

Dynamic Network Influence: The Art of Strategic Messaging

Preliminary. Comments welcome and greatly appreciated.*

Wei Li
University of British Columbia

Xu Tan
University of Washington

December, 2023

Abstract

Strategic influencers send costly messages over time to persuade agents in a network. Each influencer maximizes her total discounted payoff, which decreases in the agents' opinion deviations from her agenda. Agents update opinions by taking weighted averages of neighbors' opinions and messages from the influencers. In a single influencer benchmark, early messages are more extreme to hasten agenda adoption, followed by moderate messages to align opinions with her agenda. The single influencer is worse off in networks where weights agents attach to their own opinions are farther apart from the weights agents attach to their neighbors' opinions because such networks have uniformly larger eigenvalues than a more balanced network. With multiple influencers, if they have the same impact on agents, consensus emerges in any network as the *average agenda* of the influencers. If they have different impacts, in symmetric networks, consensus still emerges but it is closer to the agenda of the more impactful influencer. In asymmetric networks, asymmetric influencers often target

*We thank Simon Board, Rahul Deb, Jeff Ely, Alex Frankel, Alastair Langtry, Vitor Farinha Luz, Shunya Noda, Jesse Perla, Mike Peters, Marit Rehavi, Sudipta Sarangi, Euncheol Shin, Philip Solimine, Kyungchul (Kevin) Song, Wing Suen, and Mark Whitmeyer for helpful comments. We are especially grateful for extensive discussions with Ben Golub, Matt Jackson and Li, Hao. We thank James Yu for outstanding research assistance. We also thank the seminar participants at U of T, the theory lunch workshop at UBC, the International Workshop on Economic Theory, the NSF Network Science and Economics Conference, and the 2023 Asia Meeting of the Econometric Society for great feedback. Wei Li thanks SSHRC Insight Grant 435-2018-0265 for financial support. Email: wei.li@ubc.ca; tanxu@uw.edu.

different subgroups with differing intensity, generating perpetual disagreement and polarization.

JEL: D85, D83, C73.

Keywords: optimal dynamic intervention in social networks, strategic competition among influencers, consensus and perpetual disagreements, general non-zero sum N -player LQ games.

1 Introduction

Misinformation and disinformation are rampant in social networks, significantly influencing the opinions of those exposed to them. According to global data, there were 5,613 distinct misinformation stories reported from the beginning of the Covid pandemic until the end of December 2020.¹ Twitter data from the 2016 US presidential campaign revealed that Russian disinformation efforts followed a "firehose of falsehood" model, featuring high-frequency, multi-channel, and continuous messages without regard to consistency.² Fake online reviews have been employed to promote businesses, deceiving unsuspecting consumers into purchasing inferior products. For instance, Tripadvisor's 2021 report concluded that nearly 1 million reviews (3.6%) on the site were fraudulent. The impact of such misinformation and disinformation campaigns is tangible, as they influence people's decisions regarding vaccination, voting choices, purchasing behavior, and so on. This paper studies how influencers compete to persuade a network of agents over time to adopt their agendas, whether beneficial or harmful.

Two novel features set our model apart. First, influencers face a trade-off between short-term and long-term outcomes as they seek to maximize their total discounted payoffs. They send more extreme messages early on, despite their higher immediate cost, to accelerate agenda adoption and increase long-term payoffs.³ Second, we study dynamic competition among influencers, including both strategic ones and mechanical entities such as spam bots. We find two sufficient conditions under which consensus among agents is sustained. When influencers have equal impact on agents, the consensus is the average agenda of all influencers. When influencers have different impact on agents, consensus closer to the stronger influencer's agenda emerges in symmetric networks. More generally, we show perpetual disagreement and polarization often emerge when influencers have different impacts and the network is asymmetric.

As a benchmark, we study how a single influencer sends messages to the agents in each period to minimize the total (discounted) quadratic distance between agents' opinions and

¹See "Localized Misinformation in a Global Pandemic: Report on COVID-19 Narratives around the World", March 25, 2021.

²Rand: The Russian "Firehose of Falsehood" Propaganda Model: Why It Might Work and Options to Counter It.

