)} \end{aligned}$$

This makes formula (5) consistent with the definitions.

If now $M^{(t)} = (m_{i,j}^{(t)})$ and $Q^{(t)} = (Q_{ij}^{(t)})$ and we define $s = (s_i)$ and $S = \text{diag}(s)$, then equation (5) can be rewritten at once in a compact form as:

$$M^{(t)} = S Q^{(t)} \text{diag}(M^{(t-1)} \mathbf{1}) \quad (6)$$

Note that we got rid of $L_j^{(t)}$ altogether by substituting its definition.

We now have $m^{(0)} = \xi$, a vector containing all initial motivations of the users. Then $M^{(0)} = \text{diag}(m^{(0)}) = \text{diag}(\xi)$ and we can see that the equation for $t = 1$ reduces to:

$$M^{(1)} = S Q^{(1)} \text{diag}(M^{(0)} \mathbf{1}) = S Q^{(1)} M^{(0)} \quad (7)$$

This leads us to the following set of state equations:

$$\begin{cases} L^{(t)} = M^{(t)} \mathbf{1} \\ M^{(t)} = S Q^{(t)} \text{diag}(M^{(t-1)} \mathbf{1}), & \text{for } t > 1 \\ M^{(1)} = S Q^{(1)} M^{(0)} \\ M^{(0)} = \text{diag}(\xi) \end{cases} \quad (8)$$

The model displays that the evolution of motivation depends on the initial motivations ξ , on the susceptibility of individuals and on the mentions. To get the final motivation levels after t time steps, one just calculates $L^{(t)} := M^{(t)} \mathbf{1}$. The diagonal of $M^{(t)}$ contains the previous motivation levels, the rest of the elements of row i show the contribution of all other users to the motivation level of i in the last time step. Note that the M -matrices may be particularly sparse, probably with nonzero diagonal entries, but many zeros otherwise.

In the Pure Growth Model a mentionee influences others, yet the mention does not influence his own motivation. By experiments using user sets of size $|U| = 3$ and defining simple mention matrices we saw that there was a strong dependence on the initial state and motivations could grow indefinitely. We shall deal with this divergent behaviour later in this section.

4.3 Pure Growth Model: susceptibility and reactions

The previous model used only susceptibility and the increments were proportional to the level of the mentioner, $L_j^{(t)}$. The next model also incorporated reactions of a mentionee to a mention. For

every user i it is modeled by the reaction parameter r_i . If the mentionee has a high r_i , then this leads to a relatively large motivational increment. If j mentions i , not only will j be influenced by i , but i is also influenced by getting mentioned. Moreover, in this model we assume that increments are proportional to motivational difference rather than absolute motivation.

$$\begin{aligned} m_{i,j}^{(t)} &= Q_{ij}^{(t)} s_i (L_j^{(t)} - L_i^{(t)})^+ + Q_{ji}^{(t)} r_i (L_j^{(t)} - L_i^{(t)})^+ \\ &= (Q_{ij}^{(t)} s_i + Q_{ji}^{(t)} r_i) (L_j^{(t)} - L_i^{(t)})^+ \end{aligned}$$

Summing over all $j \neq i$, we get the total increment in time t for user i :

$$m_i^{(t)} := \sum_{j \in U \setminus \{i\}} m_{i,j}^{(t)} \quad (9)$$

The motivational level of user i at time t , denoted by $L_i^{(t)}$ is then calculated in the habitual recursive way.

$$\begin{aligned} L_i^{(t)} &= L_i^{(t-1)} + \sum_{j \in U \setminus \{i\}} m_{i,j}^{(t-1)} \\ L_i^{(0)} &= \xi_i \end{aligned}$$

This model exhibited some interesting properties. Here we present a few definitions and proofs that give insight in its dynamical structure.

Theorem 1 *If $L_i^{(t)} \leq L_j^{(t)}$, then $L_i^{(t)} + m_{i,j}^{(t)} \leq L_j^{(t)} + m_{j,i}^{(t)}$ for all $i, j \in U$.*

Proof: If $L_i^{(t)} \leq L_j^{(t)}$, then $m_{i,j}^{(t)} = (Q_{ij}^{(t)} s_i + Q_{ij}^{(t)} r_i) (L_j^{(t)} - L_i^{(t)})$ and $m_{j,i}^{(t)} = 0$. So then:

$$\begin{aligned} L_i^{(t)} + m_{i,j}^{(t)} &= L_i^{(t)} + (Q_{ij}^{(t)} s_i + Q_{ij}^{(t)} r_i) (L_j^{(t)} - L_i^{(t)}) \\ &\leq L_i^{(t)} + (L_j^{(t)} - L_i^{(t)}) \\ &= L_j^{(t)} \\ &\leq L_j^{(t)} + m_{j,i}^{(t)}. \end{aligned} \quad (10)$$

■

Of course, we may interchange i and j and the theorem still holds. This theorem ensures that the motivation levels of i and j cannot affect each other in such a way that one 'overtakes' the other. So in the case that only a mention between i and j occurs in some time step, $L_i^{(t)} \leq L_j^{(t)}$ implies $L_i^{(t+1)} \leq L_j^{(t+1)}$. Moreover, by symmetry, $L_i^{(t)} = L_j^{(t)}$ implies $L_i^{(t)} + m_{i,j}^{(t)} = L_j^{(t)} + m_{j,i}^{(t)}$. (The last can also be seen by noticing that $m_{i,j}^{(t)} = m_{j,i}^{(t)} = 0$ in this special case.)

Note that the above only describes microlevel level changes. It can still happen that $L_i^{(t)} < L_j^{(t)}$, but $L_i^{(t+1)} > L_j^{(t+1)}$, thanks to contributions of other nodes (to the level of i). We found an easy example for this.

Example 2 *Suppose that $|U| = 3$. $L^{(1)} = (1 \ 2 \ 3)^T$. We assume that users are highly susceptible ($s = 1$), but not reactive ($r = 0$). At time $t = 1$, user 1 now mentions 2 as well as 3,*

so:

$$Q^{(1)} = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

(Note that $Q_{ii}^{(t)} := \frac{1}{s_i}$).

We can easily calculate that $L^{(2)} = (4 \ 2 \ 3)$. So with regard to motivational level, user 1 has overtaken users 2 and 3 by these dynamics.

For the relatively simple system consisting of only two users, we can find a direct formula to calculate the motivation levels. Let $p = \operatorname{argmin}\{\xi_1, \xi_2\}$ and $q = \operatorname{argmax}\{\xi_1, \xi_2\}$. For notational convenience, define $\omega_p(t) = Q_{pq}^{(t)} s_p + Q_{qp}^{(t)} r_p$.

Theorem 3 For all $t > 1$, the motivation levels of users p and q are given by:

$$\begin{cases} L_p^{(t)} = \xi_p \prod_{\tau=1}^{t-1} (1 - \omega_p(\tau)) + \xi_q \left[1 - \prod_{\tau=1}^{t-1} (1 - \omega_p(\tau)) \right] \\ L_q^{(t)} = \xi_q \end{cases} \quad (11)$$

Proof: In all cases, because we have just two users, by theorem 1 we can see that $L_q^{(t)} \geq L_p^{(t)}$ for all t . This means that $(L_p^{(t)} - L_q^{(t)})^+ = 0$ for all t and then it follows that $m_{q,p}^{(t)} = 0$ for all t . This implies that $L_q^{(t)} = \xi_q + \sum_{t \in T} m_{q,p}^{(t)} = \xi_q$ for all t .

The formula for $L_p^{(t)}$ is found by applying the principle of mathematical induction. For $t = 2$:

$$\begin{aligned} L_p^{(2)} &= L_p^{(1)} + \omega_p(1)(L_q^{(1)} - L_p^{(1)}) \\ &= \xi_p + \omega_p(1)(\xi_q - \xi_p) \\ &= \xi_p(1 - \omega_p(1)) + \xi_q \omega_p(1), \end{aligned}$$

which equals the value of the formula given in (11) for $t = 2$.

As the induction hypothesis, we assume that for $t = k$, the formula given by (11) holds. If we can prove that the formula therefore also holds for $t = k + 1$, we are done. This is done in the following derivation:

$$\begin{aligned} L_p^{(k+1)} &= L_p^{(k)} + \omega_p(k)(L_q^{(k)} - L_p^{(k)}) \\ &= (1 - \omega_p(k))L_p^{(k)} + \omega_p(k)L_q^{(k)} \\ &= (1 - \omega_p(k)) \left\{ \xi_p \prod_{\tau=1}^{k-1} (1 - \omega_p(\tau)) + \xi_q \left[1 - \prod_{\tau=1}^{k-1} (1 - \omega_p(\tau)) \right] \right\} + \omega_p(k)\xi_q \\ &= \xi_p \prod_{\tau=1}^k (1 - \omega_p(\tau)) + \xi_q \left\{ (1 - \omega_p(k)) \left[1 - \prod_{\tau=1}^{k-1} (1 - \omega_p(\tau)) \right] + \omega_p(k) \right\} \\ &= \xi_p \prod_{\tau=1}^k (1 - \omega_p(\tau)) + \xi_q \left[1 - \prod_{\tau=1}^k (1 - \omega_p(\tau)) \right] \end{aligned}$$

By the principle of mathematical induction, we have proven the formulas for a two-user system. ■

For the two-user system with $\mathcal{T} = \mathbb{N}$, define $L^\infty = \lim_{t \rightarrow \infty} L^{(t)}$, if this limit exists.

Corollary 4 *If $\lim_{t \rightarrow \infty} \prod_{\tau=1}^t (1 - \omega_p(\tau)) = 0$, then $L^\infty = M\mathbf{1}$. If $\omega_p(t) = 0$ for all t , then $L^\infty = \xi$.*

Proof: This follows from simple substitution of the conditions in (11) and performing a limiting procedure. ■

We observe that for the two-user system, the following holds: if $\xi_p = \xi_q$, then $L_p^{(t)} = L_q^{(t)} = \xi_p$. This means that the motivation levels are constant, no matter how the users behave.

We generalized this observation. Therefore we characterize the set of initial motivation levels that are unaffected by any type of dynamics. These are fixed points for motivational systems, independent of whatever kind of dynamics occur.

Definition 5 *L is said to be a dynamics-independent fixed point (of a network of users) if $\xi = L$ implies $L^{(\tau)} = L$ for all points $\tau \in \mathcal{T}$ in time, and this holds for any sequence of mention matrices $(Q^{(t)})_{t=1}^T$, and any susceptibility and reaction vectors s and r .*

It is easy to see that when everybody has equal motivation, their motivation will not change over the entire time horizon. These can also be proven to be the only distributions that are dynamics-independent fixed points.

Theorem 6 *The set of dynamics-independent fixed points Λ for any set of users is given by:*

$$\Lambda := \{c\mathbf{1} | c \in \mathbb{R}\}$$

Here, $|\mathbf{1}| = |U|$.

Proof: The proof goes in two steps:

- $L = c\mathbf{1}$ is a dynamics-independent fixed point for every $c \in \mathbb{R}$.
- If some L is a dynamics-independent fixed point, then $L \in \Lambda$.

The first part is by mathematical induction. If $\xi = L^{(1)} = c\mathbf{1} \in \Lambda$, then by definition $m_{i,j}^{(1)} = 0$. So $L_i^{(2)} = L_i^{(1)} + \sum_{j \in U \setminus \{i\}} m_{i,j}^{(1)} = L_i^{(1)}$. This implies that $L^{(2)} = L^{(1)}$. For the induction step, let $L^{(t)} = c\mathbf{1}$. By definition, then $m_{i,j}^{(t)} = 0$. So $L_i^{(t+1)} = L_i^{(t)} + \sum_{j \in U \setminus \{i\}} m_{i,j}^{(t)} = L_i^{(t)}$ for all $i \in U$ and $L^{(t+1)} = L^{(t)}$. By the principle of mathematical induction, we find that $L^{(\tau)} = L$ for all τ . Thus for every $c \in \mathbb{R}$ we find that $L = c\mathbf{1}$ is a dynamics-independent fixed point.

The second part is by contraposition. Suppose that $L \notin \Lambda$. Then there exist users i and j such that $L_i < L_j$. If, at some time step, $L^{(t)} = L$, then we can choose $r_i = s_i = \frac{1}{2}$, and $Q_{ij}^{(t)} = Q_{ji}^{(t)} = 1$ and the rest of the Q -matrix zero. Then $m_{i,j}^{(t)} = L_j^{(t)} - L_i^{(t)}$ and, because this is the only non-zero increment, we have that $L_i^{(t+1)} = L_i^{(t)} + \sum_{j \in U \setminus \{i\}} m_{i,j}^{(t)} = L_j^{(t)} > L_i^{(t)}$. This implies that $L^{(t+1)} \neq L^{(t)}$ for this choice of dynamics. So L is not a dynamics-independent fixed point. ■

The analysis of this version of the Pure Growth Model becomes involved quite rapidly. Questions like 'when does this model converge to a dynamics-independent fixed point?' remain unanswered. Furthermore it is still possible for motivation levels to trail off to infinity, see Example 7

Example 7 Choose $|U| = 3$, $r = 1$, $s = 0$, $\xi = (1 \ 2 \ 3)^T$, and every time step the two highest-motivated users mention the lowest-motivated one.

Because the reaction is one for the lowest-motivated user, the user increases in motivation by three (a difference of 1 and a difference of 2 with the higher-motivated users). The resulting vectors of motivation levels are as follows:

$$\begin{aligned} L^{(1)} &= (1 \ 2 \ 3)^T \\ L^{(2)} &= (4 \ 2 \ 3)^T \\ L^{(3)} &= (4 \ 5 \ 3)^T \\ L^{(4)} &= (4 \ 5 \ 6)^T \end{aligned}$$

It is not difficult to see that in general $L^{(1+3n)} = (1+3n \ 2+3n \ 3+3n)^T$ for $n \in \mathbb{N}$. Thus for every $M > 0$ there exists $n \in \mathbb{N}$ such that $L_i^{(1+3n)} > M$. Because $(L_i^{(k)})_{k \in \mathbb{N}}$ is a non-decreasing sequence for every i (motivation levels cannot decrease), we have that $k \geq 1+3n$ implies $L_i^{(k)} \geq L_i^{(1+3n)} > M$ and thus $L_i^{(k)} \rightarrow \infty$ as $k \rightarrow \infty$ for all $i \in U$.

4.4 Motivational decay

The models developed in the previous subsections still lack realism, because motivation levels can grow indefinitely. Moreover, in the second model (see Section 4.3), there is a constant fixed point, which is also not an attractive feature. Therefore, we introduce an effect that counters motivational increase: negative drift. Beside its convenience as a mathematical escape route, it is also reasonable from the view of motivation mechanisms.

To see why, let us return to the formula for the next motivation level.

$$L_i^{(t+1)} = L_i^{(t)} + \sum_{j \in U \setminus \{i\}} m_{i,j}^{(t)}.$$

Because the increments are non-negative, levels are always increasing. Moreover, someone who is motivated will retain his motivation level even when he is not mentioned or mentioning at all, so says this model. However, it appeals to the intuition that such an 'isolated' individual, lacking social incentives, would gradually feel less connected or inspired to participate in the activism. We could model this, for example, with a drift function:

$$L_i^{(t+1)} = h_i(L^{(t)}) + \sum_{j \in U \setminus \{i\}} m_{i,j}^{(t)}.$$

The drift term could be dependent on the total state, but we could also choose to simply define the drift function as:

$$h_i(L^{(t)}) = \alpha_i L_i^{(t)}, \quad 0 \leq \alpha_i \leq 1. \quad (12)$$

We can even simplify this by setting $\alpha_i = \alpha$ for all $i \in U$. Choosing $\alpha = 1$ we then get the original model without drift. Choosing $\alpha = 0$ means that the motivation level at each time point is fully determined by the mentions of the preceding time instance. This definition of drift only makes sense if motivations are restricted to the positive real numbers, which is the case if we demand $\xi \geq 0$. For $\alpha \in (0, 1)$ there is an exponential decay of motivation level over time.

Unfortunately, for the models presented in this section we could not prove that a drift function of the above types puts any bound on the height of motivation levels.

The concept of motivational drift will be used in the Satiation-Deprivation Model, explored in the next section. Interestingly, by the definition of that model, the drift is not even necessary to bound the motivation levels, but is only included to capture the intuition of motivational decay in the absence of mentions.

Note: We could also have chosen to set $h_i(L^{(t)}) = L_i^{(t)} - \alpha$, which describes a constant drift of the level. This model captures a natural state-independent negative drift in motivation, while on the other hand, the level is boosted by mentioning and being mentioned. We should treat the constant drift with care; if there are no mentions ever, all motivations will drift to negative infinity. One ad hoc solution is to make a slight change to this formula and set $h_i(L^{(t)}) = (L_i^{(t)} - \alpha)^+$. The exponential drift seems more elegant, so for the rest of this research we used equation (12) with $\alpha_i = \alpha$.

5 The Satiation-Deprivation Model: Construction and Analytical Results

The Satiation-Deprivation (SD) Model forms the core of this research project and the synthesis of ideas developed throughout Section 4. In line with those models, it aims to describe a dynamic motivation process that is boosted by Twitter mentions and that drops in the absence of them. It is the analytical tractability of the SD Model and its ability to keep motivation levels contained in $[0, 1]$ that make it a more attractive alternative to the models presented in the previous section. In Section 5.1 we first develop requirements for the model; what criteria should it satisfy? Section 5.2 explains the ideas behind the model and mathematically constructs the SD Model. In Section 5.3 we find analytical results for the SD Model and analytically characterize the stationary state of a motivation process.

5.1 Model requirements

The model consists of a function that transfers mentions, located on a time axis, to a motivation level that develops through time. It should satisfy the following criteria:

1. Motivation levels should always lie between 0 and 1.
2. More mentions in the recent past should imply a higher motivation level.
3. Mentions that lie further in the past should contribute less to the current motivation level.
4. If there are more mentions in the recent past, a new mention relatively affects the motivation level less. This is in agreement with the **satiation** proposition in social exchange theory, see [3].
5. If there are less mentions in the recent past, a new mention relatively affects the motivation level more. This is in agreement with the **deprivation** proposition in social exchange theory, see [3].
6. There should be room for including a measure of the mentioner's degree of influence. A mention of a high-influence mentioner should induce a higher motivation level than a mention of a low-influence mentioner.

Note that we do not include criteria for susceptibilities or reactions. They do not appear in the SD Model, because we could not find an appropriate way of calculating these parameters based on the data. Instead, we introduce a parameter that models the mentioner's degree of influence, thus capturing individual differences.

We shall see that the criteria can be synthesized quite neatly. The next argument is essential to this synthesis.

The motivation-satiation argument The motivation level is comprised of mentions in the recent and distant past, where those in the recent past weigh heavier than those in the far past. Thus, a high motivation level implies a few recent mentions or many of them in the past, and therefore implies a high *level of satiation* as well. In other words, the higher an individual is motivated, the more difficult it becomes to increase his motivation further. New mentions do not mean much to whom is already motivated.

We shall see that this mechanism of the motivation-satiation argument helps us in keeping motivation levels between 0 and 1, thus providing a normalized motivation score for each user without an artificial normalization step.

5.2 Theoretical development of the SD Model

For the development of the motivation model we resort to an analogy with painters and paintings. First imagine that each user i has at his disposal an (initially empty) canvas. Mentioners of i are regarded as painters that are randomly covering parts of the i 's canvas in paint; the more important the mentioner, the larger the area he can paint. Paint on the canvas can be translated to motivation of user i . This means that a larger fraction of the canvas covered in paint implies a larger motivation level of the user. Painters work independently on the same canvas and can thus cover the same parts over and over again.

We use the painter analogy to develop the theoretical underpinnings of the Satiation-Deprivation Model. The result will be a stochastic recursion.

First, for the sake of exposition, let us assume that the canvas of user i is empty at time t . Now there is a set of painters, denoted by $J_i^{(t)}$, that visit the canvas of i at time t and start to paint. The probability of each point on the canvas to be painted by $j \in J_i^{(t)}$ is assumed equal to ϕ_j . We call ϕ_j the *degree of influence of user j* .

As an approximation for the expected coverage of i 's canvas, divide the canvas into n equal patches and suppose that each patch has probability ϕ_j to be painted (entirely) by j . Let $X_p^{(t)}$ denote the indicator that patch p is newly painted at time t . Because the painters work independently, it is seen that:

$$P(X_p^{(t)} = 0) = \prod_{j \in J_i^{(t)}} (1 - \phi_j) \quad \forall p \in \{1, \dots, n\}.$$

Because $P(X_p^{(t)} = 1) = 1 - P(X_p^{(t)} = 0)$, we can calculate the expectation of $X_p^{(t)}$.

$$\mathbb{E}[X_p^{(t)}] = 1 - \prod_{j \in J_i^{(t)}} (1 - \phi_j).$$

We define $L_i^{(t)}$ to be the fraction of points that is painted at time t . Clearly we had $L_i^{(t)} = 0$. Now we can calculate $L_i^{(t+1)}$. To this end, we let the number of patches go to infinity (so each patch shrinks to a point) and invoke the Law of Large Numbers:

$$L_i^{(t+1)} := \lim_{n \rightarrow \infty} \frac{\sum_{p=1}^n X_p^{(t)}}{n} = \mathbb{E}[X_p^{(t)}]. \quad \text{w.p. 1}$$

We shall now define $\Phi_i^{(t)} := 1 - \prod_{j \in J_i^{(t)}} (1 - \phi_j)$ for conciseness.

Now, to generalize, suppose $L_i^{(t)} \geq 0$. Because $X_p^{(t)}$ denotes whether patch p is *newly* painted at time t , we must only regard paint that hits the currently unpainted part of the canvas, which is a fraction $1 - L_i^{(t)}$.

$$X_p^{(t)} = \begin{cases} 1 & \text{w.p. } \Phi_i^{(t)} (1 - L_i^{(t)}) \\ 0 & \text{w.p. } 1 - \Phi_i^{(t)} (1 - L_i^{(t)}) \end{cases}.$$

Furthermore we assume that the paint on the already painted fraction decays at rate α . Using

another limiting procedure with the Law of Large Numbers, we obtain:

$$\begin{aligned}
 L_i^{(t+1)} &= \alpha L_i^{(t)} + \lim_{n \rightarrow \infty} \frac{\sum_{p=1}^n X_p^{(t)}}{n} \\
 &= \alpha L_i^{(t)} + \mathbb{E} [X_p^{(t)}] \\
 &= \alpha L_i^{(t)} + \Phi_i^{(t)} (1 - L_i^{(t)}) \\
 &= (\alpha - \Phi_i^{(t)}) L_i^{(t)} + \Phi_i^{(t)},
 \end{aligned}$$

with probability one (this follows from the first to the second line, where the Law of Large Numbers is used again). This concludes the theoretical development of the model. For references we summarize the result in a definition.

Definition 8 *The Satiation-Deprivation (SD) Recursions are defined by the recursive equations:*

$$\begin{aligned}
 L_i^{(t+1)} &= (\alpha - \Phi_i^{(t)}) L_i^{(t)} + \Phi_i^{(t)}, \\
 \Phi_i^{(t)} &= 1 - \prod_{j \in U} (1 - \phi_j Q_{ji}^{(t)}), \\
 L_i^{(0)} &= 0.
 \end{aligned} \tag{13}$$

Note that this definition shows already that the order of painters (mentioners) visiting i is irrelevant within the same time step. So if j and k visit i at the same time step, it does not matter who paints first; the motivation level will be the same.

5.3 The SD Model as a stochastic model

The stochastic model is used to derive theoretical properties of the Satiation-Deprivation Model. The assumption is that each user i mentions user j with probability q_{ij} at each point in time. Let $(Q^{(t)})_{t \in T}$ be a sequence of matrices indexed by time horizon T where each entry of each matrix is the outcome of a Bernoulli experiment:

$$Q_{ij}^{(t)} = \begin{cases} 1 & \text{w.p. } q_{ij} \\ 0 & \text{w.p. } 1 - q_{ij} \end{cases}$$

The matrices are identically distributed and independent of one another. We shall use a matrix Q as dummy stochastic matrix with the same distribution.

Define the stochastic process $\{L^{(t)}\}_{t \in \mathbb{Z}}$ by the first two equations of the SD Recursions, see Equation (13), where $Q_{ij}^{(t)}$ are stochastic. We thus make assumptions on whether persons mention each other; such mentions are independent of each other and their probability of occurrence does not change over time. Note that $0 \leq \Phi_i^{(t)} \leq 1$ always. We let $t \in \mathbb{Z}$ and disregard the equation $L_i^{(0)} = 0$; at $t = 0$ the process has already been going on for infinite time. This definition aids in finding convergence results.

Questions of interest are whether $L^{(t)}$ converges in distribution, and if yes, what is this distribution? Following the work of Brandt ([2]), we can find an analytic formula.

We use Theorem 1 of [2] and its corollary in the proof of the next Theorem.

Theorem 9 *Let $\Psi = \{(\alpha - \Phi_i^{(t)}, \Phi_i^{(t)})\}_{t=-\infty}^{\infty}$ be the sequence of coefficients in Equation (13). If either:*

- $\alpha \in (0, 1)$
- $\alpha = 0$ and $P\{\Phi_i^{(0)} < 1\} > 0$
- $\alpha = 1$ and $P\{\Phi_i^{(0)} > 0\} > 0$,

then the only stationary solution for the SD Recursions (see Equation (13)) with stochastic $(Q^{(t)})_{t \in \mathcal{T}}$ is given by:

$$l_i^{(t)}(\Psi) = \sum_{\tau=0}^{\infty} \left(\prod_{k=t-\tau}^{t-1} (\alpha - \Phi_i^{(k)}) \right) \Phi_i^{(t-\tau-1)} \quad (14)$$

Moreover, if we start at an arbitrary $L_i^{(0)}$, then:

$$\lim_{t \rightarrow \infty} L_i^{(t)} \stackrel{d}{\sim} l_i^{(0)}(\Psi)$$

It also holds that:

$$l_i^{(0)}(\Psi) \stackrel{d}{=} (\alpha - \Phi_i^{(0)}) l_i^{(0)}(\Psi) + \Phi_i^{(0)} \quad (15)$$

Proof: For given Ψ , we have that $(\alpha - \Phi_i^{(t)}, \Phi_i^{(t)})$ are i.i.d. distributed, with $\Phi_i^{(t)} \sim \Phi_i$ for all t . The sequence Ψ is stationary (which follows from the pairs being i.i.d.) and ergodicity follows from the i.i.d assumption and the Law of Large Numbers. These are two requirements on Ψ for the theorem to hold.

For the results we need still verify that $\mathbb{E}[\log |\alpha - \Phi_i^{(0)}|] < 0$ and $\mathbb{E}[\log |\Phi_i^{(0)}|]^+ < \infty$, according to [2].

Because $0 \leq \Phi_i^{(0)} \leq 1$, we have $0 \leq |\alpha - \Phi_i^{(0)}| \leq \max\{1 - \alpha, \alpha\}$. Then for $\alpha < 1$, we have $\log |\alpha - \Phi_i^{(0)}| \leq \log \max\{\alpha, 1 - \alpha\} < 0$. Secondly, if $\alpha = 0$ and $P\{\Phi_i^{(0)} < 1\} > 0$ (so that $\Phi_i^{(0)}$ is not always 1), $\log |\alpha - \Phi_i^{(0)}|$ will be either zero or smaller than zero. The latter happens with positive probability, so $\mathbb{E}[\log |\alpha - \Phi_i^{(0)}|] < 0$. Thirdly, if $\alpha = 1$, and $P\{\Phi_i > 0\} > 0$ (so that Φ_i is not always 0), $\log |\alpha - \Phi_i^{(0)}| = 0$ with probability $P\{\Phi_i = 0\} < 1$, and otherwise it is negative (with probability > 0). It then follows as well that $\mathbb{E}[\log |\alpha - \Phi_i^{(0)}|] < 0$.

Furthermore, we have $\mathbb{E}[\log |\Phi_i|]^+ = 0$, because $\log |\Phi_i| = \log \Phi_i \leq 0$.

Equation (15) holds by the i.i.d. assumption on Q and thus on Φ , and the corollary to Theorem 1 in [2]. $\blacksquare$

We find an expression relating the expectation of the motivation level to the mention probability matrix. We leave out the argument Ψ from now on because it is clear which sequence we always regard.

$$\begin{aligned} \mathbb{E}[l_i^{(t)}] &= \mathbb{E} \left[\sum_{\tau=0}^{\infty} \left(\prod_{k=t-\tau}^{t-1} (\alpha - \Phi_i^{(k)}) \right) \Phi_i^{(t-\tau-1)} \right] \\ &= \sum_{\tau=0}^{\infty} \mathbb{E} \left[\left(\prod_{k=t-\tau}^{t-1} (\alpha - \Phi_i^{(k)}) \right) \Phi_i^{(t-\tau-1)} \right] \\ &= \sum_{\tau=0}^{\infty} \prod_{k=t-\tau}^{t-1} (\alpha - \mathbb{E}[\Phi_i^{(k)}]) \mathbb{E}[\Phi_i^{(t-\tau-1)}] \\ &= \mathbb{E}[\Phi_i] \sum_{\tau=0}^{\infty} (\alpha - \mathbb{E}[\Phi_i])^\tau \end{aligned} \quad (16)$$

From the second to the third line we used independence of the $\Phi_i^{(k)}$ and $\Phi_i^{(t-\tau-1)}$. The independence follows from the independence of the $\Phi_i^{(x)}$ and $\Phi_i^{(y)}$ if $x \neq y$ and the fact that $k \neq t - \tau - 1$ in every term of the infinite sum. In the last step we used that all $\Phi_i^{(t)}$ are identically distributed and $\Phi_i^{(t)} \sim \Phi_i$. Note that the expected value does not depend on t for the stationary distribution.

The expectation of Φ_i is given by:

$$\mathbb{E}[\Phi_i] = 1 - \prod_{j \in U} (1 - \phi_j q_{ji})$$

Assuming that $\alpha \in (0, 1)$, we have $|\alpha - \mathbb{E}[\Phi_i]| < 1$ and from (16) we obtain:

$$\mathbb{E}[l_i^{(t)}] = \frac{\mathbb{E}[\Phi_i]}{1 - \alpha + \mathbb{E}[\Phi_i]}$$

This ensures also that the expectation is between 0 and 1, as $\alpha \in (0, 1)$.

For the case $\alpha = 0$, the process can be seen as restarting at each time point, and it is trivial that then $\mathbb{E}[l_i^{(t)}] = \mathbb{E}[\Phi_i]$. Here we defined $l_i^{(t)}$ to be the stationary distribution, which equals the distribution of Φ_i .

The case $\alpha = 1$ is a little more involved. If $\mathbb{E}[\Phi_i] > 0$ (so $P\{\Phi_i > 0\} > 0$), the series in (16) still converges, and we simply get $\mathbb{E}[l_i^{(t)}] = 1$. This agrees with common sense: if there is no discounting, but there are mentions of i , the process will in expectation reach one. In this case, because additionally $P(\lim_{t \rightarrow \infty} L_i^{(t)} > 1) = 0$, we must have that any process converges to 1 almost surely, otherwise we would have $\mathbb{E}[l_i^{(t)}] < 1$.

If we have $\alpha = 1$ and $\mathbb{E}[\Phi_i] = 0$, we cannot use Theorem 9 or the series. There may only be a null set for which $\Phi_i > 0$ (as $\Phi_i < 0$ cannot occur). Because the number of time points is countable, we argue that $l_i^{(t)} = L_i$ almost surely if $l_i^{(0)} = L_i$, thus $\mathbb{E}[l_i^{(t)}(L_i)] = L_i$.

Dependence of expected motivation level on Q Let $\mathbb{E}[Q_{\cdot, i}] = (q_{ji})_{j \in U}$, a vector of length $|U|$. For two vectors u and v , we write $u > v$ if $u_j \geq v_j$ for all j , and there exists at least one k such that $u_k > v_k$ strictly.

Proposition 10 *Let P and R be stochastic matrices where the elements (i, j) follow a Bernoulli distribution with parameters p_{ij} and r_{ij} respectively. Let $\alpha < 1$. If $\mathbb{E}[P_{\cdot, i}] > \mathbb{E}[R_{\cdot, i}]$, then $\mathbb{E}[l_i^{(t)} | Q = P] \geq \mathbb{E}[l_i^{(t)} | Q = R]$.*

Proof: Substituting into the expectations as derived in this section yields the result.

6 The Satiation-Deprivation Model: Sequential Approach and Extensions

In Section 5 we constructed the Satiation-Deprivation Model using an analogy with painters. In this section the analogy is continued in order to develop a sequential approach to the problem and suggest various extensions of the model. The algorithms in this section give the Satiation-Deprivation Model a broader scope at the cost of analytical tractability. For the simplest version of the algorithm we prove the equivalence with the Satiation-Deprivation Recursions, see Equation (13). We shall begin with this simplest version and build to more complex algorithms.

6.1 Notations for the sequential approach

Some new notations are added to the ones used in previous sections. These allow for generalization of the model. To this end, let C be a *characteristics set*. Each user $i \in U$ has an associated *characteristics distribution* $R_i = (R_{i,c})_{c \in C}$, where $R_{i,c} \geq 0$ and $R_i = \sum_{c \in C} R_{i,c} = 1$. Furthermore, there is a *characteristic motivation level* $L_{i,c}^{(t)}$ associated with each character trait of each person, developing over the time horizon $\mathcal{T}$. It holds that $0 \leq L_{i,c}^{(t)} \leq R_{i,c}$ for all $t \in \mathcal{T}$. The total *motivation level* of user i is defined as $L_i^{(t)} := \sum_{c \in C} L_{i,c}^{(t)}$. As before, each user has a degree of influence given by $\phi_i \in [0, 1]$. Let $J_i^{(t)} := \{j \in U : Q_{ji}^{(t)} = 1\}$.

6.2 The sequential approach for $|C| = 1$

The case $|C| = 1$ actually excludes the entire characteristics set and is therefore the simplest case. It implies that $R_{i,c} = R_i = 1$ for all $(i, c) \in U \times C$.

We now return to the painter analogy. We can visualize the process as follows. Each user has an empty (white) canvas with area $R_i = 1$. Independently from that, the user also has a can of paint. The (red) paint of user j can cover an area ϕ_j of any canvas with paint whenever j visits the canvas. (Note that this interpretation is slightly different from the one in Section 5.2, where we used patches). Suppose that user j visits the canvas of user i at time t , so $j \in J_i^{(t)}$. User j will now use all his paint to cover a random area of the canvas. Of course this covers an area ϕ_j of the canvas.

Now it can be that at the same moment $k \in J_i^{(t)}$ as well ($k \neq j$), so user k visits the canvas of user i and also paints a random area of the canvas. By the Law of Large Numbers a fraction ϕ_j of his paint will hit paint-covered places, and thus not contribute to the area. A fraction $(1 - \phi_j)$ will cover white canvas. The total painted surface will equal: $1 - (1 - \phi_j)(1 - \phi_k)$. Note that double layers do not have any meaning for this process - we should regard paint that hits already painted areas as being lost instantly.

We can see (also from Section 5.2) that it does not matter in which order we treat j and k ; the resulting motivation level will be the same. The sequential approach can be given in an algorithmic form, see Algorithm 1. In the algorithm $L_i^{(t, counter)}$ keeps track of the motivation level when various painters visit the canvas during the same period, but in some order. We compute the resulting change in motivation level by using the current level $L_i^{(t)}$ and the degree of influence from i 's mentioners/painters.

Algorithm 1 Satiation-Deprivation Algorithm for $|C| = 1$

Input: U, T, ϕ, α
 $J_i^{(t)}$ for all $i \in U, t \in [T]$
Output: $L_i^{(t)}$ for all $i \in U, t \in [T]$

$L_i^{(0)} = 0 \quad \forall i \in U$
for $t \in \{1, \dots, T\}$ **do**
 for $i \in U$ **do**
 $counter = 0$
 $\tilde{L}_i^{(t, counter)} = L_i^{(t)}$
 for $j \in J_i^{(t)}$ **do**
 $counter = counter + 1$
 $m_{i,j}^{(t)} = \phi_j(1 - \tilde{L}_i^{(t, counter-1)})$
 $\tilde{L}_i^{(t, counter)} = \tilde{L}_i^{(t, counter-1)} + m_{i,j}^{(t)}$
 end for
 $L_i^{(t+1)} = \alpha L_i^{(t)} + \left(\tilde{L}_i^{(t, counter)} - L_i^{(t)} \right)$
 end for
end for

6.3 The equivalence of the SD Recursions and the sequential approach for $|C| = 1$

In this section we show that the Satiation-Deprivation Recursions ((13)) and the sequential approach for $|C| = 1$ are equivalent in the sense that they yield the same sequences of motivation levels $(L^{(t)})_{t \in \mathcal{T}}$ given $(Q^{(t)})_{t \in \mathcal{T}}$, α and $(\phi_i)_{i \in U}$.

Theorem 11 *Let $|C| = 1$. The SD Recursions given by (13) are equivalent with the sequential approach.*

Proof: For the proof we rewrite the recursive part of the SD Recursions in a perhaps slightly more revealing way:

$$L_i^{(t+1)} = \alpha L_i^{(t)} + \Phi_i^{(t)} (1 - L_i^{(t)})$$

We do mathematical induction on t , and for every t we do mathematical induction on the number of mentioners.

For $t = 0$, the result is trivial.

Now for some $t \geq 0$ suppose we have calculated $L_i^{(t)}$ in the sequential approach and also with the SD Recursions and they are shown to be equal. We prove that $L_i^{(t+1)}$ from (13) is then indeed given by the sequential approach. Let $J_i^{(t)} = \{j_1, \dots, j_{n_i}\}$ be the set of persons that mention i at time t , and $n_i = |J_i^{(t)}|$. Clearly, if $n_i = 0$, then $\Phi_i^{(t)} = 0$ and $L_i^{(t+1)} = \alpha L_i^{(t)}$. This agrees with the Satiation-Deprivation Model.

Now suppose $n_i = 1$. We then see that $\Phi_i^{(t)} = \phi_{j_1}$, and $L_i^{(t+1)} = \alpha L_i^{(t)} + \phi_{j_1}(1 - L_i^{(t)})$, which also agrees with the sequential approach. Now suppose that the statement holds for $n_i = k$. We now prove the statement for $n_i = k + 1$.

So let $n_i = k + 1$. Define $m_{i,[k]}^{(t)} := \Phi_{i,k}^{(t)}(1 - L_i^{(t)})$, where $\Phi_{i,k}^{(t)} = 1 - \prod_{j \in U - j_{k+1}} (1 - \phi_j Q_{ji}^{(t)})$. This is the motivational increment of i at time t directly after the first k painters have done their job, but before painter $j_{k+1} = j_{n_i}$ does his job. Note that this formula holds for both approaches by

the induction hypothesis. Next, painter $j_{k+1} = j_{n_i}$ comes in, so that:

$$m_{i,[k+1]}^{(t)} = m_{i,[k]}^{(t)} + \phi_{j_{k+1}} \left(1 - L_i^{(t)} - m_{i,[k]}^{(t)} \right),$$

according to the sequential approach. So, according to the algorithm, the eventual motivation level becomes:

$$\begin{aligned} L_i^{(t+1)} &= L_i^{(t)} + m_{i,[k+1]}^{(t)} \\ &= L_i^{(t)} + m_{i,[k]}^{(t)} + \phi_{j_{k+1}} \left(1 - L_i^{(t)} - m_{i,[k]}^{(t)} \right) \\ &= L_i^{(t)} + \Phi_{i,k}^{(t)} \left(1 - L_i^{(t)} \right) + \phi_{j_{k+1}} \left(1 - \Phi_{i,k}^{(t)} \right) \left(1 - L_i^{(t)} \right) \\ &= L_i^{(t)} + \left(\Phi_{i,k}^{(t)} + \phi_{j_{k+1}} - \phi_{j_{k+1}} \Phi_{i,k}^{(t)} \right) \left(1 - L_i^{(t)} \right). \end{aligned}$$

Therefore we need to show that for $n_i = k + 1$, we have that $\Phi_i^{(t)} = \Phi_{i,k}^{(t)} + \phi_{j_{k+1}} - \phi_{j_{k+1}} \Phi_{i,k}^{(t)}$ by the SD Recursions. This is verified by the following calculations:

$$\begin{aligned} \Phi_i^{(t)} &= 1 - \prod_{j \in U} \left(1 - \phi_j Q_{ji}^{(t)} \right) \\ &= 1 - (1 - \phi_{j_{k+1}}) \prod_{j \in U - j_{k+1}} \left(1 - \phi_j Q_{ji}^{(t)} \right), \quad \text{because } Q_{j_{k+1}i}^{(t)} = 1 \\ &= 1 - \prod_{j \in U - j_{k+1}} \left(1 - \phi_j Q_{ji}^{(t)} \right) + \phi_{j_{k+1}} - \phi_{j_{k+1}} \left(1 - \prod_{j \in U - j_{k+1}} \left(1 - \phi_j Q_{ji}^{(t)} \right) \right) \\ &= \Phi_{i,k}^{(t)} + \phi_{j_{k+1}} - \phi_{j_{k+1}} \Phi_{i,k}^{(t)} \end{aligned}$$

This concludes the proof that the sequential approach is equivalent with the SD Recursions. $\blacksquare$

6.4 The sequential approach for $|C| \geq 1$

If we let $|C| > 1$, we can let the painters be more dynamic. They could have a preference for painting certain parts of the canvas. Dividing the canvas of user i into (not necessarily equal) patches of sizes $R_{i,1}, \dots, R_{i,|C|}$, we can let the painter's interaction with the canvas depend on his own canvas. The painter treats every patch separately. For each patch he covers a fraction ϕ_j of this 'patch' with paint, but he can not use more paint for the patch than would be necessary to paint his own patch of the same characteristic. If his canvas is partitioned exactly the same as that of i , this is no restriction at all. On the other hand, if the partitions differ a lot, the painter will not be able to paint a lot: a bigger patch on his own canvas does not give more freedom, because he is limited to the patch size of i . Besides this effect, it also means that other patches must be smaller, which limit the painter.

The sequential approach for $|C| \geq 1$ can also be given by an algorithm, see Algorithm 2. Note that this algorithm can also be used for the case $|C| = 1$.

We note that for this algorithm the order of users in $J_i^{(t)}$ can have influence on the resulting motivation level, because of the more complex interactions of different painters.

This complexer version of the model may be useful for the Movember research in later stages, when we know more about the users themselves and their identities. Specifically, Anna Priante is working on measuring identities on Twitter with regard to the Movember campaign. The more

Algorithm 2 Satiation-Deprivation Algorithm for $|C| \geq 1$

Input: U, C, T, ϕ, α
 $R_{i,c}$ for all $i \in U, c \in [C]$
 $J_i^{(t)}$ for all $i \in U, t \in [T]$
Output: $L_{i,c}^{(t)}$ for all $i \in U, c \in [C], t \in [T]$

$L_{i,c}^{(0)} = 0 \quad \forall (i, c) \in U \times C$
for $t \in \{1, \dots, T\}$ **do**
 for $i \in U$ **do**
 $counter = 0$
 $\tilde{L}_{i,c}^{(t,k)} = L_{i,c}^{(t)}$ for all $c \in C$
 for $j \in J_i^{(t)}$ **do**
 $counter = counter + 1$
 for $c \in C$ **do**
 $m_{i,c,j}^{(t)} = \phi_j \min\{R_{i,c} - \tilde{L}_{i,c}^{(t,counter-1)}, R_{j,c}\}$
 $\tilde{L}_{i,c}^{(t,counter)} = \tilde{L}_{i,c}^{(t,counter-1)} + m_{i,c,j}^{(t)}$
 end for
 end for
 for $c \in C$ **do**
 $L_{i,c}^{(t+1)} = \alpha L_{i,c}^{(t)} + \left(\tilde{L}_{i,c}^{(t,counter)} - L_{i,c}^{(t)} \right)$
 end for
 end for
end for

general model could then be used for including these identity traits. For example, it could make users that are more alike (in terms of identity) affect each other more than users with very different character traits.

6.5 Extensions

An alternative to patches as a means to express preference for certain parts of the canvas would be a probability distribution over the canvas, assigning a density as where to drop paint. Using such a probability model implies that painters that carry similar probability distributions (which form a mathematical representation of their identity) will try to cover similar parts of the canvas, thereby increasing the odds that a smaller region will become painted. In that sense, being mentioned by a broad spectrum of identities would lead to a higher motivation (relating identity one-to-one to a probability distribution). However, we think that this is a far more difficult approach than the patch approach: we would have to relate identity to a probability distribution, which is hard to obtain from data. We would also have to find a way to sample an area based on the distribution for simulation purposes.

Combining the above ideas could lead to a model where person i is motivated the most when:

- He is mentioned by many people
- He is mentioned by people to whom he can relate
- He is mentioned by a variation of people

We note that the case $|C| > 1$ has only been introduced to show possible extensions of the model. In this research we shall restrict ourselves to using the $|C| = 1$ case for which we can also use the SD Recursions.

7 Donation potential model

The Satiation-Deprivation Model provides us with a mathematical framework to give each user a motivation level for each point in time based on Twitter mentions. Next, we need a framework to relate motivation to donation. We shall speak about the *donation potential* of a motivation level. This is the probability that a user will make a donation, given his motivation level.