³This approach diverges from the large network learning literature, which primarily focuses on the long-run outcomes (see [DeMarzo, Vayanos, and Zwiebel \(2003\)](#), [Golub and Jackson \(2010\)](#), among many others).

her agenda. For example, a government tries to persuade agents that a vaccine is effective. In addition to listening to the government’s messages, agents follow the DeGroot learning rule by taking a weighted average of her neighbors’ opinions, and the weights form a network matrix. We find that the government’s optimal message in each period is linear in the agents’ current opinions. She sends the most extreme—and most costly—messages initially to push against the agents’ initial opinions. The messages become increasingly moderate and converge to zero as opinions converge to her agenda. To obtain these results, we present the government’s payoff as a quadratic function of initial opinions, weighted by an endogenous matrix capturing her total discounted payoff given her optimal strategy.⁴ While equilibrium characterization is obtained, the interdependence between the agents’ opinions and the influencers’ messages complicates its economic interpretation. Methodologically, we simplify the analysis by suitably decomposing the network matrix and transforming the problem into one with independence among dimensions.⁵

Simple comparative statics follow our decomposition immediately. In symmetric networks, the government’s early messages become more extreme, and later messages less so, if the discount rate is higher, or if the cost is lower. Moreover, divergence of initial opinions from her agenda such as the average opinion, and to a lesser extent, differences in initial opinions, reduces her payoff. Higher eigenvalues of the network matrix lead to more persistent initial opinions, which in turn diminish her payoff. For instance, consider an imbalanced network in which the weights agents attach to their own opinions are far away from the weights they attach to their neighbors. Then, either the agents are too stubborn that initial opinions persist, or the agents are too impressionable that opinions fluctuate, both slowing down agenda adaption and reducing the government’s payoff. The opposite is true in a balanced network.

Simple policy analysis also follows. We compare the government’s optimal messages and payoffs across different interventions: one-shot intervention (initial message only), myopic dynamic intervention (maximizing the next-period payoff, due to reelection pressure) and optimal dynamic intervention. In symmetric networks, she sends the most extreme initial message in one-shot intervention, followed by optimal dynamic intervention and then myopic intervention. Her later messages are less moderate in optimal dynamic intervention

⁴Formally, this matrix is the solution to the influencer’s discrete-time algebraic Ricatti equation, a nonlinear fixed point equation.

⁵Specifically, we use eigendecomposition for symmetric networks or singular value decomposition for general networks, which is especially useful when a matrix has low ranks.

than those in myopic intervention. Interestingly, myopic intervention always outperforms one-shot intervention, with the difference being largest at intermediate cost levels.

Next, we examine strategic interactions among multiple influencers who have equal impacts, as measured by the weight agents assign to an influencer's messages. Remarkably, a consensus emerges wherein all agents believe in the average agenda of the influencers, even in asymmetric networks where some agents are more important than others. In this scenario, each influencer's strategy remains an affine function of the current opinions: the slope determines opinion convergence speed while the constant term steers the opinions toward an influencer's agenda. Equal impact means that all influencers' strategies share the same slope, with only the constant terms differing based on their agendas. While the influencers do target more important agents with more extreme messages, equal impacts makes them do so uniformly. In the limit, their messages offset each other and their agendas are equally weighted in the consensus. As a result, influencers with agendas closer to the average get higher payoffs and vice versa. In our example, the government trying to persuade agents about vaccine efficacy is unlikely to succeed despite paying a high cost if all other influencers push extreme anti-vaccine agendas. But it may do well when facing a diverse group of influencers with varied agendas.

Our result also shows why like-minded influencers may flock to the same network. In the limit, messages from influencers whose agendas are on the same side of the average agenda are strategic substitutes; otherwise, they are strategic complements. In the case when all influencers have the same agenda, their messages are identical and serve as strategic substitutes. Persuading the agents then becomes a public good, and each influencer wants to free ride on the others. As the number of such influencers increases, everyone can send less extreme (and cheaper) messages while accelerating opinion convergence, leading to a higher payoff.

When influencers have different impacts, the influencer with a stronger impact adopts an optimal strategy with a steeper slope, pushing harder against the agents' opinions in each period and accelerating convergence. In symmetric networks, opinions eventually reach a consensus that aligns more closely with the agenda of the stronger influencer. As agents are equally important in symmetric networks, they are treated equally by each influencer in the limit, and the stronger influencer prevails. In contrast, consensus is often impossible to achieve in general asymmetric networks, with agents disagreeing with each other and both influencers. We demonstrate this result in a uniform opinion leader network, where every

agent receives the same opinion weight from all agents, and the agent receiving the highest weight is the opinion leader. In this scenario, the limit opinions of all agents align more closely with the agenda of the stronger influencer, with the opinion leader being the closest. The distance between an agent’s limit opinion and the stronger influencer’s agenda increases as the agent becomes less important in the network. Intuitively, both influencers find it optimal to send more extreme messages to the opinion leader, but the stronger influencer has a comparative advantage in persuading him. Recognizing that competing with the stronger influencer on the opinion leader is not optimal, the weaker influencer diverts more effort on the opinion followers. As a result, the agents are differentially "targeted" by the influencers, leading to permanent disagreements. Such disagreement arises *endogenously* and does not hinge on stubborn agents with fixed opinions.