We look at regions of motivation levels. For each of the regions we obtain a value related to the potential of the region when it comes to yielding donations. More specifically, let $U \times T$ be the set of measure points (users and time in days, both discrete). For each user-day pair $(i, t) \in U \times T$ we know the associated motivation level $L_i^{(t)} \in [0, 1]$ and we know whether or not a donation was made, given by the indicator $d_i^{(t)} \in \{0, 1\}$. The *data-induced* donation potential of a region $M \subseteq [0, 1]$ is defined to be:

$$D(M) = \frac{\sum_{(i,t) \in U \times T} d_i^{(t)} \mathbb{1}\{L_i^{(t)} \in M\}}{\sum_{(i,t) \in U \times T} \mathbb{1}\{L_i^{(t)} \in M\}}$$

Although the region M can be any subset of $[0, 1]$ we will only be concerned with cases where M is an interval. Note that $D(M)$ is only defined if there exists a pair (i, t) for which $L_i^{(t)} \in M$, so if there is at least one measure point with a motivation level in M .

Next we use a sliding-window approach to obtain a function defined for every motivation level, the data-induced donation potential of a point $m \in [0, 1]$.

$$\begin{aligned} M(m, \gamma) &:= \left[\max\{m - \frac{1}{2}\gamma, 0\}, \min\{m + \frac{1}{2}\gamma, 1\} \right] \\ D_\gamma(m) &:= \begin{cases} D(M(m, \gamma)) & \text{if } L_i^{(t)} \in M(m, \gamma) \text{ for some } (i, t) \\ 0 & \text{otherwise} \end{cases} \quad \text{for } m \in [0, 1] \end{aligned} \tag{17}$$

The choice of γ depends on the data set. It is a trade-off between data quantity and locality; the more data is available, the smaller γ could be chosen.

If motivation is highly positively correlated with donation potential, we expect D_γ to be monotonically increasing in m .

8 Matching procedure

We wish to match Twitter users to their donations, so that the Satiation-Deprivation Model and the donation-potential model can be applied in practice. We use the donation data, where each donation has a Movember ID and name associated. Furthermore we have the `DISPLAYNAME` and `PREFERREDUSERNAME` from every Twitter user. We look into the tweets to see whenever a Movember profile is linked to it. Next, we can check if the Twitter users and the Movember profile match on name (at least partly).

In the data of The Netherlands in 2014, we saw that only 19 persons provide a link to a Movember profile in their Tweets, of which only three carry a name similar to the Twitter account. 184 persons are referring to `nl.movember.com` overall. This indicates that it may be better to do the matching in another way.

Matching strategy: The objective is to match as many Twitter users as possible to their Movember profile. We can do this as follows:

- Match Twitter users in a 'firstname-lastname' fashion to Movember profiles. This means that we consider a Twitter profile and a Movember profile a match if the first and last name appear in both of them.
- Manually inspect the set of matched accounts. Discard if they are incorrect, keep if they seem to match. We also mark dubious ones, which we can include if the set of found matches is very small.
- Create a csv-file with Twitter profiles and Movember IDs. This file contains accepted matches.

We applied the above strategy on the data for The Netherlands. There are 10 586 Movember profiles. There are 4 204 Twitter users. We split the Movember names, like 'Jan de Wit', into lowercase strings: 'jan', 'de' and 'wit'. Then we look if there are Twitter users whose preferred username or displayname (both put in lowercase) contain the first or the last string. We included a 'Movember user-Twitter user' pair for further research if this was the case. We found 358 784 such user pairs. Many of them seem to be unrelated, other than the equality in substrings. In fact, many pairs can be excluded without looking at the profiles themselves at all. Think of the number of hits that 'jan' will have; a following last name can easily exclude most 'jan's from the manual inspection.

In order to restrict this number of 'matches', we demanded that the first and the last string should both be present in a Twitter account before we accept them. This still leaves 19 214 potential matching user pairs. We also notice that there are some ambiguous Movember profiles, such as with the name 'd' which gives many hits. We should leave those out. Allowing only comparison with strings with at least two symbols, we get only 862 potential matching user pairs. If we then, additionally, also require that the two strings are not equal (thus getting rid of persons like 'peter peter' or 'jan') then we are left with 471 users.

To summarize: The first rigorous match was based on whether the first and the last non-identical lowercase strings of a Movember profile were both present in either the (lowercase) preferred username or the (lowercase) display name on Twitter.

We manually inspect this set of users. The results can be found in figure 13. Here we explain the requirements for a pair to belong to some matching class:

- A match occurs if a Twitter user says he is providing a link to his own account, if the Twitter user is referring to the same team as the Movember user belongs to, or when the Twitter profile contains pictures of the same person as displayed on the Movember profile picture.

code	definition
match	Movember and Twitter profile have the same author
uncertain_match	It is uncertain, but not unlikely, that the Movember and Twitter profile have the same author. There might be a match on name, but the profile content does neither confirm or reject agreement of authorship.
match on name	Movember and Twitter profile have an author of the same name, but no additional information is available
no_match	Movember and Twitter profile do not have the same author or it is unreasonable to suppose they do.
inconclusive	The names are sort of similar, but neither the names nor the profiles give sufficient information

Table 4: Matching codes

- Uncertain matches are matches where for example there is a match on name and there is also a Movember profile available, yet comparison of the Movember and the Twitter profile does neither confirm nor reject the claim that these are indeed from the same person.
- A match on name occurs. However, we may not have the Movember profiles of these people, no picture and/or corresponding team name, so we cannot be a hundred percent certain that they are indeed profiles of the same person. It may be noted that many of these names are so specific, exotic at times, that it is highly likely that the profiles belong to a single author.
- A potential match is a 'no match' if we have compared the two profiles and concluded that these are not from the same author. However, we also name it a 'no match' if it is unreasonable to suppose that they belong to the same author, even though profile information is unavailable. For example, if a Movember user is named "jan, wit de", the raw match may find 'jandegroot' as a potential match, because it contains "jan" and "de". But even without profile comparison, it is evident that this is not a correct match.
- Inconclusive means that the names are similar, but neither the names nor the profiles (if found) give sufficient evidence that these profiles belong to a single author.

The codings we used for these options can be found in Table 4.

Sweden For Sweden we performed a similar matching procedure. Because there were about twice as much potential candidates found in the rigorous match (but also more of them were found to be inconclusive) we did not perform the entire comparison. Instead, we were satisfied with 56 matches users. This data set we used as an independent data set to make our results more robust.

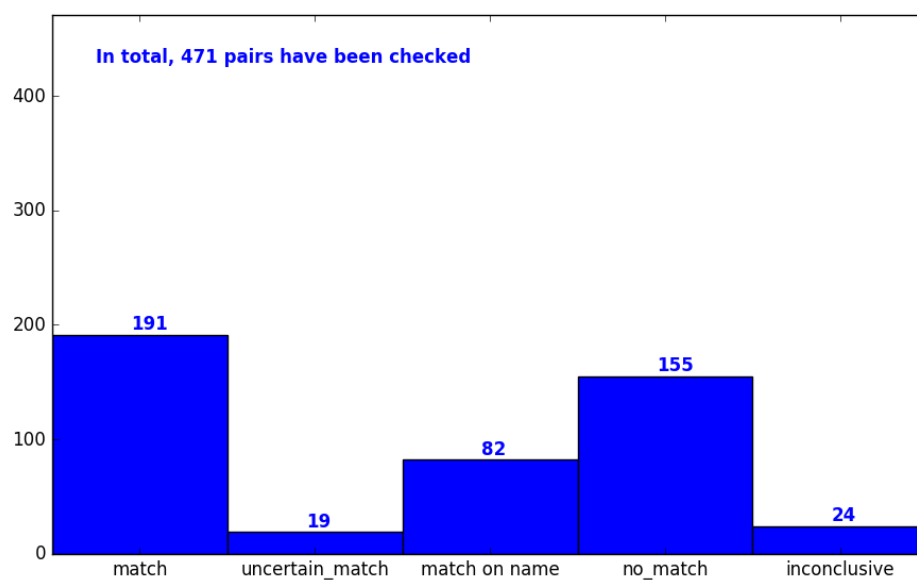


Figure 13: Matching results for The Netherlands 2014

9 Numerical results for Twitter and Movember data

9.1 Modeling ϕ for Twitter users

We have to decide on the parameters before applying the algorithm. Firstly, we choose $\alpha = 0.9$ to reflect that the effect of mentions will carry through time fairly strongly.

The definition of ϕ_i is a bit more involved. We choose to base the definition of ϕ_i on the number of followers of user i . Let the number of followers of user i be denoted by f_i . ϕ_i will now be defined as a normalized ranking based on f_i .

$$\phi_i := \frac{1}{|U|} \left(1 + \left| \{j \in U \mid f_j < f_i\} \right| + \frac{1}{2} \left| \{j \in U \mid f_j = f_i, j \neq i\} \right| \right)$$

The first '1' is added to make even the contribution of the least-followed person count. If there are ties, the ranking procedure considers half of them to have a smaller amount of followers and they all get the average rank.

Note that basing ϕ on the centralities in the mention network seems an unwise idea, because this is the same network as used for the Q -matrices and this could lead to undesirable interference. For example, in a star network the central mentioner will be important and boost many people's motivations at the same time. It is like stating that his mentions are important because he mentions so often.

9.2 Visual results for Satiation-Deprivation Model and donation potential

For the 191 matching users, we ran the Satiation-Deprivation Algorithm. We look at how often and when they were mentioned, see Figure 14 . Over time, we see how the motivation levels are evolving.

We also plotted the distributions of motivation levels for measure points associated with donations and for all measure points (donation-related or not), see Figure 15 for the absolute numbers and Figure 16 for the normalized numbers.

The graph for the data-induced donation potential for the Netherlands in 2014 can be found in Figure 17 . There is a distinct positive relation between motivation level and donation potential. Between motivation levels 0 and 1 there is about an 11% difference in donation potential.

We also made a graph for the data-induced donation potential for Sweden in 2014. This graph can be found in Figure 18. The positive relation is seen once again. For this data set the line appears to be steeper; there is about a 26% difference in donation potential between motivation levels 0 and 1.

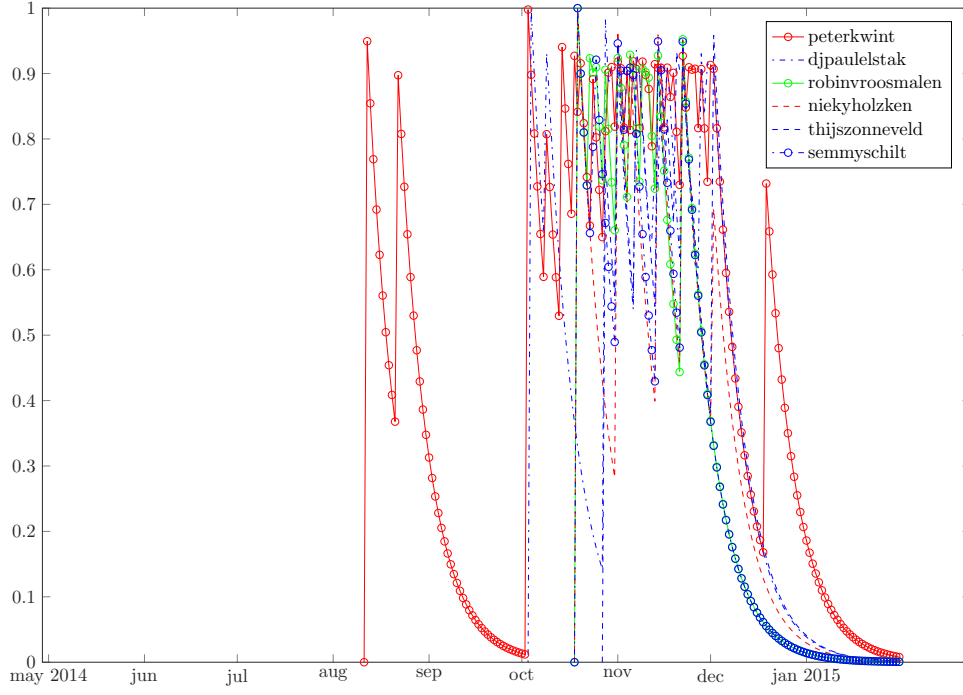


Figure 14: Satiation-Deprivation Model evolution for matched persons (191 persons considered, those with highest average level are shown). Users had to be perfectly matched to be included. $\alpha = 0.9$

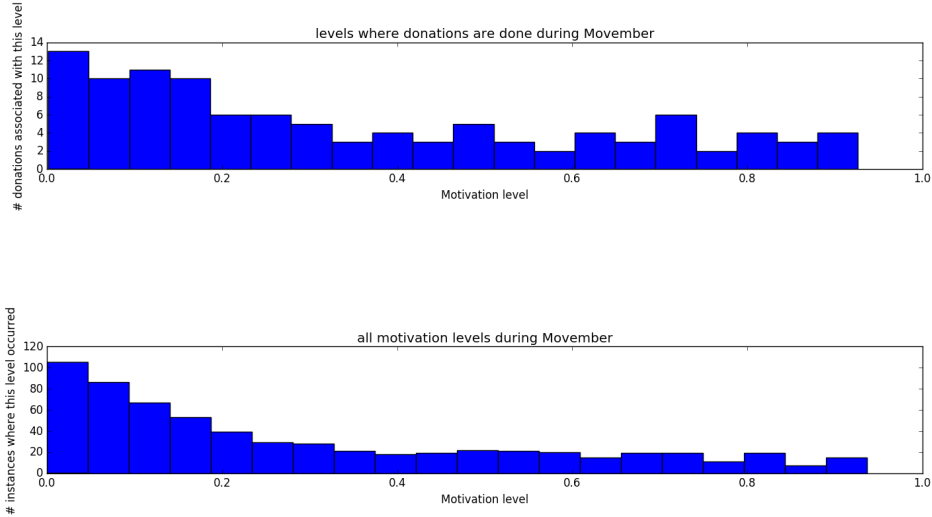


Figure 15: Scores according to the Satiation-Deprivation Model. In the upper histogram a motivation level is counted if the person under consideration made a donation at the same time point. In the lower histogram, all motivation levels are counted. Zero levels are excluded

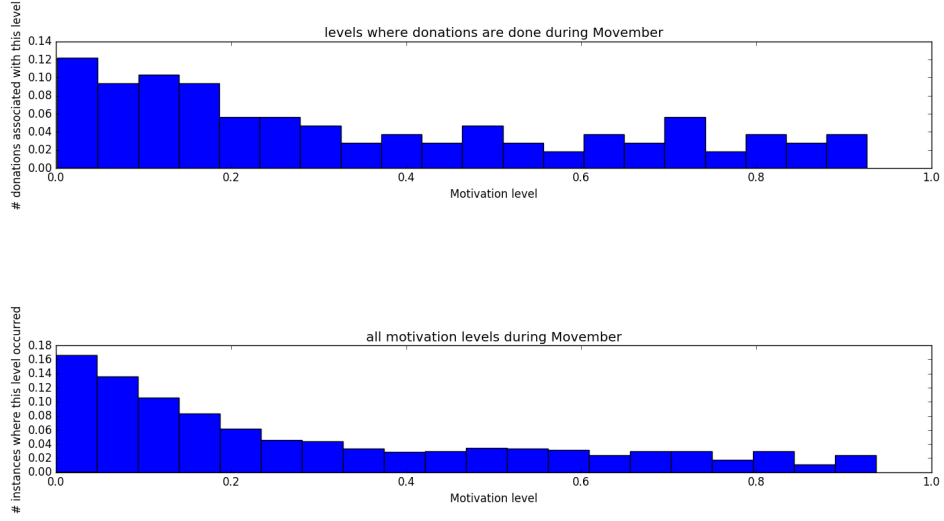


Figure 16: Scores according to the Satiation-Deprivation Model. In the upper histogram a motivation level is counted if the person under consideration made a donation at the same time point. In the lower histogram, all motivation levels are counted. Zero levels are excluded. The bars are normalized so that they total to length 1.

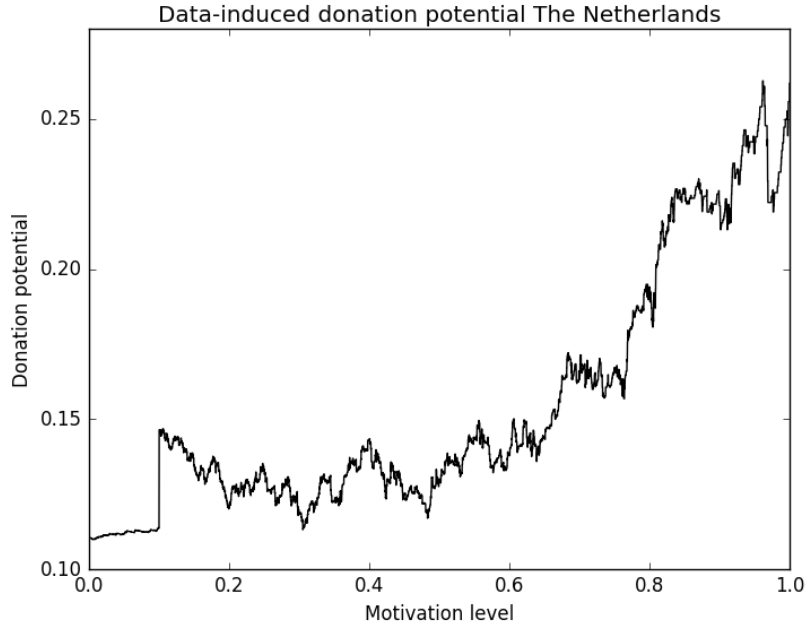


Figure 17: Data-induced donation potential of each motivation level for The Netherlands, derived from 191 users. Here we used a sliding-window approach. $D_{0.1}(x)$ is plotted. See Equation 17.

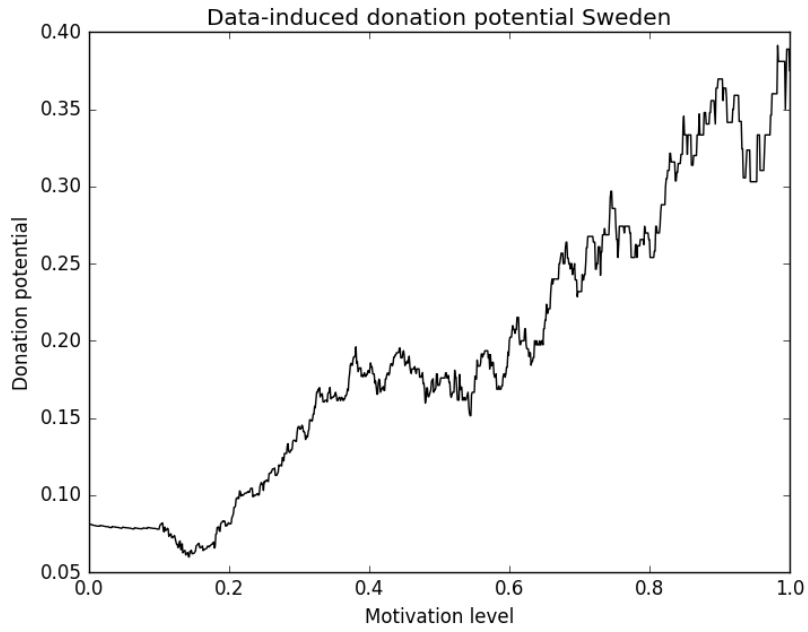


Figure 18: Data-induced donation potential of each motivation level for Sweden, derived from 56 users. Here we used a sliding-window approach. $D_{0.1}(x)$ is plotted. See Equation 17.

10 Statistical justification of Satiation-Deprivation Model

We have seen that there is a positive relation between motivation and donation if we follow the Satiation-Deprivation Model. In this part we wish to establish the contribution of the Satiation-Deprivation Model compared to baseline approaches. For the statistical tests we use a random set of 40 matched Twitter users in The Netherlands, 2014.

In Section 10.1 we use the Kolmogorov-Smirnov Test. This test shows that motivation levels and donation occurrences are related.

In Sections 10.2 and 10.3 we show that the SD Model performs better than a baseline approach in describing the positive relationship between mentions and donations. To this end we use randomization of mentions in the Movember period. This means that for each user we count the number of mentions made to him or her during November and reassign each mention with uniform probability to some day in November. We shall show in Section 10.2 that the number of co-occurrences of a donation to i and a mention to i on the same day could also have been achieved if mentions were randomly made, so if there was not truly a relationship between mentions and donation occurrence. On the other hand, we shall show that donations co-occur with high motivation levels more often than could be expected if mentions were randomized, so that a positive relationship between high motivation levels and donation occurrence is likely.

In Section 10.4 we investigate the effect of the order of mentioners on donation occurrence. We make permutations of mentioners and calculate the median motivation level at donation days of the mentionee. Observing that this median level can get as high as the median level for the original order of mentioners, we shall show that the influence of the mentioners' order is not statistically significant.

10.1 Goodness-of-Fit: The Kolmogorov-Smirnov Test

We use the Kolmogorov-Smirnov Test (see [5]) to see if motivation scores at donation dates could have been drawn from the same motivation distribution as the motivation levels at non-donation dates. Let M_D be the motivation distribution of donations and M_N the motivation distribution of all other time-user instances (none-donation). Then:

$$\begin{aligned} H_0 &: M_D \stackrel{d}{=} M_N \\ H_1 &: M_D \stackrel{d}{\neq} M_N \end{aligned}$$

Further we let n_D be the number of donation dates (all users combined) and n_N the number of non-donation dates of all users combined. Define as the test statistic:

$$D(n_D, n_N) = \sup_x |F_{n_D}(x) - G_{n_N}(x)|,$$

where F_{n_D} is the empirical distribution function of motivation scores found at donation instances, and G_{n_N} is the empirical distribution function of motivation scores found at all none-donation instances.

Using the table provided in the Appendix 26 we see that, as an approximation, H_0 should be rejected if:

$$D(n_D, n_N) > c(\alpha) \sqrt{\frac{n_D + n_N}{n_D n_N}}, \quad (18)$$

with $c(0.10) = 1.22$, $c(0.05) = 1.36$ and $c(0.01) = 1.63$.

In our data set, we have $n_D = 172$ and $n_N = 1028$. We find that $D(172, 1028) = 0.1585$ and $c(0.01) \sqrt{\frac{172+1028}{172 \cdot 1028}} = 0.1343$. So equation (18) is satisfied and the null hypothesis should be rejected,

even on the 1% level.

Using Python's `SCIPY.STATS.KS_2SAMP` we also find that the p-value is approximately 0.001, in agreement with the above calculations.

Based on this test, it is quite safe to reject the null hypothesis: apparently motivation scores related to donation instances are distributed significantly differently from the overall distribution. However, we cannot conclude from this test only that this means that the scores are higher, only that the distributions are different.

10.2 Co-occurrence of mentions and donations

The co-occurrence of mentions and donations is used as a baseline for the Satiation-Deprivation method. In this section we show that the Movember data suggests that donations are often, but not significantly often, made on the same day that the sponsor is mentioned. Therefore we use a quite simple mathematical model.

Let n_i be the number of mentions that user i obtains during the Movember period. T is the time horizon and equals 30 (days) for the analysis. Let D_i be the set of days on which a donation has been made by user i . Furthermore, $y_i = |D_i|$ is the number of days on which a donation is made to user i . Now let X_{ij} be the day that user i obtains mention $j \in \{1, \dots, n_i\}$. Suppose that X_{ij} is uniformly distributed over $\{1, \dots, T\}$. Let Z_{ij} be a co-occurrence indicator:

$$Z_{ij} = \begin{cases} 1 & \text{if } X_{ij} \in D_i \\ 0 & \text{otherwise} \end{cases}$$

So $Z_{ij} = 1$ if a mention occurs on one of the y_i donation days. Trivially $P(Z_{ij} = 1) = \frac{y_i}{T}$ if the uniformity assumption holds. Note that mentions are assumed to be made independently of one another within this model. Next, define the number of co-occurrences S to be:

$$S = \sum_{i \in U} \sum_{j \in [n_i]} Z_{ij}.$$

Then it can easily be seen that:

$$\begin{aligned} \mathbb{E}[S] &= \sum_{i \in U} \sum_{j \in [n_i]} \mathbb{E}[Z_{ij}] \\ &= \sum_{i \in U} \sum_{j \in [n_i]} \frac{y_i}{T} \\ &= \frac{1}{T} \sum_{i \in U} n_i y_i \end{aligned}$$

We shall use this expected number of co-occurrences for comparison with the true amount of co-occurrences in the data.

Implementation We regard a random set of 40 matched users and their parameters. Of course, we know how often they are mentioned and how many donations they made. Thus we can easily calculate $\mathbb{E}[S] = 17.1$. This is the expected number of donation-mention co-occurrences based on independent and uniformly distributed mentions. We construct a statistical test as follows. Let S_{Movember} be defined similarly as above, in terms of X_{ij} . The null hypothesis would be that the X_{ij} are independently and uniformly distributed over the campaign days. In this case we would

have $\mathbb{E}[S] = 17.1$. The alternative hypothesis is that mentions are more likely to occur on donation days. Formally, the hypotheses are as follows.

$$\begin{aligned} H_0 &: \mathbb{E}[S_{November}] = 17.1 \\ H_1 &: \mathbb{E}[S_{November}] > 17.1 \end{aligned}$$

In the data we find that the realized number of co-occurrences is 27. This suggests that donations are often made on the same day that the donator is mentioned.

We look whether this is also significant at the 10% significance level. By running the probabilistic null model a million times we find a p-value of about 0.164 for the event that the number of co-occurrences equals or exceeds 27. This means that our findings are not statistically significant on the 10% level, although they still suggest some correlation between the occurrence of donations and mentions. On a side note, with 30 co-occurrences there would have been statistical significance.

10.3 Co-occurrence of high motivation levels and donations

In this section we investigate whether high motivation levels co-occur with donations more often than could be expected by chance. The idea is again to keep the donations fixed. Under the null hypothesis, we assume that mentions are randomly distributed over the campaign days.

It is important to note that we disregard mentions made outside of the Movember period, even though these mentions can contribute to the motivation level of a user. It is far more likely that mentions are made in November than in any other month. Therefore, including days from October or December would in our eyes violate the uniformity assumption of the null hypothesis already too severely to make statistical results very reliable.

Let $median(\bar{R})$ be the median of the realization, where $\bar{R}$ is a list of motivation scores. A motivation level for user i at time t is included in the list if i made a donation at this day. We chose to work with medians here, although the mean might have worked fine as well. Notice that using a mean here would have no relation with the mean number of occurrences calculated in the baseline (see previous section), so the choice does not really affect the result.

Let X_{ij} be uniformly distributed over $\{1, \dots, 30\}$. If $X_{ij} = t$ this means that person i is being mentioned at day t by the j -th mention, where j is just a mention label. With each X_{ij} there is a ϕ_{ij} associated, the degree of influence of the j -th mentioner of i . We suppose that n_i mentions are being made to i , and for the simulations n_i is taken to be equal to the number of mentions in the data. Also the same ϕ 's are used, so ϕ_{ij} refers to the number of followers of some actual mentioner of i during the campaign period. The pseudocode can then be written as follows:

Algorithm 3 Pseudocode for co-occurrence of mentions and motivation levels

```

 $S_k = \emptyset$  for all  $k \in [N_{sim}]$ 
for  $k \in [N_{sim}]$  do
    Make  $k$ -th realization of  $\{X_{ij}\}$  of mention days
    Perform an SD-simulation according to this realization of mention days
    For each user and day, if the day is a donation day for this user, add motivation score to list  $S_k$ 
end for

```

The null hypothesis is that the $median(\bar{R})$ is a realization of $median(S)$ (a stochastic variable associated with the above process). The alternative hypothesis is that $median(\bar{R}) > median(S)$ stochastically.

$$\begin{aligned} H_0 &: median(\bar{R}) \stackrel{d}{=} median(S) \\ H_1 &: median(\bar{R}) > median(S) \text{ stochastically} \end{aligned}$$

At the 10% significance level we reject the null hypothesis if:

$$\frac{1}{N_{sim}} \sum_{k=1}^{N_{sim}} \mathbb{1}\{median(S_k) > median(\bar{R})\} < 0.10. \quad (19)$$

Implementation Implementing the above, with $N_{sim} = 10^4$ we find:

$$\frac{1}{N_{sim}} \sum_{k=1}^{N_{sim}} \mathbb{1}\{median(S_k) > median(\bar{R})\} = 0.034.$$

So the median of the realization of motivation scores gives a higher score than we would expect if the mentions were random. The found result is sufficient to refute the null hypothesis. So the motivation scores indeed seem to carry information about donation occurrence. Moreover, these scores seem more indicative of donation occurrence than just the number of mentions.

10.4 The effect of the order of mentioners on donation occurrence

Apart from the distribution of mentions over the timeline, another factor of influence on the motivation level is the order of the mentioners. Until now we have provided evidence that the timing of mentions influences the occurrence of donations. The next step is to see if the order of mentioners has a positive relation with donation occurrence. We use forty matched users from the data set of The Netherlands 2014 to check this statistically.

Keeping the timing of mentions fixed (and equal to what occurs in the data) we perform permutations on the order in which a person is mentioned by his different mentioners. The difference with our earlier methods is that the timing is now fixed and differences in motivation level can only occur because of differences in the order of mentioners. For example, if user i is mentioned by j and k , there are two orders in which he could be mentioned, both leading to a different motivation trajectory depending on the degrees of influence of j and k (namely ϕ_j and ϕ_k). The idea is that, given a donation date t_d of the user, the trajectory and accompanying order that has the highest motivation level at t_d is more likely to have occurred than the other orders, if our hypothesis is true.

We describe this in a statistical model, where the null hypothesis is that the median of donation-day motivation levels (according to the data) is a random draw from the medians resulting from all possible mentioner permutations. Here we require that all other data characteristics such as mention dates and donation dates are fixed. If we can reject this hypothesis (based on the p-value of some suitable test statistic), we have the alternative hypothesis that the order of the mentioners as found in the data gives rise to higher donation-day motivation levels than (most) other potential orders, and thus that the order significantly contributes to donation potential (and motivation). Suppose that R is a list of motivation levels, where a motivation level at a user-date instance is included if the user made a donation on this date. Because the data is subject to a stochastic process of mentioning, R is a stochastic list. For the data set of The Netherlands 2014, with 40 matched users, and including only the data of the month November, the realization of R is denoted by $\bar{R}$.

Furthermore, let S also be a list of motivation levels, where a motivation level at a user-date instance is included if the user made a donation on this date. However, in this case the mentioner order is uniformly randomized - although each mentionee still has the same set of mentioners. Let (S_k) be a sequence of such lists. Each list S_k is made by first drawing a sequence of mentioners from a uniform distribution over all possible mentioner orders for each of the forty users and consequently using the SD Model to calculate motivation levels.

The hypotheses are as follows:

$$\begin{aligned} H_0 &: \text{median}(S) \stackrel{d}{=} \text{median}(R) \\ H_1 &: \text{median}(S) < \text{median}(R) \text{ stochastically} \end{aligned}$$

The test statistic is defined to be:

$$t((\bar{S}_k)) = \frac{1}{N} \sum_{k=1}^N \mathbf{1} \{ \text{median}(\bar{S}_k) > \text{median}(\bar{R}) \}.$$

We reject H_0 if $t((\bar{S}_k)) < \alpha$ for some fixed $\alpha \in [0, 1]$. For our purposes we choose $\alpha = 0.10$. For our data where $\text{median}(\bar{R}) = 0.052$, we find $t((\bar{S}_k)) = 0.151$ for $N = 1000$. It appears suggestive: only 15 percent of the random orders give a median as large (at least as high) as the actual data does. It is, however, not statistically significant.

11 Discrete-Time Strategic Mentioning

We have established a positive relationship between motivation level and donation potential. The next question is what strategy the Movember Foundation could adopt on Twitter to increase donations. The goal in this and the next section is to increase the donation potential. Because of the positive relationship with motivation level, the goal becomes to maximize the average motivation level. We shall restate the research question, as given in Section 1.2 :

Research question 2: *What mentioning strategy should user i adopt to maximize the average motivation level of user j ?*

It seems trivial; i should mention j as often as possible. Indeed, following the model, every additional mention will increase (at least not decrease) the average motivation level. However, this is a highly unrealistic scenario. Not only does it take effort to mention someone, but also the credibility of the mentioner will decrease if he overdoes the job. Such a mentioner will be regarded as annoying rather than stimulating. Note that this is not part of the model, but rather common sense.

For this reason we restrict i 's mentions. He is allowed to make at most M mentions over the horizon $\{1, \dots, T-1\}$ and at most one every day. The latter requirement is quite natural. Let $\tau \in \{0, 1\}^{T-1}$ be a mentioning policy vector. We understand $\tau^{(t)} = 1$ to mean that user i mentions user j at time t . The goal is to find the optimal policy τ^{OPT} that maximizes $\sum_{t=1}^T L_j^{(t)}(\tau)$, where $L_j^{(1)} = \lambda$ is the given initial motivation level. $L_j^{(t)}$ is a function of τ as the motivation level depends on the mentioning strategy used. The optimization problems is as follows:

$$\begin{aligned}
 \tau^{OPT} &= \arg \max_{\tau \in \{0,1\}^T} \sum_{t=1}^T L_j^{(t)}(\tau) \\
 \text{subject to } &\sum_{t=1}^{T-1} \tau^{(t)} \leq M \\
 &L_j^{(1)} = \lambda \\
 &L_j^{(t+1)} = \left(\alpha - \Phi_j^{(t)}\right) L_j^{(t)} + \Phi_j^{(t)} \quad \forall t \in \{1, \dots, T-1\} \\
 &\Phi_j^{(t)} = 1 - \left(1 - \phi_i \tau^{(t)}\right) \prod_{k \in U-i} \left(1 - \phi_k Q_{k,j}^{(t)}\right) \quad \forall t \in \{1, \dots, T-1\}
 \end{aligned}$$

Some challenges with this definition of the problem immediately become clear. $Q^{(t)}$ is not known before time t , which makes deciding on a policy impossible in a direct and deterministic way. On the other hand, using the stochastic model would lead to a very complex model and at the same time it is not realistic to assume that we know the mentioning probabilities in practice.

For these reasons we assume mentions to be scarce in general. A policy is decided upon under the assumption that no mentions other than those of user i will occur in the future. The scarcity assumption is in part justified by the data, see Figure 19 and 20. It can be seen that most users do not get mentioned even once. If we look at users specifically mentioned by MOVEMBERNL we see that these people get mentioned more in general, see Figure 21. They might be more well-known users, like celebrities and politicians, that also caught the attention of Movember. We shall see that the scarcity assumption does play a role in the application of our results later on.

Although we ignore future mentions, we do take into account observed mentions made in the past. This will lead to an algorithm that can work online. It will calculate the optimal policy for the

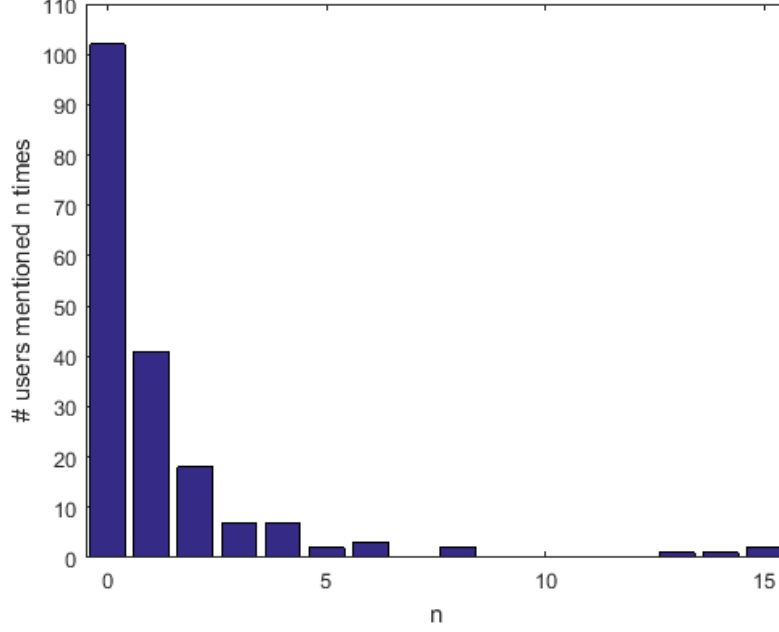


Figure 19: Histogram displaying how often users are mentioned during November in The Netherlands 2014. We looked at the 191 matched users.

remaining time until the horizon is reached, assuming no mentions of others will occur.

11.1 State space, action set and value function

The state space is formally defined as $S := \{0, \dots, T-1\} \times [0, 1] \times \{0, \dots, M\}$. The tuple $(x, L, m) \in S$ denotes the state where there are x time instances left to T (so it defines a state at time $T-x$), L is the motivation level and m is the number of remaining mentions. The action set is simply $A = \{0, 1\}$, where choosing action 1 means that i decides to mention j . For notational convenience we define $\phi = \phi_i$.

We define the optimal value function $V : S \rightarrow \mathbb{R}$ for $m \in \{1, \dots, x-1\}$. $V(x, L, m)$ gives the maximal value for $\sum_{t=T-x+1}^T L_j^{(t)}$ under the restriction that i can still give at most m mentions to j and given that $L_j^{(T-x)} = L$. $V(T, L, m) = \sum_{t=1}^T L_j^{(t)}$ is called the optimal total motivation level. It can be recursively defined by optimality equations, as follows:

$$V(x, L, m) := \alpha L + \max\{V(x-1, \alpha L, m), \phi(1-L) + V(x-1, \alpha L + \phi(1-L), m-1)\} \quad (20)$$

Note that the first argument in "max" corresponds to not making a mention at time $T-x$, whereas the second corresponds with making a mention. We still need to define boundary values for V so that it is completely defined. These boundaries are easily found for $m = 0$ and $m \geq x$, as shown in the next theorem.

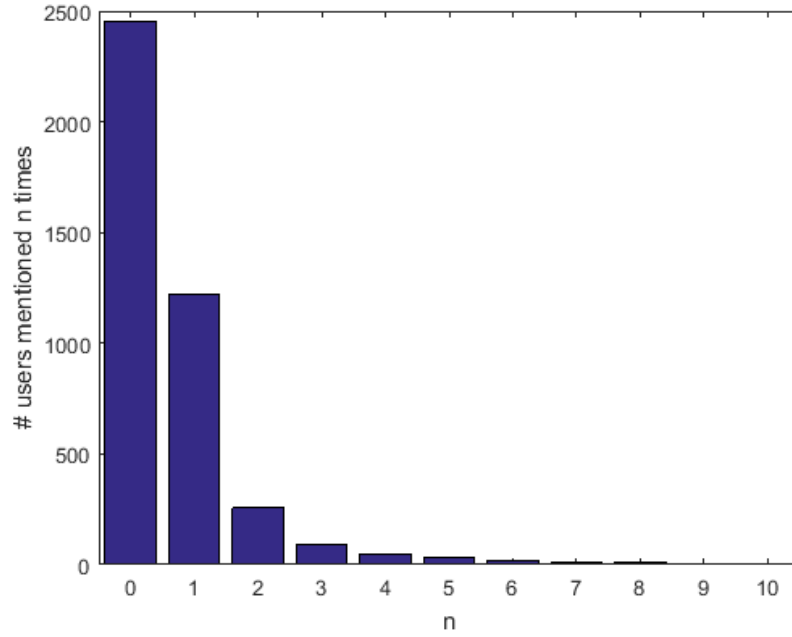


Figure 20: Histogram displaying how often users are mentioned during November in The Netherlands 2014. We looked at all 4204 Dutch users in the data set.

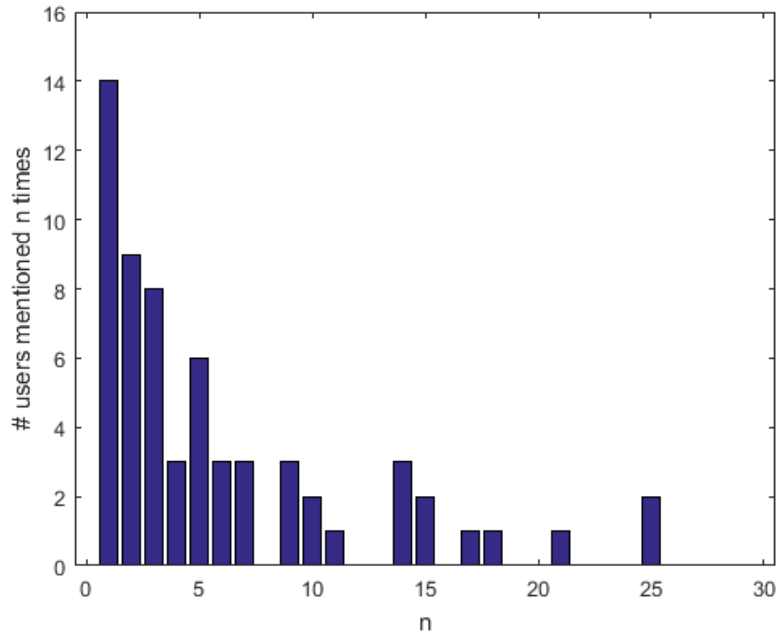


Figure 21: Histogram displaying how often users are mentioned during November in The Netherlands 2014. We looked at the 70 individuals that got mentioned at least once by MOVEMBERNL.

Theorem 12 *Let $K \in \mathbb{N}$. The following three formulas always hold for $\alpha < 1$:*

$$\begin{aligned} V(x, L, 0) &= L \frac{\alpha - \alpha^{x+1}}{1 - \alpha} \\ V(x, L, x) &= \sum_{k=1}^x \left[(\alpha - \phi)^k L + \phi \sum_{l=1}^k (\alpha - \phi)^{k-l} \right] \\ V(x, L, x + K) &= V(x, L, x) \end{aligned}$$

Proof: If we are in state $(x, L, 0)$ we have x time instances left where the action 0 is always made, since the action 1 is not available. This means that the motivation level will only decay at rate α :

$$\begin{aligned} \sum_{t=T-x+1}^T L_i^{(t)} &= (\alpha + \alpha^2 + \dots + \alpha^x) L \\ &= L \sum_{k=1}^x \alpha^k \\ &= L \frac{\alpha - \alpha^{x+1}}{1 - \alpha} \end{aligned}$$

If we are in state $(x, L, x + K)$ for $K \in \mathbb{N}$ we can mention someone every day. In the end we will be left with K mentions that we did not use, so we may just as well assume that we are in state (x, L, x) to begin with, as the value will be the same. Thus $V(x, L, x + K) = V(x, L, x)$. Then we have:

$$V(x, L, x) = (\alpha - \phi) L + \phi + V(x - 1, (\alpha - \phi) L + \phi, m - 1),$$

where $(\alpha - \phi) L + \phi$ is the motivation level of j at time $T - x + 1$, and the second term contains the sum of the motivation levels from $T - x + 2$ until T . Note that also $x - 1 = m - 1$ so that we can repeat the argument. In general, then, the level at time $T - x + n$ is recursively given by:

$$\begin{aligned} L_j^{(T-x+n)} &= (\alpha - \phi) L_j^{(T-x+n-1)} + \phi \\ L_j^{(T-x)} &= L \end{aligned}$$

The general solution for this first-order difference equation is:

$$L_j^{(T-x+n)} = (\alpha - \phi)^n L + \phi \sum_{l=1}^n (\alpha - \phi)^{n-l}.$$

Then it is trivially seen that:

$$V(x, L, x) = \sum_{k=1}^x L_j^{(T-x+k)} = \sum_{k=1}^x \left[(\alpha - \phi)^k L + \phi \sum_{l=1}^k (\alpha - \phi)^{k-l} \right]$$

This proves the theorem. ■

The theorem gives formulas for the cases $m = 0$ and $m \geq x$. In the first case there was only one policy available because i could not make mentions. In the latter case the optimal policy was trivial, namely to mention j every day until the horizon is reached.

Now, note how in the optimality equation (Equation 20) we either decrease x or we decrease x and m both by the same amount if we enter a recursion. So if $m < x$ to begin with, then eventually either $m = 0$ or $m = x$ occurs and in the meantime $m < x$ holds for every sub-recursion. We can therefore recursively define the function V and eventually we end up in a state for which we have a closed-form solution by Theorem 12. So V is now fully and uniquely defined by the optimality equations and the boundary values.

11.2 Brute-force algorithm for optimal strategies

Although it would be great to find a closed-form formula to calculate the optimal policy, we note that the state space is small enough to handle horizons of until length $T = 15$ reasonably using a brute-force approach. This means that we try all possible $\binom{T}{M}$ policies and keep the policy τ for which the total motivation level is maximal.

In the algorithm we not only keep λ and x as an input parameter, but α and ϕ as well, to allow for flexibility. Here we give the algorithm to calculate the value of a given policy, see Algorithm 4. Deciding on which policy is optimal is trivial: just pick the one that gives the highest value.