Throughout this paper, we use the simplest model possible to highlight the economic insights in the messaging game among influencers, in Appendix B, we explore a considerably more general model that allows for asymmetric payoff functions and general impact matrices (not just a scalar) for each influencer. In this extended analysis, we characterize the necessary equilibrium conditions and derive the equations governing the evolution of opinions.

Our paper draws from and contributes to several large literature. Most closely related to our model is the recent and growing literature on network intervention/targeting. [Galeotti, Golub, and Goyal \(2020\)](#) consider a static game in which the social planner optimally intervenes by changing agents’ private returns to investment, which exhibit strategic spillovers in a network. Similarly, [Jeong and Shin \(2022\)](#) explore how a designer optimally implants initial opinions to align the network’s limit opinions with her agenda. We differ from them by considering dynamic intervention, allowing influencers to send messages to agents over time. [Grabisch, Mandel, Rusinowska, and Tanimura \(2018\)](#) study two strategic agents, each targeting one agent in a network to influence agents’ limit beliefs. [Sadler \(forthcoming\)](#) examines influence campaigns within networks, and [Vohra \(2023\)](#) studies how two strategic influencers, concerned about the network’s limit belief, choose the probabilities to influence the agents. In contrast, in our model, each influencer aims to maximize her total discounted payoff, caring about both short-term and long-term opinions. [Galeotti and Goyal \(2009\)](#) investigate the intervention of a single strategic agent who knows only the distribution of agents’ degrees in a network, while [Bloch and Shabayek \(2023\)](#) study optimal targeting when the planner lacks knowledge of the identities of agents occupying different network positions. In our model, the influencer possesses precise knowledge of the network’s structure and the

agents’ identities, enabling her to target each agent with a different message.

Our paper is also related to both the classic literature on linear quadratic (LQ) optimal control problems, as summarized by [Anderson and Moore \(1989\)](#) and [Bertsekas \(2017\)](#), and the more recent literature on LQ network games. For instance, [Ballester, Calvo-Armengol, and Zenou \(2006\)](#) demonstrate in a network game with LQ payoff functions that the Nash equilibrium action of each player is proportional to her Bonacich centrality. In contrast, in our model, the game is among the strategic influencers who are not part of the network. Various papers, including [Papavassilopoulos and Olsder \(1984\)](#), [Freiling, Jank, and Abou-Kandil \(1996\)](#), and [Engwerda \(1998\)](#), have studied two-player or zero-sum LQ games. However, their focus is primarily on the existence of equilibrium and computation algorithms, rather than exploring strategic interactions and economic implications as in our model.

In addition, our paper is related to the vast literature on learning in networks and opinion dynamics. In our model, agents learn naively and update their opinions according to the learning rule proposed by [DeGroot \(1974\)](#).⁶ [Gale and Kariv \(2003\)](#) and [Mossel, Sly, and Tamuz \(2015\)](#) consider the theoretical benchmark of fully rational agents updating by Bayes’ rule, and many recent papers explore quasi-Bayesian learning rules in which agents are boundedly rational.⁷ Since quasi-Bayesian learning can be cognitively and computationally demanding, agents in lab and field experiments often exhibit very limited cognitive ability.⁸ Given that our focus is on how strategic influencers compete among themselves to influence agents over time, we assume agents in the network learn naively for tractability.

2 A Model of Influencers

We introduce our main model and defer further discussions to Section 2.1. Agents are connected in a network $G = (\mathcal{N}, A)$, where $\mathcal{N} = \{1, \dots, N\}$ represents the set of agents and $A = [A_{ij}]$ is a real-valued $N \times N$ row stochastic matrix. Each element A_{ij} represents

⁶Variants of DeGroot learning has been analyzed by [Friedkin and Johnsen \(1990\)](#), [DeMarzo, Vayanos, and Zwiebel \(2003\)](#), [Golub and Jackson \(2010\)](#), and [Ghaderi and Srikant \(2014\)](#). See Chapter 7 of [Jackson \(2008\)](#).

⁷See [Bala and Goyal \(1998\)](#), [Alatas, Banerjee, Chandrasekhar, Hanna, and Olken \(2016\)](#), [Molavi, Tahbaz-Salehi, and Jadbabaie \(2018\)](#), [Levy and Razin \(2018\)](#), [Mueller-Frank and Neri \(2019\)](#), [Li and Tan \(2020\)](#), [Li and Tan \(2021\)](#), and [Della Lena \(2022\)](#), among many others.