Algorithm 4 Value of a mentioning policy

Input: $x, \tau, \lambda, \alpha, \phi$
Output: $V(x, \tau, \lambda, \alpha, \phi)$

```

if  $1 \leq \sum_{t=1}^x \tau^{(t)} < x$  then
  if  $\tau_1 == 0$  then
     $v = \alpha L + V(x - 1, \tau(2 : x), \alpha L, \alpha, \phi)$ 
  else
     $v = \alpha L + \phi(1 - L) + V(x - 1, \tau(2 : x), \alpha L + \phi(1 - L), \alpha, \phi)$ 
  end if
else
  if  $\sum_{t=1}^x \tau^{(t)} == 0$  then
     $v = L \frac{\alpha - \alpha^{x+1}}{1 - \alpha}$ 
  else
     $v = \sum_{k=1}^x \left( (\alpha - \phi)^k L + \phi \sum_{l=1}^k (\alpha - \phi)^{k-l} \right)$ 
  end if
   $V(x, \tau, \lambda, \alpha, \phi) = v$ 
end if

```

11.3 Results

Simulations One might expect it to be optimal to distribute mentions uniformly over the time horizon. However, using simulations, the algorithm shows that even if $\lambda = 0$, this is not always the case. In many cases the uniform mention distribution was found to yield quite high values, but for example the case $T = 6$, $M = 3$, $\lambda = 0$, $\alpha = 0.9$ and $\phi = 0.4$ comes with the optimal policy $\pi^{OPT} = (1 \ 1 \ 1 \ 0 \ 0 \ 0)$ and forms a counterexample for the statement.

Applying the on-line strategy We regarded 70 users that were mentioned by MOVEMBERNL during November. Denote this set of users by U_M . We set $T = 30$, with the time points designating the days of November. We call the set of mentioning sequences in the data the *data strategy*, which can be seen as a baseline. Applying the on-line strategy means that for each person $i \in U_M$ we redistribute the number of received mentions from MOVEMBERNL over the month. Mentions from other persons are assumed to be independent from this redistribution and remain at their original days. We used the data from The Netherlands 2014 to see how the on-line strategy performs. For the data strategy we found that:

$$\frac{1}{|U_M|} \sum_{i \in U_M} \frac{1}{T} \sum_{t=1}^T L_i^{(t)} = 0.3947,$$

where t indexes the days of November. This is the average motivation level over all users in U_M and over all days of November 2014.

Next, we start applying the on-line strategy, where we use that $\phi_{NovemberNL} = 0.7686$ and $\alpha = 0.9$. First we decide on a horizon $h \in \mathbb{N}$. Then, for all days $t \in \{1, 2, \dots, 30\}$ we regard the period $\{t, t+1, \dots, \min(t+h, 30)\}$. Per user i , we look at how many mentions are still available for i , what is the current motivation level $L_i^{(t)}$ (used as λ) and then we determine an optimal strategy τ^{OPT} based on this information. To this end we assume absence of future external mentions (mentions from others) and make use of Algorithm 4. If $\tau_1^{OPT} = 1$, then we choose to make a mention at time t and we reduce the number of available mentions for i by one. If $\tau_1^{OPT} = 0$, then we do not make a mention at time t and the number of available mentions for i remains the same.

Choosing h large has both advantages and disadvantages. The advantage is that we take into account a larger period of time, so we are less myopic in our decision-making. For example, for $h = 1$, we shall always make a mention now if one is available, because it maximizes the $L_i^{(t+1)}$. However, if we choose h to be large ($h = 30$, for example), then not only does the computation time increase, but there is also a higher probability that other mentions occur in the period that is being considered. The algorithm does not take these mentions into account and seemingly optimal mentioning strategies could actually lead to a worse realized average motivation level if the scarcity assumption is violated.

We define two performance parameters, depending on h , that give us insight in the performance of Algorithm 4 on the data set. For each h , set $x = h$. The strategy that follows from the algorithm is denoted by $\tau^{str}(h)$. The performance parameter compares $\tau^{str}(h)$ with the data strategy τ^{data} . Let $D_{0.1}(\cdot)$ be the donation potential function with $\gamma = 0.1$. For our purposes we set $T = 30$, where $t = 1$ corresponds with November 1st.

$$\begin{aligned}\Delta^L(h) : &= \frac{1}{|U_M|} \sum_{i \in U_M} \frac{1}{T} \sum_{t=1}^T \left[L_i^{(t)}(\tau^{str}(h)) - L_i^{(t)}(\tau^{data}) \right] \\ \Delta^D(h) : &= \frac{1}{|U_M|} \sum_{i \in U_M} \frac{1}{T} \sum_{t=1}^T \left[D_{0.1}(L_i^{(t)}(\tau^{str}(h))) - D_{0.1}(L_i^{(t)}(\tau^{data})) \right]\end{aligned}$$

$\Delta^L(h)$ denotes the average increase in average motivation level with respect to strategy τ^{data} by using strategy $\tau^{str}(h)$. $\Delta^D(h)$ denotes the average increase in average donation potential with respect to strategy τ^{data} by using strategy $\tau^{str}(h)$. To calculate the donation potential, we refer to Equation 17.

In Figures 22 and 23 we plot $\Delta^L(h)$ and $\Delta^D(h)$ respectively. We see that using a mentioning strategy algorithm indeed leads to a higher average motivation level than the data strategy. Surprisingly, even using a horizon $h = 1$ leads to an increase; this is unexpected, because this strategy just entails making all mentions in the first days of the campaign (with the restriction of making one mention per day at most). As our decision-making takes into account larger horizons, and thus becomes less myopic, the values become slightly better. However, making the horizon larger than two weeks is even slightly detrimental to the obtained values, possibly because the scarcity assumption is too badly violated. We therefore suggest to choose horizons of about two weeks. This gives the best performance according to Figure 22 and it is computationally not very expensive (for example, with $M = 2$ there are $\binom{14}{2} = 91$ policies to compare).

In Figure 23 we see how using different horizons affects the average individual increase in average donation potential. Remarkably, choosing horizons larger than $h = 2$ does not seem to lead to improvements in the donation potential. Overall, an improvement of about 19% in donation potential could be reached by strategic mentioning. This means that, if the brute-force mentioning strategy was used, the average user in U_M would have been 19% more likely to make a donation at a given campaign day, than he was in the realized setting (where the data strategy was used).

In the next section we explore strategies in a continuous-time setting and we shall find that the optimal distribution of mentions largely depends on ϕ .

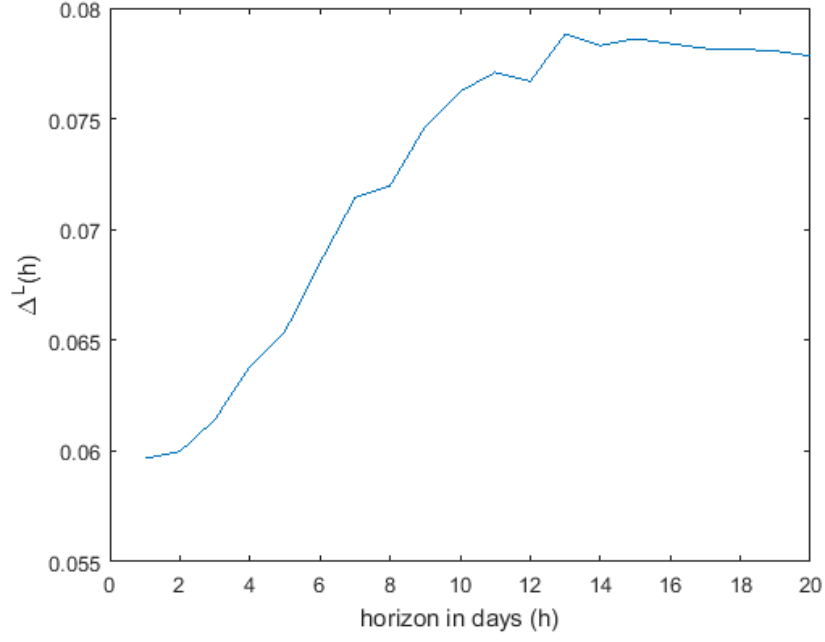


Figure 22: Average individual increase in average value by using an on-line mentioning strategy, with respect to different horizons h . The set of users consists of those mentioned by MOVEMBERNL.

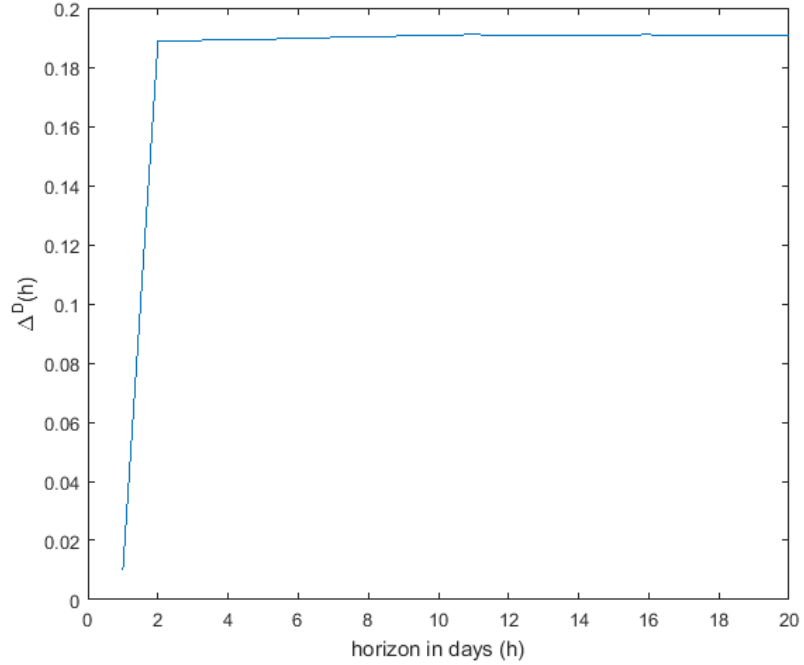


Figure 23: Average individual increase in average donation potential by using an on-line mentioning strategy, with respect to different horizons h . The set of users consists of those mentioned by MOVEMBERNL.

12 Continuous-Time Strategic Mentioning

Until now we have always supposed that a single mention can be given any day. This is a discrete time decision model. A natural extension is to look at a continuous time decision process. A user has M mentions and a continuous time horizon $[0, T]$ and wishes to maximize the average motivation level by giving the mentions at strategic points in time. For the current purposes, a continuous version of the Satiation-Deprivation Model is easily defined, mainly because no other mentions are assumed. The decay is captured by assuming exponential decay with rate α per day. We shall assume $\phi > 0$ throughout this section, as for $\phi = 0$ the strategy does not influence the motivation level.

The total motivation level is now an integral:

$$V := \int_{t=0}^T L(t) dt.$$

In the absence of mentions (or for $\phi = 0$), the initial motivation λ yields the following total motivation level:

$$V = \lambda \int_{t=0}^T \alpha^t dt.$$

If a user with degree of influence ϕ mentions the user at times $\tau = (\tau_1 \ \tau_2 \ \dots \ \tau_M)$, the formula is more complex. Let $g_k(\tau_k) := \lim_{t \rightarrow \tau_k^-} L(t)$ be the motivation level just before τ_k (or equivalently, the level at τ_k if the user were not mentioned at that time point). Next, $f_k(\tau_k) := \lim_{t \rightarrow \tau_k^+} L(t) + \phi(1 - \lim_{t \rightarrow \tau_k^-} L(t))$ is the motivation level just after τ_k . We also introduce the notations $\tau_0 = 0$, $f_0(\tau_0) = \lambda$ and $\tau_{M+1} = T$. With $L(0) = \lambda$, we can easily see that:

$$\begin{aligned} g_k(\tau_k) &= f_{k-1}(\tau_{k-1}) \alpha^{\tau_k - \tau_{k-1}} \quad \text{for all } k \in \{1, \dots, M\} \\ f_k(\tau_k) &= g_k(\tau_k) + \phi(1 - g_k(\tau_k)) \quad \text{for all } k \in \{1, \dots, M\} \end{aligned}$$

The value of mentioning strategy τ for $\alpha > 0$ is then given by:

$$\begin{aligned} V(\tau) &= \sum_{m=0}^M f_m(\tau_m) \int_{t=\tau_m}^{\tau_{m+1}} \alpha^{t-\tau_m} dt \\ &= \frac{1}{\ln(\alpha)} \sum_{m=0}^M f_m(\tau_m) (\alpha^{\tau_{m+1}-\tau_m} - 1) \end{aligned}$$

The optimal mentioning strategy is the maximizer of the above function.

For $k = 1 \dots M$, the partial derivatives of $V(\tau)$ are given by:

$$\begin{aligned} \frac{\partial V(\tau)}{\partial \tau_k} &= \frac{1}{\ln(\alpha)} \left[\frac{\partial}{\partial \tau_k} f_{k-1}(\tau_{k-1}) (\alpha^{\tau_k - \tau_{k-1}} - 1) + \frac{\partial}{\partial \tau_k} f_k(\tau_k) (\alpha^{\tau_{k+1} - \tau_k} - 1) \right] \\ &= \frac{1}{\ln(\alpha)} \left[f_{k-1}(\tau_{k-1}) \alpha^{\tau_k - \tau_{k-1}} \ln(\alpha) + (\alpha^{\tau_{k+1} - \tau_k} - 1) \frac{\partial}{\partial \tau_k} f_k(\tau_k) - f_k(\tau_k) \alpha^{\tau_{k+1} - \tau_k} \ln(\alpha) \right] \\ &= f_{k-1}(\tau_{k-1}) \alpha^{\tau_k - \tau_{k-1}} - f_k(\tau_k) \alpha^{\tau_{k+1} - \tau_k} + \frac{1}{\ln(\alpha)} (\alpha^{\tau_{k+1} - \tau_k} - 1) \frac{d}{d\tau_k} f_k(\tau_k). \end{aligned}$$

12.1 Optimal mentioning for $M = 1$

For the case that $M = 1$, we can now find the optimal assignment of τ_1 . First, suppose that $\lambda = 0$. Then $f_0(\tau_0) = 0$ and $f_1(\tau_1) = \phi$. This means that:

$$\begin{aligned} V(\tau) &= \frac{1}{\ln(\alpha)} \phi (\alpha^{T-\tau_1} - 1) \\ \frac{\partial V(\tau)}{\partial \tau_1} &= \frac{1}{\ln(\alpha)} \phi (-\alpha^{T-\tau_1} \ln(\alpha)) \\ &= -\phi \alpha^{T-\tau_1} \end{aligned}$$

This means that $\frac{\partial V(\tau)}{\partial \tau_1} < 0$ always, so $\tau_1 = 0$ is optimal. If $\lambda > 0$, then we have:

$$\begin{aligned} f_0(\tau_0) &= \lambda \\ g_1(\tau_1) &= \lambda \alpha^{\tau_1} \\ f_1(\tau_1) &= \lambda \alpha^{\tau_1} + \phi (1 - \lambda \alpha^{\tau_1}) \\ \frac{d}{d\tau_k} f_k(\tau_k) &= \ln(\alpha) (1 - \phi) \lambda \alpha^{\tau_1} \\ \frac{\partial V(\tau)}{\partial \tau_1} &= \lambda \alpha^{T-1} - f_1(\tau_1) \alpha^{T-\tau_1} + \frac{1}{\ln(\alpha)} (\alpha^{T-\tau_1} - 1) \frac{d}{d\tau_k} f_k(\tau_k) \\ &= \lambda \alpha^{\tau_1} - [(1 - \phi) \lambda \alpha^{T-1} + \phi] \alpha^{T-\tau_1} + (\alpha^{T-\tau_1} - 1) (1 - \phi) \lambda \alpha^{\tau_1} \\ &= \lambda \alpha^{\tau_1} - (1 - \phi) \lambda \alpha^T - \phi \alpha^{T-\tau_1} + (1 - \phi) \lambda \alpha^T - (1 - \phi) \lambda \alpha^{\tau_1} \\ &= \lambda \alpha^{\tau_1} - \phi \alpha^{T-\tau_1} - (1 - \phi) \lambda \alpha^{\tau_1} \\ &= \phi \lambda \alpha^{\tau_1} - \phi \alpha^{T-\tau_1} \end{aligned}$$

Setting $\frac{\partial V(\tau)}{\partial \tau_1} = 0$, then:

$$\begin{aligned} \phi \lambda \alpha^{\tau_1^*} &= \phi \alpha^{T-\tau_1^*} \\ \lambda \alpha^{\tau_1^*} &= \alpha^{T-\tau_1^*} \\ \lambda (\alpha^{\tau_1^*})^2 &= \alpha^T \\ (\alpha^{\tau_1^*})^2 &= \frac{\alpha^T}{\lambda} \\ \alpha^{\tau_1^*} &= \sqrt{\frac{\alpha^T}{\lambda}} \\ \tau_1^* &= \log_\alpha \sqrt{\frac{\alpha^T}{\lambda}} \\ &= \frac{1}{2} (T - \log_\alpha(\lambda)) \end{aligned}$$

We also calculate the second derivative to see that we indeed have a maximum:

$$\frac{\partial^2 V(\tau)}{\partial^2 \tau_1} = \ln(\alpha) (\phi \lambda \alpha^{\tau_1} + \phi \alpha^{T-\tau_1})$$

Because $\ln(\alpha) < 0$ and $\phi \lambda \alpha^{\tau_1} + \phi \alpha^{T-\tau_1} > 0$ always, we have $\frac{\partial^2 V(\tau)}{\partial^2 \tau_1} < 0$ always.

Note that if $\lambda = 1$, then $\tau_1^* = \frac{1}{2}T$.

Also note that it could be that $\tau_1^* < 0$. This violates the restriction that $\tau_1 > 0$ always, so this critical point of V is not the optimal point if we restrict τ_1 to $[0, T]$. Because the unique maximum

is obtained to the left of zero, the value function must be decreasing on $[0, T]$, so $\tau_1 = 0$ maximizes $V(\tau)$.

This means that the maximizing τ_1 is always given by the formula:

$$\tau_1^{OPT} = \begin{cases} \max\{0, \frac{1}{2}(T - \log_\alpha(\lambda))\} & \text{if } \lambda > 0 \\ 0 & \text{if } \lambda = 0 \end{cases}$$

So if $\lambda = 1$ then it is optimal to mention at $\frac{1}{2}T$, otherwise the mention should be made before this time.

It is worth noting that the optimal strategy does not depend on the degree of influence of the mentioner in this case.

12.2 Optimal mentioning for $M = 2$

For $M = 2$ we shall aim to find optimality equations. The goal is to lay the foundations for generalization to general M . For every given τ_2 we have to place τ_1 somewhere in the interval $[0, \tau_2]$. The derivation will go just as above and the optimality equation for τ_1^{OPT} equals:

$$\tau_1^{OPT} = \begin{cases} \max\{0, \frac{1}{2}(\tau_2^{OPT} - \log_\alpha(\lambda))\} & \text{if } \lambda > 0 \\ 0 & \text{if } \lambda = 0 \end{cases}$$

We also want a second equation so that we obtain a system of two unknowns and two equations. Therefore we use $\frac{\partial V(\tau)}{\partial \tau_2}$.

For $\lambda = 0$, we obtain:

$$\begin{aligned} f_1(\tau_1) &= \phi \\ g_2(\tau_2) &= \phi \alpha^{\tau_2 - \tau_1} \\ f_2(\tau_2) &= \phi \alpha^{\tau_2 - \tau_1} + \phi (1 - \phi \alpha^{\tau_2 - \tau_1}) \\ &= (1 - \phi) \phi \alpha^{\tau_2 - \tau_1} + \phi \\ \frac{d}{d\tau_2} f_2(\tau_2) &= \ln(\alpha) (1 - \phi) \phi \alpha^{\tau_2 - \tau_1} \\ \frac{\partial V(\tau)}{\partial \tau_2} &= \phi \alpha^{\tau_2 - \tau_1} - [(1 - \phi) \phi \alpha^{\tau_2 - \tau_1} + \phi] \alpha^{T - \tau_2} + (\alpha^{T - \tau_2} - 1) (1 - \phi) \phi \alpha^{\tau_2 - \tau_1} \\ &= \phi \alpha^{\tau_2 - \tau_1} - \phi \alpha^{T - \tau_2} - (1 - \phi) \phi \alpha^{\tau_2 - \tau_1} \\ &= \phi^2 \alpha^{\tau_2 - \tau_1} - \phi \alpha^{T - \tau_2} \end{aligned}$$

Setting $\frac{\partial V(\tau)}{\partial \tau_2} = 0$ we can find the critical point τ_2^* given τ_1 :

$$\begin{aligned} \alpha^{\tau_2^* - \tau_1} &= \frac{1}{\phi} \alpha^{T - \tau_2^*} \\ \alpha^{\tau_2^* - \tau_1} &= \alpha^{T - \tau_2^* - \log_\alpha(\phi)} \\ \tau_2^* &= \frac{1}{2} (T + \tau_1 - \log_\alpha(\phi)) \end{aligned}$$

Since $\tau_1^{OPT} = 0$ (because $\lambda = 0$), we obtain $\tau_2^* = \frac{1}{2} (T - \log_\alpha(\phi))$. For $\phi = 1$, this means the optimal strategy is $\tau = (0, \frac{1}{2}T)$. This strategy is also optimal for $\alpha = \phi$. However, if $\alpha \neq \phi$, then the second mention point should be made earlier. Notice that $\tau_2^* < 0$ if $\phi < \alpha^T$. Under the assumption that the critical point is indeed a maximum, we should choose $\tau_2^{OPT} = 0$ in these cases.

For $\lambda > 0$ we obtain:

$$\begin{aligned}
 \frac{\partial V(\tau)}{\partial \tau_2} &= \lambda \alpha^{\tau_2} + \phi \alpha^{\tau_2 - \tau_1} - \phi \lambda \alpha^{\tau_2} - (1 - \phi)^2 \lambda \alpha^T - (1 - \phi) \phi \alpha^{T - \tau_1} - \phi \alpha^{T - \tau_2} \\
 &\quad + (\alpha^{T - \tau_2} - 1) ((1 - \phi)^2 \lambda \alpha^{\tau_2} + (1 - \phi) \phi \alpha^{\tau_2 - \tau_1}) \\
 &= \lambda \alpha^{\tau_2} + \phi \alpha^{\tau_2 - \tau_1} - \phi \lambda \alpha^{\tau_2} - (1 - \phi)^2 \lambda \alpha^T - (1 - \phi) \phi \alpha^{T - \tau_1} - \phi \alpha^{T - \tau_2} \\
 &\quad + (1 - \phi)^2 \lambda \alpha^T + (1 - \phi) \phi \alpha^{T - \tau_1} - (1 - \phi)^2 \lambda \alpha^{\tau_2} - (1 - \phi) \phi \alpha^{\tau_2 - \tau_1} \\
 &= \lambda \alpha^{\tau_2} + \phi \alpha^{\tau_2 - \tau_1} - \phi \lambda \alpha^{\tau_2} - \phi \alpha^{T - \tau_2} - (1 - \phi)^2 \lambda \alpha^{\tau_2} - (1 - \phi) \phi \alpha^{\tau_2 - \tau_1}
 \end{aligned}$$

Setting $\frac{\partial V(\tau)}{\partial \tau_2} = 0$, we can find the critical point:

$$\begin{aligned}
 (\lambda + \phi \alpha^{-\tau_1} - \phi \lambda - (1 - \phi)^2 \lambda - (1 - \phi) \phi \alpha^{-\tau_1}) \alpha^{\tau_2} &= \phi \alpha^{T - \tau_2} \\
 (\lambda + \phi \alpha^{-\tau_1} - \phi \lambda - (1 - \phi)^2 \lambda - (1 - \phi) \phi \alpha^{-\tau_1}) (\alpha^{\tau_2})^2 &= \phi \alpha^T \\
 (\alpha^{\tau_2})^2 &= \frac{\phi \alpha^T}{(1 - \phi) \lambda + \phi^2 \alpha^{-\tau_1} - (1 - \phi)^2 \lambda} \\
 &= \frac{\phi \alpha^T}{-\phi \lambda + \phi^2 \alpha^{-\tau_1} + 2\phi \lambda - \phi^2 \lambda} \\
 &= \frac{\alpha^T}{-\lambda + \phi \alpha^{-\tau_1} + 2\lambda - \phi \lambda} \\
 &= \frac{\alpha^T}{(1 - \phi) \lambda + \phi \alpha^{-\tau_1}} \\
 \alpha^{\tau_2} &= \sqrt{\frac{\alpha^T}{(1 - \phi) \lambda + \phi \alpha^{-\tau_1}}}.
 \end{aligned}$$

This means that the unique critical point, for given τ_1 , is given by:

$$\tau_2^* = \frac{1}{2} [T - \log_\alpha ((1 - \phi) \lambda + \phi \alpha^{-\tau_1})]$$

However, $\tau_2^* < 0$ occurs whenever $(1 - \phi) \lambda + \phi \alpha^{-\tau_1} > \alpha^T$. Assuming that the critical point is a maximum, we should choose $\tau_2^{OPT} = 0$ in these cases.

Note that this means that for $\phi = 1$, we have $\tau_2^* = \frac{1}{2}(T + \tau_1)$ so the second mention should be made halfway between τ_1 and T .

It also means that $\tau_1^{OPT} = \max\{0, \frac{1}{4}T + \frac{1}{4}\tau_1^{OPT} - \log_\alpha(\lambda)\}$, implying $\tau_1^{OPT} = \frac{1}{3}T - \frac{4}{3}\log_\alpha(\lambda)$. For $\lambda = 1$ we thus obtain $\tau^{OPT} = (\frac{1}{3}T, \frac{2}{3}T)$, given that $\phi = 1$. Note that $\tau_1^{OPT} < \frac{1}{3}T$ in case that $\alpha \neq \lambda$. In general, the optimal solution for $\lambda = 0$ is given by:

$$\begin{cases} \tau_1^{OPT} = 0 \\ \tau_2^{OPT} = \begin{cases} \frac{1}{2}(T - \log_\alpha(\phi)) & \text{if } \phi \geq \alpha^T \\ 0 & \text{if } \phi < \alpha^T \end{cases} \end{cases}$$

The optimal solution for $\lambda > 0$ is more involved:

$$\begin{cases} \tau_1^{OPT} = \begin{cases} \frac{1}{2}(\tau_2^{OPT} - \log_\alpha(\lambda)) & \text{if } \lambda \leq \alpha^{\tau_2^{OPT}} \\ 0 & \text{if } \lambda > \alpha^{\tau_2^{OPT}} \end{cases} \\ \tau_2^{OPT} = \begin{cases} \frac{1}{2} [T - \log_\alpha ((1 - \phi) \lambda + \phi \alpha^{-\tau_1^{OPT}})] & \text{if } (1 - \phi) \lambda + \phi \alpha^{-\tau_1^{OPT}} > \alpha^T \\ 0 & \text{if } (1 - \phi) \lambda + \phi \alpha^{-\tau_1^{OPT}} \leq \alpha^T \end{cases} \end{cases},$$

but it might be easier to compute by not explicitly describing the restrictions:

$$\begin{cases} \tau_1^{OPT} = \max\{0, \frac{1}{2}(\tau_2^{OPT} - \log_\alpha(\lambda))\} \\ \tau_2^{OPT} = \max\{0, \frac{1}{2} [T - \log_\alpha ((1 - \phi) \lambda + \phi \alpha^{-\tau_1^{OPT}})]\} \end{cases}$$

This system can be easily solved by available solvers in Matlab. We used FSOLVE for this.

12.3 Optimal mentioning for general M

The aim in this section is to find optimality equations for general M . We first regard the case $\phi = 1$ in this section, which is a special case for which direct formulas can be found. For the case $\phi < 1$ we make use of the structure of the derivations of the previous sections.

12.3.1 The case $\phi = 1$

The case $\phi = 1$ turns out to be a simpler case.

By definition of f_k , it holds that $f_k(\tau_k) = 1$ for all τ_k and $k > 0$. Then it also holds that $\frac{d}{d\tau_k} f_k(\tau_k) = 0$. So then, for $k > 1$:

$$\frac{\partial V(\tau)}{\partial \tau_k} = \alpha^{\tau_k - \tau_{k-1}} - \alpha^{\tau_{k+1} - \tau_k}.$$

Setting $\frac{\partial V(\tau)}{\partial \tau_k} = 0$ we obtain:

$$\begin{aligned} \alpha^{\tau_k^* - \tau_{k-1}} &= \alpha^{\tau_{k+1} - \tau_k^*} \\ \tau_k^* - \tau_{k-1} &= \tau_{k+1} - \tau_k^* \\ \tau_k^* &= \frac{1}{2} (\tau_{k+1} + \tau_{k-1}). \end{aligned}$$

However, for $k = 1$, we have $f_0(\tau_0) = \lambda$. For $\lambda = 0$, we obtain $\frac{\partial V(\tau)}{\partial \tau_1} = -\alpha^{\tau_2 - \tau_1}$. This function is always negative, so for every τ_2 it is optimal to choose $\tau_1 = 0$. We then find a critical point by solving:

$$\begin{cases} \tau_1^* = 0 \\ \tau_k^* = \frac{1}{2} (\tau_{k+1}^* + \tau_{k-1}^*) & \text{for } k \geq 1 \\ \tau_{M+1}^* = T \end{cases}$$

It is clear that a solution to the system is given by:

$$\tau_k^* = \frac{k-1}{T}. \quad (21)$$

Since the system consists of $M+1$ unknowns and $M+1$ linearly independent equations, the given solution is also unique. It means that the strategy where mentions are made uniformly over time is a critical point.

For $\lambda > 0$:

$$\frac{\partial V(\tau)}{\partial \tau_1} = \lambda \alpha^{\tau_1} - \alpha^{\tau_2 - \tau_1},$$

and thus:

$$\begin{aligned} \frac{\partial V(\tau)}{\partial \tau_1} &= 0 \\ \lambda \alpha^{\tau_1^*} &= \alpha^{\tau_2 - \tau_1^*} \\ \alpha^{\tau_1^* + \log_\alpha(\lambda)} &= \alpha^{\tau_2 - \tau_1^*} \\ \tau_1^* + \log_\alpha(\lambda) &= \tau_2 - \tau_1^* \\ \tau_1^* &= \frac{1}{2} (\tau_2 - \log_\alpha(\lambda)). \end{aligned}$$

The critical point for $\lambda > 0$ is then found by solving the system:

$$\begin{cases} \tau_1^* = \frac{1}{2} (\tau_2^* - \log_\alpha(\lambda)) \\ \tau_k^* = \frac{1}{2} (\tau_{k+1}^* + \tau_{k-1}^*) & \text{for } 1 < k \leq M \\ \tau_{M+1}^* = T \end{cases} \quad (22)$$

Lemma 13 For $k \in \{1, 2, \dots, M\}$, τ_k^* from Equation 22 can be recursively calculated by:

$$\tau_k^* = \frac{k}{k+1} \tau_{k+1}^* - \frac{1}{k+1} \log_\alpha(\lambda) \quad (23)$$

Proof: By induction. For $k = 1$ the statement holds immediately.

Suppose that for $k = \kappa$, τ_κ^* is can be calculated by using Equation 23. Then for $k = \kappa + 1$, we have:

$$\begin{aligned} \tau_{\kappa+1}^* &= \frac{1}{2} (\tau_\kappa^* + \tau_{\kappa+2}^*) \\ &= \frac{1}{2} \left(\frac{\kappa}{\kappa+1} \tau_{\kappa+1}^* - \frac{1}{\kappa+1} \log_\alpha(\lambda) + \tau_{\kappa+2}^* \right) \\ \left(1 - \frac{\kappa}{2\kappa+2} \right) \tau_{\kappa+1}^* &= \frac{1}{2} \tau_{\kappa+2}^* - \frac{1}{2\kappa+2} \log_\alpha(\lambda) \\ \frac{\kappa+2}{2\kappa+2} \tau_{\kappa+1}^* &= \frac{1}{2} \tau_{\kappa+2}^* - \frac{1}{2\kappa+2} \log_\alpha(\lambda) \\ \tau_{\kappa+1}^* &= \frac{\kappa+1}{\kappa+2} \tau_{\kappa+2}^* - \frac{1}{\kappa+2} \log_\alpha(\lambda) \end{aligned}$$

Thus the statement also holds for $k = \kappa + 1$ and by the Principle of Mathematical Induction the statement is true for all $k \in \{1, 2, \dots, M\}$. ■

Proposition 14 For $k \in \{1, 2, \dots, M\}$, τ_{M-k}^* from Equation 22 can be directly calculated by:

$$\tau_{M-k}^* = \frac{M-k}{M+1} T - \frac{k+1}{M+1} \log_\alpha(\lambda)$$

Proof: This follows from solving the first-order difference equation given by Equation 23 and using $\tau_{M+1}^* = T$ as boundary condition.

Another approach is to verify that the proposed solution satisfies Equation 22. As this is a system of $M+1$ unknowns and $M+1$ linearly independent equations, it is the unique solution. ■

12.3.2 The case $\phi < 1$: deriving fixed-point equations

For this case we shall find a fixed-point equation. It will help us in finding optimal strategies. We begin with a theorem that gives a closed-form definition of $f_k(\cdot)$ for every $k > 0$.

Theorem 15 For $k > 0$, $f_k(\tau_k)$ is given by:

$$f_k(\tau_k) = \lambda(1-\phi)^k \alpha^{\tau_k} + \phi \sum_{i=0}^{k-1} (1-\phi)^i \alpha^{\tau_k - \tau_{k-i}} \quad (24)$$

Proof: For $k = 1$, we have $f_1(\tau_1) = \lambda\alpha^{\tau_1} + \phi(1 - \alpha^{\tau_1}) = (1 - \phi)\lambda\alpha^{\tau_1} + \phi$. This is in agreement with formula (24).

Suppose that the formula holds for $k = \kappa$. Then for $k = \kappa + 1$ we have:

$$\begin{aligned} f_{\kappa+1}(\tau_{\kappa+1}) &= (1 - \phi)f(\tau_{\kappa})\alpha^{\tau_{\kappa+1}-\tau_{\kappa}} + \phi \\ &= \lambda(1 - \phi)^{\kappa+1}\alpha^{\tau_{\kappa+1}} + (1 - \phi)\phi \left(\sum_{i=0}^{\kappa-1} (1 - \phi)^i \alpha^{\tau_{\kappa}-\tau_{\kappa-i}} \right) \alpha^{\tau_{\kappa+1}-\tau_{\kappa}} + \phi \\ &= \lambda(1 - \phi)^{\kappa+1}\alpha^{\tau_{\kappa+1}} + \phi \sum_{j=0}^{\kappa} (1 - \phi)^j \alpha^{\tau_{\kappa+1}-\tau_{\kappa+1-j}}, \end{aligned}$$

where we used the index transformation $j = i + 1$. So clearly formula (24) holds for all k . ■

Corollary 16 $\frac{d}{d\tau_k} f_k(\tau_k)$ is given by:

$$\frac{d}{d\tau_k} f_k(\tau_k) = \ln(\alpha) (f_k(\tau_k) - \phi)$$

Proof: The proof is given by the following derivation.

$$\begin{aligned} \frac{d}{d\tau_k} f_k(\tau_k) &= \lambda(1 - \phi)^k \alpha^{\tau_k} \ln(\alpha) + \phi \sum_{i=1}^{k-1} (1 - \phi)^i \alpha^{\tau_k-\tau_{k-i}} \ln(\alpha) \\ &= \ln(\alpha) (f_k(\tau_k) - \phi). \end{aligned}$$

Notice how the first term of the sum ($i = 0$) drops because $\alpha^{\tau_k-\tau_k} = 1$, and thus a constant with respect to τ_k . ■

Above observations are used to rewrite the formula for the critical points. We use this rewritten version to find a nonlinear system for the critical points of V . The partial derivatives of V are found in the next derivation:

$$\begin{aligned}
\frac{\partial V(\tau)}{\partial \tau_k} &= f_{k-1}(\tau_{k-1}) - f_k(\tau_k) \alpha^{\tau_{k+1}-\tau_k} + (\alpha^{\tau_{k+1}-\tau_k} - 1) (f_k(\tau_k) - \phi) \\
&= f_{k-1}(\tau_{k-1}) \alpha^{\tau_k-\tau_{k-1}} - \phi \alpha^{\tau_{k+1}-\tau_k} - f_k(\tau_k) + \phi \\
&= \left[\lambda(1-\phi)^{k-1} \alpha^{\tau_{k-1}} + \phi \sum_{i=0}^{k-2} (1-\phi)^i \alpha^{\tau_{k-1}-\tau_{k-1-i}} \right] \alpha^{\tau_k-\tau_{k-1}} - \phi \alpha^{\tau_{k+1}-\tau_k} \\
&\quad - \lambda(1-\phi)^k \alpha^{\tau_k} - \phi \sum_{i=1}^{k-1} (1-\phi)^i \alpha^{\tau_k-\tau_{k-i}} \\
&= \lambda(1-\phi)^{k-1} \alpha^{\tau_k} + \phi \alpha^{\tau_k} \sum_{i=1}^{k-2} (1-\phi)^i \alpha^{-\tau_{k-1-i}} + \\
&\quad \phi \alpha^{-\tau_{k-1}} \alpha^{\tau_k} - \phi \alpha^{\tau_{k+1}-\tau_k} - \lambda(1-\phi)^k \alpha^{\tau_k} - \phi \alpha^{\tau_k} \sum_{i=1}^{k-1} (1-\phi)^i \alpha^{-\tau_{k-i}} \\
&= \lambda \phi (1-\phi)^{k-1} \alpha^{\tau_k} + \phi \alpha^{-\tau_{k-1}} \alpha^{\tau_k} - \phi \alpha^{\tau_{k+1}} \alpha^{-\tau_k} + \\
&\quad \phi \alpha^{\tau_k} \left\{ \sum_{i=2}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}} - \sum_{i=1}^{k-1} (1-\phi)^i \alpha^{-\tau_{k-i}} \right\} \\
&= \lambda \phi (1-\phi)^{k-1} \alpha^{\tau_k} + \phi \alpha^{-\tau_{k-1}} \alpha^{\tau_k} - \phi \alpha^{\tau_{k+1}} \alpha^{-\tau_k} + \\
&\quad \phi \alpha^{\tau_k} \left\{ \frac{\phi}{1-\phi} \sum_{i=2}^{k-1} (1-\phi)^i \alpha^{-\tau_{k-i}} - (1-\phi) \alpha^{-\tau_{k-1}} \right\} \\
&= \lambda \phi (1-\phi)^{k-1} \alpha^{\tau_k} + \phi \alpha^{-\tau_{k-1}} \alpha^{\tau_k} - \phi \alpha^{\tau_{k+1}} \alpha^{-\tau_k} + \\
&\quad \phi \left[\frac{\phi}{1-\phi} \left(\sum_{i=2}^{k-1} (1-\phi)^i \alpha^{-\tau_{k-i}} \right) - (1-\phi) \alpha^{-\tau_{k-1}} \right] \alpha^{\tau_k} \\
&= \lambda \phi (1-\phi)^{k-1} \alpha^{\tau_k} + \phi^2 \alpha^{-\tau_{k-1}} \alpha^{\tau_k} - \phi \alpha^{\tau_{k+1}} \alpha^{-\tau_k} + \frac{\phi^2}{1-\phi} \left(\sum_{i=2}^{k-1} (1-\phi)^i \alpha^{-\tau_{k-i}} \right) \alpha^{\tau_k} \\
&= \left(\lambda \phi (1-\phi)^{k-1} + \phi^2 \sum_{i=1}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}} \right) \alpha^{\tau_k} - \phi \alpha^{\tau_{k+1}} \alpha^{-\tau_k}
\end{aligned}$$

Theorem 17 *The critical points of V are given by the following system of equations:*

$$\left\{ \tau_k^* = \frac{1}{2} \left(\tau_{k+1}^* - \log_{\alpha} \left(\lambda(1-\phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}^*} \right) \right) \quad \text{for all } k \in \{1, \dots, M\}. \right. \quad (25)$$

Note: $\tau_{M+1} := T$.

Proof: We simply set $\frac{\partial V(\tau)}{\partial \tau_k} = 0$ and solve for τ_k :

$$\begin{aligned}
 \phi \alpha^{\tau_{k+1}} \alpha^{-\tau_k^*} &= \left(\lambda \phi (1 - \phi)^{k-1} + \phi^2 \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}} \right) \alpha^{\tau_k^*} \\
 \phi \alpha^{\tau_{k+1}} &= \left(\lambda \phi (1 - \phi)^{k-1} + \phi^2 \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}} \right) (\alpha^{\tau_k^*})^2 \\
 (\alpha^{\tau_k^*})^2 &= \frac{\phi \alpha^{\tau_{k+1}}}{\lambda \phi (1 - \phi)^{k-1} + \phi^2 \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}}} \\
 \alpha^{\tau_k^*} &= \sqrt{\frac{\phi \alpha^{\tau_{k+1}}}{\lambda \phi (1 - \phi)^{k-1} + \phi^2 \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}}}} \\
 \tau_k^* &= \frac{1}{2} \log_{\alpha} \left(\frac{\phi \alpha^{\tau_{k+1}}}{\lambda \phi (1 - \phi)^{k-1} + \phi^2 \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}}} \right) \\
 &= \frac{1}{2} \left(\log_{\alpha} (\alpha^{\tau_{k+1}}) - \log_{\alpha} \left(\lambda (1 - \phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}} \right) \right) \\
 &= \frac{1}{2} \left(\tau_{k+1} - \log_{\alpha} \left(\lambda (1 - \phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}} \right) \right)
 \end{aligned}$$

So if all other τ_l ($l \neq k$) are fixed, the above gives τ_k^* as the critical point for V as a function of τ_k . In the critical point of $V(\tau)$ this condition hold for every k and the system given by Equation (25) follows. $\blacksquare$

In the following we shall prove that $\tau_1^* \leq \tau_2^* \leq \dots \leq \tau_M^*$ in many cases. To this end we need the next lemma.

Lemma 18 *If $\phi \geq \lambda$, then for all $1 \leq k \leq M$ it holds that:*

$$\tau_k^* \leq \frac{1}{2} \left(\tau_{k+1}^* + \max_{i=1,2,\dots,k-1} \tau_i^* \right) \tag{26}$$

Proof: Because $\alpha < 1$, we have that $\log_{\alpha}(\cdot)$ is monotonically decreasing in its argument. τ_k^* is thus upper-bounded by $\frac{1}{2} (\tau_{k+1} - \log_{\alpha}(x))$ when $x \geq \lambda (1 - \phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1 - \phi)^{i-1} \alpha^{-\tau_{k-i}^*}$. The

following derivation results in such an x .

$$\begin{aligned}
 \lambda(1-\phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}^*} &\leq \phi(1-\phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}^*} \\
 &= \phi \sum_{i=1}^k (1-\phi)^{i-1} \alpha^{-\tau_{k-i}^*} \\
 &\leq \phi \max_{i=1,2,\dots,k-1} \alpha^{-\tau_{k-i}^*} \sum_{i=1}^k (1-\phi)^{i-1} \\
 &= \phi \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} \sum_{i=0}^{k-1} (1-\phi)^i \\
 &= \phi \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} \frac{1-(1-\phi)^k}{\phi} \\
 &= \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} (1-(1-\phi)^k) \\
 &\leq \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*}.
 \end{aligned}$$

Now we use that $\alpha^{-p} \geq \alpha^{-q}$ implies $p \geq q$. This means that:

$$\begin{aligned}
 \tau_k^* &\leq \frac{1}{2} \left(\tau_{k+1}^* + \log_{\alpha} \left(\max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} \right) \right) \\
 &= \frac{1}{2} \left(\tau_{k+1}^* + \max_{i=1,2,\dots,k-1} \tau_i^* \right)
 \end{aligned}$$

■

Proposition 19 *If $\phi \geq \lambda$, then for the critical points of $V(\tau)$ it holds that $\tau_1^* \leq \tau_2^* \leq \dots \leq \tau_k^* \leq \dots \leq T$.*

Proof: We use Equation (26) and a total induction argument on k .

As a basic step, we have $\tau_1^* \leq \frac{1}{2}\tau_2^*$, so $\tau_1^* \leq \tau_2^*$. Next, suppose that $\tau_1^* \leq \tau_2^* \leq \dots \leq \tau_{\kappa-1}^* \leq \tau_{\kappa}^*$. Then we also have:

$$\begin{aligned}
 \tau_{\kappa}^* &\leq \frac{1}{2} \left(\tau_{\kappa+1}^* + \max_{i=1,2,\dots,\kappa-1} \tau_i^* \right) \\
 &= \frac{1}{2} (\tau_{\kappa+1}^* + \tau_{\kappa-1}^*) \\
 &\leq \frac{1}{2} (\tau_{\kappa+1}^* + \tau_{\kappa}^*) \\
 \frac{1}{2}\tau_{\kappa}^* &\leq \frac{1}{2}\tau_{\kappa+1}^*
 \end{aligned}$$

So $\tau_{\kappa} \leq \tau_{\kappa+1}$ and by the Principle of Mathematical Induction we thus have $\tau_1^* \leq \tau_2^* \leq \dots \leq \tau_k^* \leq \dots \leq T$. ■

Lemma 20 *If $\tau_1^* > 0$, then Lemma 18 and Corollary 19 hold for all ϕ and λ .*

Proof: $\alpha^{-x} > 1$ holds for $x > 0$, so if $\tau_1^* > 0$ then $\max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} > 1$ for all k . Then we get:

$$\begin{aligned}
& \lambda(1-\phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}^*} \\
& \leq \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} \left\{ \lambda(1-\phi)^{k-1} + \phi \sum_{i=0}^{k-2} (1-\phi)^i \right\} \\
& = \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} \left\{ \lambda(1-\phi)^{k-1} + (1-(1-\phi)^{k-1}) \right\} \\
& = \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*} \left\{ 1 - (1-\lambda)(1-\phi)^{k-1} \right\} \\
& \leq \max_{i=1,2,\dots,k-1} \alpha^{-\tau_i^*}
\end{aligned}$$

The rest of the proofs follows as in Lemma 18. ■

A natural extension (from the $M = 1$ and $M = 2$ cases) of an optimal strategic mentioning system would be:

$$\left\{ \tau_k^{OPT} = \max \left\{ 0, \frac{1}{2} \left(\tau_{k+1}^{OPT} - \log_{\alpha} \left(\lambda(1-\phi)^{k-1} + \phi \sum_{i=1}^{k-1} (1-\phi)^{i-1} \alpha^{-\tau_{k-i}^{OPT}} \right) \right) \right\} \quad \text{for all } k \in \{1, \dots, M\}. \right. \quad (27)$$

Note that substituting $\phi = 1$, interpreting 0^0 as 1 and defining $\sum_{i=1}^0 \xi(i) := 0$ yields Equation 22. This means that we can use the fixed-point approach for all $\phi \leq 1$, even though its original derivation assumed $\phi \neq 1$.

Solving this fixed-point problem would intuitively give a strategy that maximizes the average motivation level. We shall show in the next subsections that this is often, but not always, the case. Using Equation 27 for finding suitable mentioning strategies will be referred to as the fixed-point approach in the rest of this research.

12.3.3 Numerical approach

As an alternative to the fixed-point approach, we can use available optimization tools to find strategies that optimize V given the number of mentions M . In Matlab, we can use FMINCON to achieve this. The function for which we want to find a *minimizer* is $-V(\tau|T, \alpha, \lambda, \phi, M)$. Our starting point is a random feasible strategy:

$$\tau^{(0)} = T \times \text{sort}(\text{rand}(M, 1)).$$

Each variable is restricted to lie in $[0, T]$. Furthermore we demand that $\tau_k - \tau_{k+1} \leq 0$ for all $k \in \{1, 2, \dots, M\}$. They are added as linear inequality constraints in FMINCON.

12.3.4 Optimality of fixed-point equation

We performed simulations with Equation (27) and it seems to give good and intuitively sensible results. It could be possible to prove optimality by using concavity properties, similar to what we did in the $M = 1$ case. However, for now we choose another approach, because finding concavity

properties appears to be involved. We set $T = 30$, in agreement with the number of days in November. Then we choose parameter values for M , α , λ and ϕ and we sample $N_R = 1000$ random strategies. We count how often a random strategy gives a value at least as high as the value obtained from solving Equation (27), for this choice of parameters. We checked all combinations of the following parameter values (thus creating four nested for-loops):

$$\begin{aligned} M &\in \{1, 2, 3, 4, 5\}, \\ \alpha &\in \{0.1, 0.3, 0.5, 0.7, 0.9\}, \\ \lambda &\in \{0, 0.2, 0.4, 0.6, 0.8, 1\}, \\ \phi &\in \{0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9, 1\}. \end{aligned}$$

Let $R(M, \alpha, \lambda, \phi)$ be the set of N_R random strategies drawn for the parameter values M , α , λ and ϕ . We increased a counter if a random strategy $\tau \in R(M, \alpha, \lambda, \phi)$ satisfied the condition that $V(\tau|M, \alpha, \lambda, \phi) > V(\tau^{OPT}|M, \alpha, \lambda, \phi) + \epsilon$, where τ^{OPT} is calculated by Equation (27). We chose $\epsilon = 10^{-4}$. ϵ was introduced to deal with computational errors; strategies with a value equal to $V(\tau^{OPT})$ were sometimes falsely regarded as giving a (slightly) higher value due to these errors. Formally, the percentage of random strategies which yields higher results than τ^{OPT} is given by:

$$F(M, \alpha, \lambda, \phi) := \frac{|\{\tau \in R(M, \alpha, \lambda, \phi) \mid V(\tau|M, \alpha, \lambda, \phi) > V(\tau^{OPT}|M, \alpha, \lambda, \phi) + \epsilon\}|}{N_R} \times 100.$$

For $M = 1$, optimality of the fixed-point approach is again shown by these simulations, so that 0% of the random strategies had a value exceeding that of τ^{OPT} .

For $M = 2$, solving Equation (27) does not always yield optimizing strategies. However, there are many vectors $(M, \alpha, \lambda, \phi)$ for which $F(M, \alpha, \lambda, \phi) = 0$, often for low values of α ($= 0.1, 0.3$). With $\alpha = 0.7$, we always find $F(2, 0.7, \lambda, \phi) < 0.1$. For $\alpha = 0.9$ we show in Table 5 the effect of different λ and ϕ on F . We note that with a higher ϕ , the fixed-point approach is more likely to yield optimal results. The effect of λ is ambiguous, although the table shows that the fixed-point approach was never outperformed in the case that $\lambda = 0$.

Note that even though the derivations for Equation (27) were restricted to $\phi < 1$, we included $\phi = 1$ in simulations using the fixed-point solution. The table shows that this approach also gives desirable results for $\phi = 1$, as the resulting strategies were not outperformed by any random strategy. It agrees with the side-note we made right after posing Equation (27).

For $M = 3$ the fixed-point approach performs quite well as long as $\phi > 0.5$. In this case not more than 3% of the random strategies outperform the fixed-point strategy. The fixed-point approach works best when λ is small.

In general we hypothesize that the fixed-point approach works best when ϕ is large and λ is small. Otherwise

12.4 Effect of parameters on optimal average motivation values and optimal strategies

The quantity of interest in the simulations is the average motivation level:

$$\hat{V}(\tau) = \frac{1}{T} V(\tau).$$

$\phi \backslash \lambda$	0	0.2	0.4	0.6	0.8	1
0.1	0	11.9	28.4	23.1	21.6	17.1
0.2	0	8.2	23.0	16.9	14.5	13.4
0.3	0	3.4	16.1	11.4	10.9	11.1
0.4	0	1.9	9.6	8.7	6.8	6.3
0.5	0	0.9	3.5	4.4	4.5	4.6
0.6	0	0.8	2.8	3.2	2.6	3.2
0.7	0	0.3	1.4	1.2	1.6	1.0
0.8	0	0	0.4	0.9	0.3	0.3
0.9	0	0	0	0.1	0	0
1.0	0	0	0	0	0	0

Table 5: Percentage of random solutions outperforming solution of Equation (27), for $M = 2$, $\alpha = 0.9$.

Trivially, $\tau = \tau^{OPT}$ is an maximizer for $\hat{V}(\tau)$ if and only if τ^{OPT} is an maximizer of $V(\tau)$. $\hat{V}(\tau)$ always lies between 0 and 1 and makes it suitable for comparison of strategies and the effect of parameters. We shall use Equation (27) for finding strategies with the optimal average motivation level. We compare its performance with the approach with FMINCON as described in Section 12.3.3.

In Figure 24 we see how different fixed parameters, such as degree of influence ϕ and initial motivation level λ , affect the optimal average motivation level that can be achieved by using Equation (27). We see that the influence of λ becomes smaller if we can make more mentions. A higher degree of influence also leads to higher motivation levels, as expected.

In Figure 25 we see the results when we apply Matlab's function FMINCON for constrained optimization. Qualitatively, the results are similar: higher degrees of influence lead to higher optimal average motivation levels and the effect of the initial motivation level becomes smaller as M grows larger. The contour plots show that the fixed-point approach can find strategies that are near-optimal, although it is outperformed by FMINCON if M grows larger.

Given λ , α , T , ϕ and M , denote the found optimizers respectively by $\tau^{fixed-point}(\lambda, \alpha, T, \phi, M)$ and $\tau^{fmincon}(\lambda, \alpha, T, \phi, M)$. For $M = 1$, $T = 30$ and $\alpha = 0.9$ the algorithms were tested for all $[\lambda, \phi] \in \{0.1, 0.2, \dots, 1\} \times \{0.1, 0.2, \dots, 1\}$, yielding 100 strategies.

Let $\|\tau\|_1 := \sum_{k \in [M]} |\tau_k|$. We found that:

$$\max_{[\lambda, \phi]} \|\tau^{fixed-point}(\lambda, \alpha, T, \phi, M) - \tau^{fmincon}(\lambda, \alpha, T, \phi, M)\|_1 < 4.0 \times 10^{-4},$$

showing that there is virtually no difference between the strategies found by both algorithms. However, on average the FMINCON takes 5.45 times longer to find this solution than the fixed-point approach does. On average, the fixed-point approach takes 0.020 seconds for an instance of $M = 1$

For $M = 2$ and repeating the same procedure, we found that:

$$\max_{[\lambda, \phi]} \|\tau^{fixed-point}(\lambda, \alpha, T, \phi, M) - \tau^{fmincon}(\lambda, \alpha, T, \phi, M)\|_1 \approx 8.4,$$

suggesting that the found strategies can differ between the two approaches. As stated before, this suggests that the fixed-point approach does not yield optimal results for $M > 2$. However, for $M = 2$ using FMINCON takes on average 11.6 times longer than using the fixed-point approach. On average, the fixed-point approach takes 0.024 seconds for an instance of $M = 2$.

Generalizing the above results, we conclude that the approach using FMINCON yields better results at the cost of longer computation times. However, all computation times we encountered stayed well below a second, so we advise to use FMINCON for instances where $M \geq 2$.

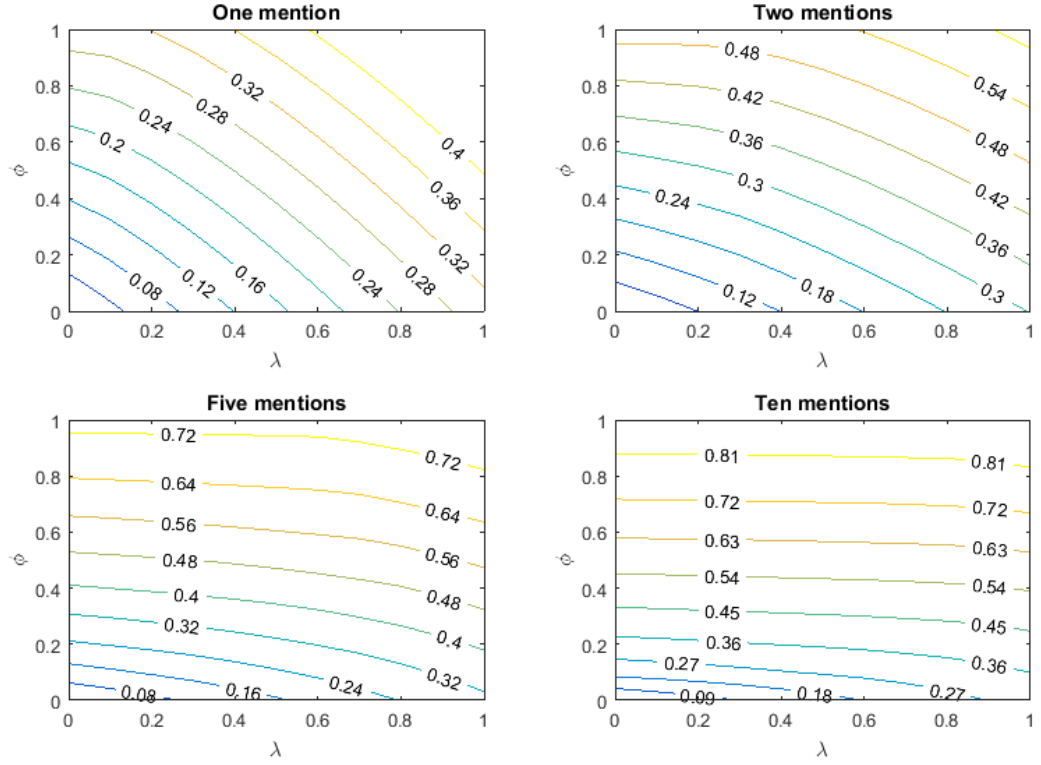


Figure 24: Optimal strategy values for different numbers of mentions. Level curves give the optimal average motivation level on the (λ, ϕ) plane when applying the fixed-point equation. $T = 30$, $\alpha = 0.9$.

We used the approach with FMINCON in the rest of this section, because it is better in finding high-valued strategies.

From Figures 27 through 30, shown in Appendix C, we can deduce that mentioning should be done early in time when the degree of influence is small and the initial motivation level is low. The higher the degree of influence, the better it is to distribute the mentions uniformly over the horizon. The higher the initial motivation level, the better it is to wait with the first mention. For $M = 2$, the results at $\phi = 1$ could also be found by using Equation 21 (for $\lambda = 0$) and Equation 23 (for $\lambda = 0.5$).

We also found feasible strategies by using Equation (27). These are sub-optimal for $M > 1$. The differences between the strategies become evident when looking at Figure 31, compared to Figure 30. On the other hand, it should be noted that even though the resulting strategies differ greatly, their average values do not, see Figures 24 and 25.

12.5 Sensitivity of value to suboptimal strategies

In this section we consider again continuous-time strategies and use the approach with FMINCON, because we found this approach to yield strategies with maximal average motivation level.

First, let $\|\cdot\|_1$ denote the 1-norm, so that $\|\tau\|_1 := \sum_{k \in [M]} |\tau_k|$. If $\|\tau^{OPT} - \tau\|_1 = \epsilon$, it means that strategy τ has its mentions positioned in such a way that in total there has been a shift of

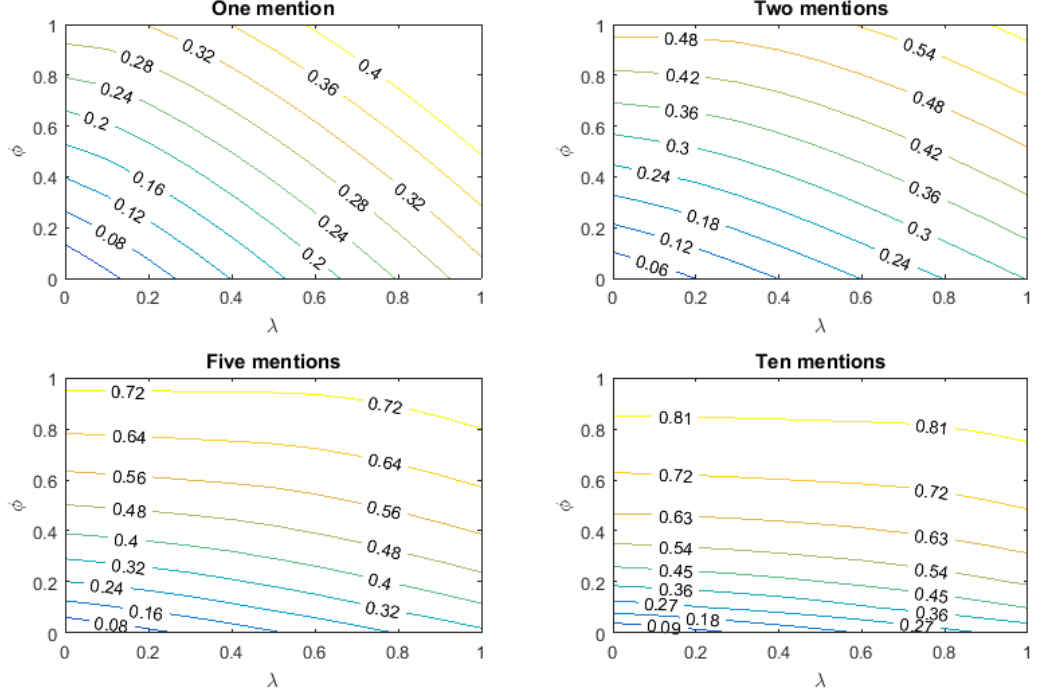


Figure 25: Optimal strategy values for different numbers of mentions. Level curves give the optimal average motivation level on the (λ, ϕ) plane when applying the strategy found by FMINCON. $T = 30$, $\alpha = 0.9$.

ϵ days compared to τ^{OPT} . For example, if τ is derived from the optimal strategy by making one mention a day earlier, and another mention two days later, then $\epsilon = 3$.

Let $\Omega := \{\tau \mid 0 \leq \tau_1 \leq \tau_2 \leq \dots \leq \tau_M \leq T\}$ be the set of feasible strategies. The *random strategy value* is defined to be $V^R := \mathbb{E}[\hat{V}(\tau^R)]$, where τ^R is uniformly drawn from Ω . This expected value can easily be approximated by Monte Carlo simulations.

The idea is that strategies that deviate slightly from the optimal strategy (with optimal average value V^{OPT}) still have good values. However, if deviations are allowed to be big, then their value will not always exceed the random strategy value any more.

Let $B_\epsilon(\tau^{OPT}) := \{\tau \mid \|\tau^{OPT} - \tau\|_1 < \epsilon\}$. Next, define $V_\epsilon(\tau^{OPT}) := \mathbb{E}[\hat{V}(\tau) \mid \tau \in B_\epsilon(\tau^{OPT}) \cap \Omega]$ where τ is uniformly distributed over $B_\epsilon(\tau^{OPT}) \cap \Omega$. This expectation can also be approximated by Monte Carlo simulation. Note that not all τ for which the distance to τ^{OPT} is less or equal to ϵ are allowed; they must also be non-decreasing vectors in order to be feasible strategies. In the simulation we draw τ uniformly from $B_\epsilon(\tau^{OPT})$ and discard the draw if $\tau \notin \Omega$.

For all experiments we set the number of feasible (non-discarded) to-be-drawn Monte Carlo vectors to be $N_{MC} = 1000$, for the random strategy value as well as $V_\epsilon(\tau^{OPT})$ for every ϵ . We set $\phi = 0.8$ in the analysis, because $\phi_{November_{NL}} \approx 0.8$. Furthermore we let $T = 30$ (the days of November) and $\alpha = 0.9$.

For the numerical results we let λ and M vary, see Figures 32 through 35 in Appendix C for examples. In the graphs we see that the expected value does not become worse very quickly. Even for $\epsilon = 10$ quite reasonable values can be expected. If the allowed number of mentions increases, there is more room for error (ϵ may be chosen bigger) and the disturbed strategy will, in expectation, still yield a value higher than the random strategy value.

13 Conclusions

In this report we examined the relation between mentions on Twitter and donation raised for the Movember Foundation. It was found that looking at mentions is by itself not enough to establish a positive relation between them. In order to model the relationship in a more advanced way, we developed the Satiation-Deprivation (SD) Model, a dynamic motivation model. We arrived at a recursive definition, see Equation (13), suitable for analytical purposes. On the other hand we made an implementation for the SD Model that allows for generalization, see Algorithms 1 and 2. We were able to describe a positive relation between motivation level and donation potential for The Netherlands and Sweden, visualized in Figures 17 and 18 respectively. Donation occurrence was shown to be related to the timing of mentions. On the other hand, the influence of the order of mentioners on donation occurrence was not shown to be significant.

Under the assumption that mentions are scarce and motivation levels of individuals follow the SD Model, we developed a brute-force algorithm to maximize the average motivation level over a discrete time axis, see Algorithm 4. Maximizing the average motivation level is related to maximizing average donation potential, according to the above results. The algorithm can distribute mentions over a given time horizon, under the assumption that no mentions of others occur. We found that using a horizon of about two weeks yielded the highest motivation levels on the data set; with larger horizons no improvements were found. However, the average donation potential did not appear to increase for horizons larger than two days. A 19% increase in donation potentials of MovemberNL's mentionees could be achieved if MOVEMBERNL adopts the mentioning algorithm. Next, we switched to a continuous time axis and analyzed efficient ways to find optimal strategies. We showed that the critical points of the value function (total motivation level as a function of mentioning strategies) give feasible strategies provided that the first mention is made in positive time. We proposed a heuristic for finding optimal strategies based on analytically obtained critical points of the the total motivation level function in continuous time, see Equation (27). The strategies obtained by solving these fixed-point equations were not proven to give the optimal strategy for $M > 1$, but yielded strategies with high motivation levels in simulations. Moreover, with respect to computation times the fixed-point approach outperformed a more straight-forward approach.

We studied the influence of different model parameters on the average motivation level. Simulations showed that

- the mentioner's degree of influence affects the average motivation level strongly;
- the effect of the initial motivation level is smaller when more mentions are allowed;
- users with small influence should make their mentions close to each other in time;
- users with large influence should distribute mentions uniformly over the available time.

Finally, slight deviations from the optimal strategy were shown to have only a small effect on the resulting average motivation level.

14 Recommendations for Movember

Our research has shown that Twitter activity is important for donations. Therefore we suggest that more research be undertaken on how to increase activity and engagement on Twitter, which could be performed by researchers in the fields of psychology and sociology.

Making mentions strategically throughout time could lead to more donations. As MOVEMBERNL has many followers on Twitter, the model suggests that distributing mentions evenly over the Movember period (for each mentionee) yields a higher donation potential than for example making them close to each other in time. For The Netherlands, Figure 17 suggests that special attention should be paid to those with middle-range motivation levels; for these users an increase in motivation level comes with a relatively high increase in donation potential. However, we should also keep in mind that mentions should be made appropriate to the context, not exclusively because they are mathematically strategic; this falls beyond the scope of this research.

Technologically, it is possible to stream Twitter data to Movember. The SD Model could then be used to compute motivation levels during future campaigns. In addition, following the methods in this research, other countries could be evaluated in similar ways. When the donation potential graphs for the countries are compared, we might find that in some countries the effect of mentions is much stronger than in others. A steeper increase could indicate that Twitter mentions are more effective in one country than in another.

15 Future research

In this research, we had information about the number of followers of mentioners and used this for defining the degree of influence. The degree of influence could be redefined by other centralities on other networks. For example, if the follower network structure from Twitter is available, we could employ established centrality measures on this network to redefine the degree of influence. In addition, it would be of interest to look at how donation height corresponds with motivation level. Moreover, we would like to define activity potential and mentioning potential in a similar way to donation potential and see if positive relations with the motivation levels exists in these cases as well.

In our search for an optimal mentioning strategy we employed a scarcity assumption. Figure 22 showed us that this assumption can limit the applicability of Algorithm 4 when large horizons are considered. Therefore we would like to find a method that can make predictions about future mentioning behavior and include these predictions in the optimization. To this end, we could think of ways to derive mentioning probabilities from the data, and employ the results we found for stochastic mentioning (see Section 5.3).

16 Acknowledgments

I would like to give a warm thanks to my first supervisor, Nelly Litvak. Not only has she provided me with invaluable guidance throughout the development of the models in this work; she has also provided me with the courage to persevere in my investigations during this project and during my internship. Moreover, I am indebted for her clear and insightful feedback.

My thanks go to my supervisor Tijs van den Broek as well. He has provided me with Movember and Twitter data and contributed with ideas and directions in which to develop the models.

I would like to thank the assessment committee, consisting of my supervisors, Richard Boucherie and Anton Stoorvogel, for devoting their time to reading this work. Additionally, I would also like to thank Ariana Need, mainly for her suggestion to take a look at strategic mentioning. It became a large piece of this work.

I have enjoyed the time with my colleagues from the various offices I have been working in, as well as those encountered at the coffee machine and lunch breaks.

A final word of thanks goes to my family and my friends. Their contributions to this work are less direct, yet indispensable.

References

- [1] Paolo Boldi and Sebastiano Vigna. Axioms for centrality. *Internet Mathematics*, 10(3-4):222–262, 2014.
- [2] Andreas Brandt. The stochastic equation $y_{n+1} = a_n y_n + b_n$ with stationary coefficients. *Advances in Applied Probability*, 18(1):211–220, 1986.
- [3] Richard M Emerson. Social exchange theory. *Annual review of sociology*, pages 335–362, 1976.
- [4] Adam M Grant. Does intrinsic motivation fuel the prosocial fire? motivational synergy in predicting persistence, performance, and productivity. *Journal of Applied Psychology*, 93(1):48–58, 2008.
- [5] Frank J Massey Jr. The kolmogorov-smirnov test for goodness-of-fit. *Journal of the American Statistical Association*, 46(253):68–78, 1951.
- [6] Tijs van den Broek, Ariana Need, Michel Ehrenhard, Anna Priante, and Djoerd Hiemstra. The influence of prosocial norms and online network structure on prosocial behavior: an analysis of movember’s twitter campaign in 24 countries. 2015.
- [7] Han van der Veen, Djoerd Hiemstra, Tijs van den Broek, Michel Ehrenhard, and Ariana Need. Determine the user country of a tweet. *arXiv preprint arXiv:1508.02483*, 2015.
- [8] Sebastiano Vigna. A weighted correlation index for rankings with ties. In *Proceedings of the 24th International Conference on World Wide Web*, pages 1166–1176. International World Wide Web Conferences Steering Committee, 2015.

A Kendall's weighted τ

Large parts of this appendix have been drawn from [8].

Let $U = \{1, \dots, n\}$ be the user set and r and s are vectors containing two different centrality-induced rankings, where the highest ranking is 0 and it is given to the user with the highest centrality. r and s correspond to different centralities, for example in-degree and PageRank.

Let ρ be another ranking function; for simplicity we say it is the *true* ranking of the users (given by an external source). The ranked-weight product is now defined as:

$$\langle r, s \rangle_{\rho, w} := \sum_{u < v} \text{sgn}(r_u - r_v) \text{sgn}(s_u - s_v) w(\rho(u), \rho(v)) \quad (28)$$

The corresponding norm is given by $\|r\|_w = \sqrt{\langle r, r \rangle}$. The weight function w is related to the idea that difference in high-centrality persons are more important than differences in low-centrality persons. In this work we opt for a hyperbolic additive weighting scheme. This means that we use the formulas:

$$\begin{aligned} w(\rho(u), \rho(v)) &= f(\rho(u)) + f(\rho(v)), \\ f(i) &= \frac{1}{i + 1}. \end{aligned}$$

Basically, the weight takes its highest values if the two users are both highly ranked by ρ . The sum in (28) becomes large if the centralities of highly ranked individuals (according to the true ρ) are in disagreement.

Kendall's weighted τ , with true ranking ρ , is defined to be:

$$\tau_{\rho, w}(r, s) := \frac{\langle r, s \rangle}{\|r\|_w \|s\|_w}.$$

However, we are lacking a ground truth, a 'true' ranking ρ . Instead we can only use r and s themselves for an approximation, which provide us with the only available rankings. These centralities should be treated symmetrically, so we use both $\rho_{r, s}$, the ranking function that orders the users lexicographically on r and then on s if a tie occurs, and $\rho_{s, r}$, which is defined vice versa. Then Kendall's weighted τ is defined to be:

$$\tau_w(r, s) := \frac{\tau_{\rho_{r, s}, w}(r, s) + \tau_{\rho_{s, r}, w}(r, s)}{2}.$$

This is the definition as we have used it in this research.

B Critical values for the Two-sample Kolmogorov-Smirnov test

Critical Values for the Two-sample Kolmogorov-Smirnov test (2-sided)

Table gives critical D -values for $\alpha = 0.05$ (upper value) and $\alpha = 0.01$ (lower value) for various sample sizes. * means you cannot reject H_0 regardless of observed D .

$n_2 \backslash n_1$	3	4	5	6	7	8	9	10	11	12
1	*	*	*	*	*	*	*	*	*	*
2	*	*	*	*	*	16/16	18/18	20/20	22/22	24/24
3	*	*	15/15	18/18	21/21	21/24	24/27	27/30	30/33	30/36
4		16/16	20/20	20/24	24/28	28/32	28/36	30/40	33/44	36/48
5			*	24/30	30/35	30/40	35/45	40/50	39/55	43/60
6				30/36	30/42	34/48	39/54	40/60	43/66	48/72
7				36/36	36/42	40/48	45/54	48/60	54/66	60/72
8					42/49	40/56	42/63	46/70	48/77	53/84
9					42/49	48/56	49/63	53/70	59/77	60/84
10						48/64	46/72	48/80	53/88	60/96
11						56/64	55/72	60/80	64/88	68/96
12							54/81	53/90	59/99	63/108
							63/81	70/90	70/99	75/108
								70/100	60/110	66/120
								80/100	77/110	80/120
									77/121	72/132
									88/121	86/132
										96/144
										84/144

For larger sample sizes, the approximate critical value D_α is given by the equation

$$D_\alpha = c(\alpha) \sqrt{\frac{n_1 + n_2}{n_1 n_2}}$$

where the coefficient is given by the table below.

α	0.10	0.05	0.025	0.01	0.005	0.001
$c(\alpha)$	1.22	1.36	1.48	1.63	1.73	1.95

Examples: (1) At $\alpha = 0.05$ and samples sizes 5 and 8, $D_\alpha = 30/40 = 0.75$.
 (2) At $\alpha = 0.01$ and samples sizes 15 and 28, $D_\alpha = 1.63 \sqrt{\frac{15+28}{15 \cdot 28}} = 0.522$.

Figure 26

Source 'Critical Values for the Two-sample Kolmogorov-Smirnov test (2-sided)':

(https://www.webdepot.umontreal.ca/Usagers/angers/MonDepotPublic/STT3500H10/Critical_KS.pdf)

C Graphs

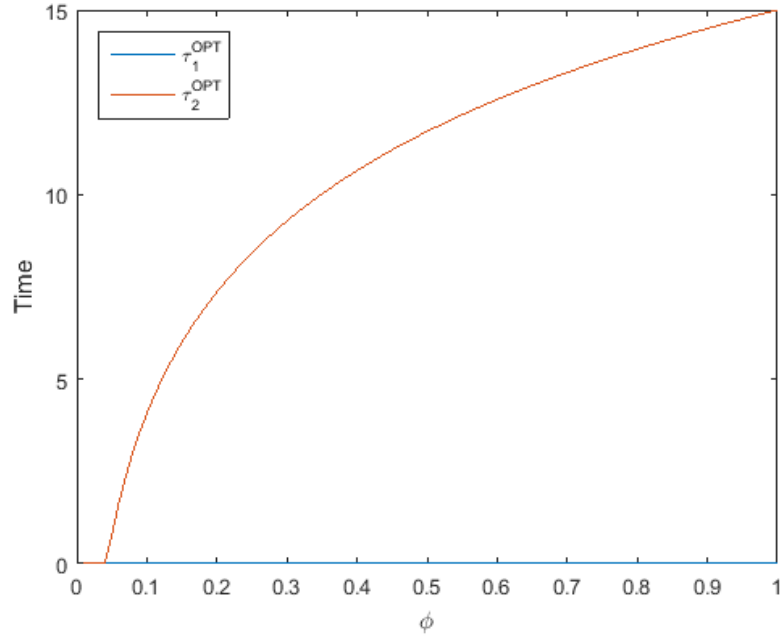


Figure 27: Optimal mentioning times, depending on ϕ . Parameters: $T = 30$, $\alpha = 0.9$, $\lambda = 0$, $M = 2$

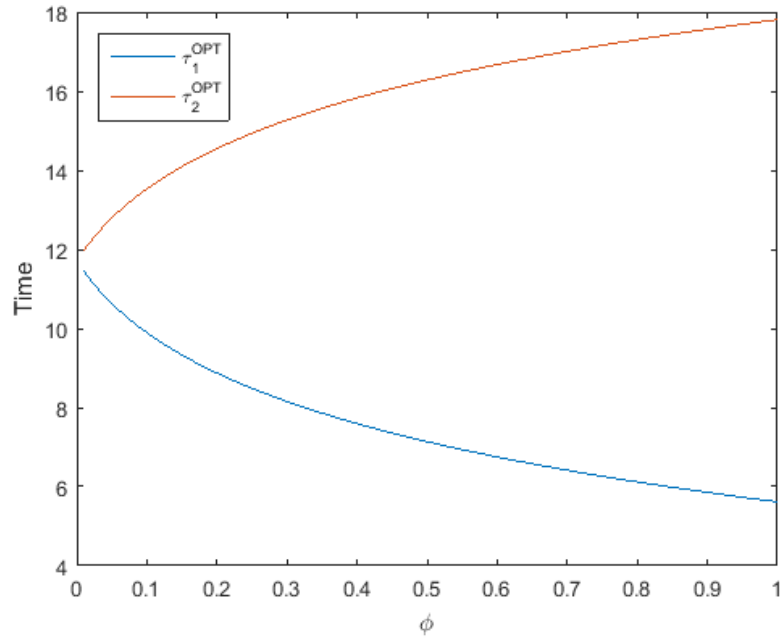


Figure 28: Optimal mentioning times, depending on ϕ . Parameters: $T = 30$, $\alpha = 0.9$, $\lambda = 0.5$, $M = 2$

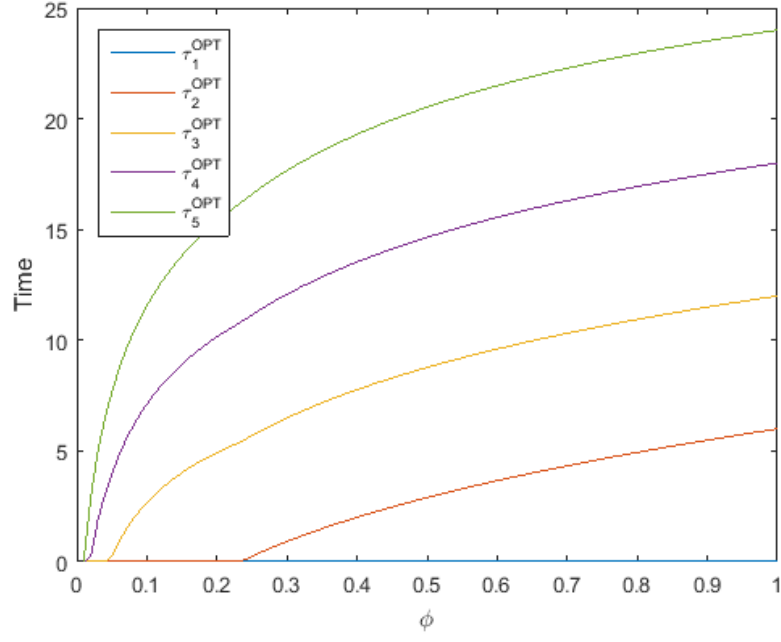


Figure 29: Optimal mentioning times, depending on ϕ . Parameters: $T = 30$, $\alpha = 0.9$, $\lambda = 0$, $M = 5$

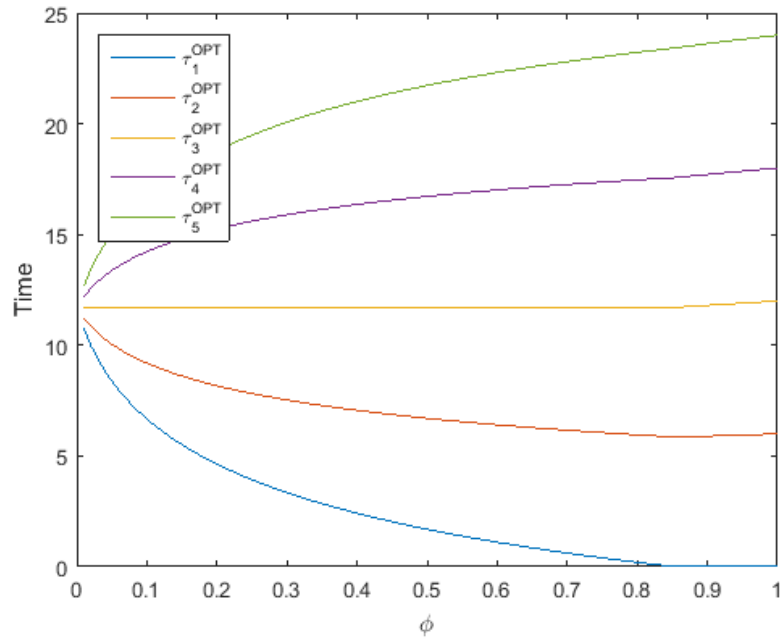


Figure 30: Optimal mentioning times, depending on ϕ . Parameters: $T = 30$, $\alpha = 0.9$, $\lambda = 0.5$, $M = 5$

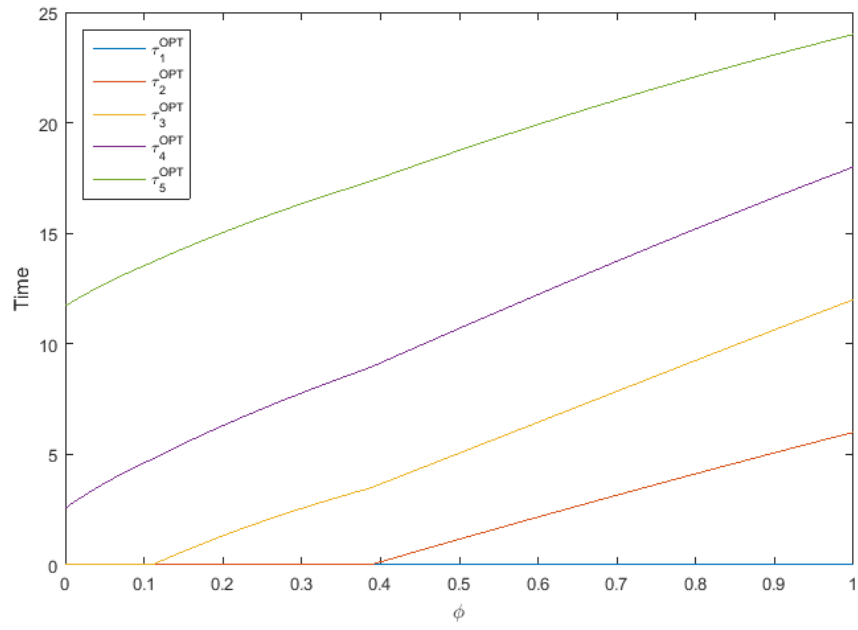


Figure 31: Heuristic mentioning times, depending on ϕ . Parameters: $T = 30$, $\alpha = 0.9$, $\lambda = 0.5$, $M = 5$. We use the fixed-point strategy, which yields suboptimal results

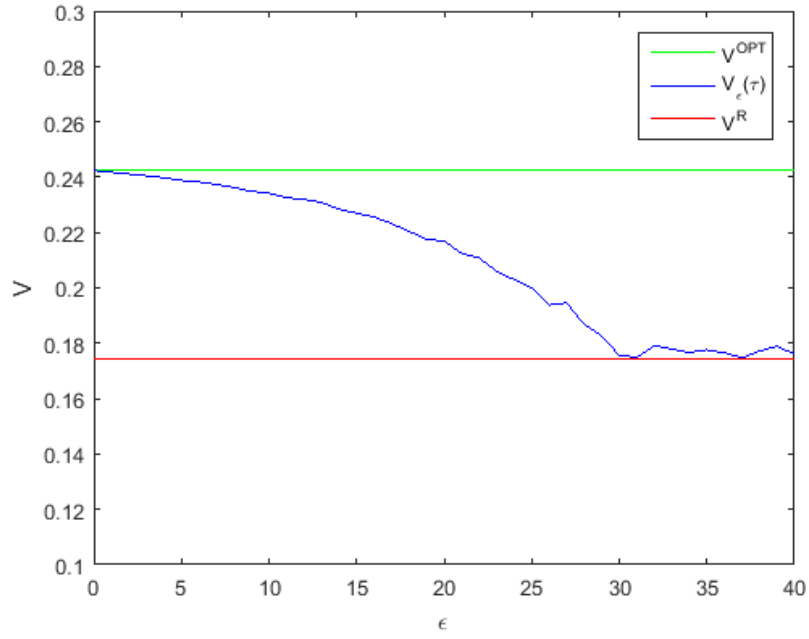


Figure 32: Sensitivity of the average motivation value to deviations from the optimal strategy. $T = 30$, $\alpha = 0.9$, $\phi = 0.8$, $M = 1$ and $\lambda = 0$

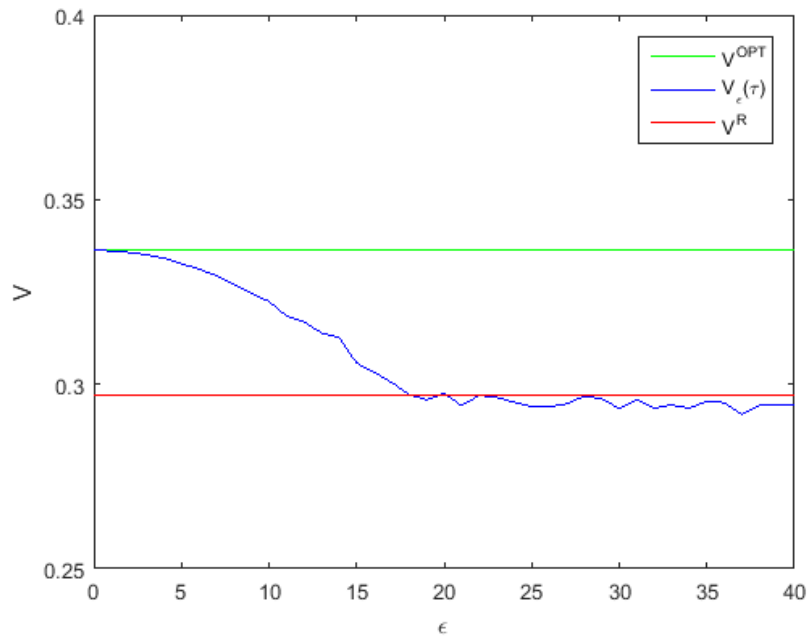


Figure 33: Sensitivity of the average motivation value to deviations from the optimal strategy. $T = 30$, $\alpha = 0.9$, $\phi = 0.8$, $M = 1$ and $\lambda = 0.5$

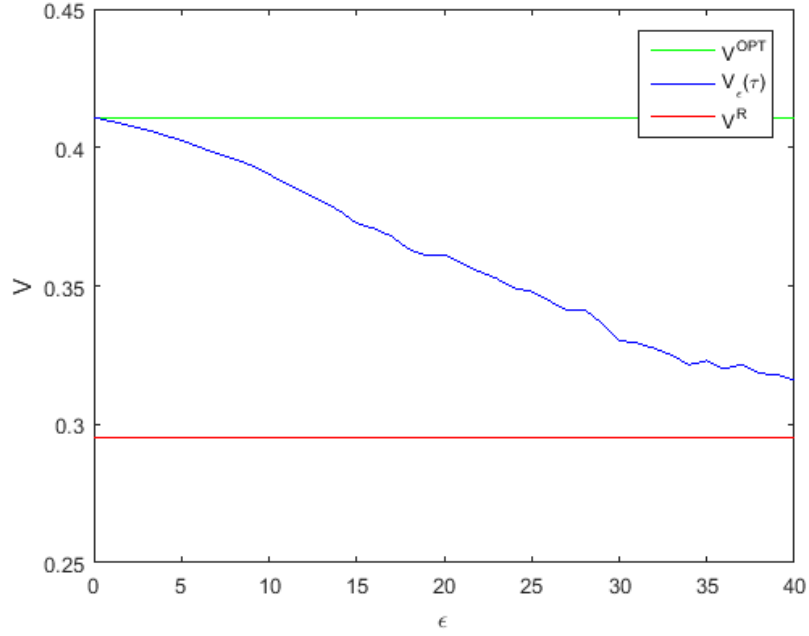


Figure 34: Sensitivity of the average motivation value to deviations from the optimal strategy. $T = 30$, $\alpha = 0.9$, $\phi = 0.8$, $M = 2$ and $\lambda = 0$

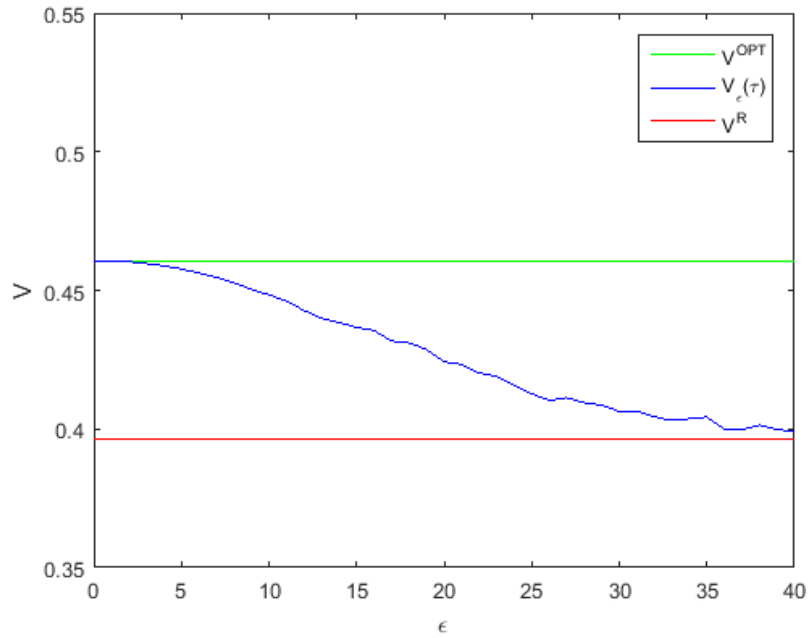


Figure 35: Sensitivity of the average motivation value to deviations from the optimal strategy. $T = 30$, $\alpha = 0.9$, $\phi = 0.8$, $M = 2$ and $\lambda = 0.5$