⁸See [Anderson and Holt \(1997\)](#), [Celen and Kariv \(2004\)](#), [Alevy, Haigh, and List \(2007\)](#), [Cai, Chen, and Fang \(2009\)](#), [Mobius, Phan, and Szeidl \(2015\)](#), [Bai, Golosov, Qian, and Kai \(2015\)](#), [Enke and Zimmermann \(2019\)](#), [Chandrasekhar, Larreguy, and Xandri \(2020\)](#), and [Grimm and Mengel \(2020\)](#), among others.

the weight assigned by agent i to the opinion of agent j . The weights satisfy $0 \leq A_{ij} \leq 1$ and $\sum_j A_{ij} = 1$. Time is discrete: $t = 0, \dots, T$ and $T \leq \infty$, indicating that the time horizon can be finite or infinite. Each agent holds an initial opinion $x_0^i \in \mathbb{R}$ at $t = 0$. In every period t , the vector of all agents' opinions is denoted by $\mathbf{x}_t \in \mathbb{R}^N$, where each element x_t^i is agent i 's opinion at period t .⁹ Agents are naive in that they update their opinions using the classic DeGroot updating rule. Each agent, referred to generically as *he* throughout the paper, updates opinions based on computing a weighted average of the opinions of the agents he listens to, as represented by matrix A .

Moreover, each agent also listens to $M \geq 1$ influencers from outside the network A denoted by $m \in \{1, \dots, M\}$. In period t , influencer m , referred to generically as *she* throughout the paper, sends messages $\mathbf{r}_t^m = (r_t^{m,1}, r_t^{m,2}, \dots, r_t^{m,N})' \in \mathbb{R}^N$. Each element $r_t^{m,i}$ is the (distinct) message influencer m sends to agent i in period t . Starting with any initial opinions $\mathbf{x}_0$, the agents' opinions are updated using both the opinions of all agents and the messages from influencers such that

$$\mathbf{x}_{t+1} = A\mathbf{x}_t + \alpha_m \sum_{m=1}^M \mathbf{r}_t^m, \quad (1)$$

where $\mathbf{x}_t$ is the vector of opinions and $\mathbf{r}_t^m$ is the vector of influencer m 's messages at t . The scalar $\alpha_m > 0$ captures the impact of influencer m 's messages. A higher value of α_m indicates greater impact. For simplicity, we assume that each influencer's messages have a uniform impact of α_m on all agents.¹⁰

Our model studies the messaging game among M influencers, each with a known agenda denoted as $b^m \in \mathbb{R}$, $m = 1, \dots, M$. The influencers send costly messages to align the agents' opinions with their respective agendas, subject to agents updating opinions according to (1). In each period t , influencer m 's stage payoff is given by

$$u_t^m(\mathbf{x}_t) = -(\mathbf{x}_t - \mathbf{b}^m)'(\mathbf{x}_t - \mathbf{b}^m) - c(\mathbf{r}_t^m)' \mathbf{r}_t^m. \quad (2)$$

Here, $\mathbf{b}^m = b^m \mathbf{1}$ is the vector with all entries equal to b^m , and $c(\mathbf{r}_t^m)' \mathbf{r}_t^m$ represents the cost

⁹We use boldface lowercase for vectors and capital letters for matrices. All vectors are column vectors.

¹⁰We use the following rule to simplify notations: the time index t is always presented as a subscript; the influencer index m is presented as a superscript (except for the impact α_m , as it frequently appears squared); and the agent index i is presented as a superscript when accompanied by t (e.g., $r_t^{m,i}$), and as a subscript when accompanied by m but without t (e.g., b_i^m).

of sending messages.¹¹ The discount rate is $\delta \in (0, 1)$, and influencer m maximizes her total discounted payoff by choosing messages $\mathbf{r}_t^m$ in each period t .

We focus on $T = \infty$ to be succinct. When $M \geq 2$, influencers choose messages ($\mathbf{r}_t^m$ for influencer m) simultaneously and independently in each period t , leading to a dynamic game of complete but imperfect information. Let $\mathbf{r}_t = \{\mathbf{r}_t^1, \dots, \mathbf{r}_t^M\}$. The history of the play at period $t \geq 1$ is $h_t = \{\mathbf{x}_0, \mathbf{r}_0, \dots, \mathbf{x}_{t-1}, \mathbf{r}_{t-1}\}$, and $h_0 = \{\mathbf{x}_0\}$. Let H_t be the collection of such histories at t . Then, influencer m 's strategy is a sequence of mappings $(\sigma_t^m)_{t=0}^\infty$ where $\sigma_t^m : H_t \rightarrow \mathbb{R}^N$ for all $t \geq 0$. For any given h_t , the continuation payoff of influencer m is

$$-\sum_{\tau=t}^{\infty} \delta^\tau ((\mathbf{x}_\tau - \mathbf{b}^m)'(\mathbf{x}_\tau - \mathbf{b}^m) + c(\sigma_\tau^m(h_\tau))' \sigma_\tau^m(h_\tau)), \quad (3)$$

where for all $\tau \geq t$, $h_{\tau+1} = \{h_\tau, \mathbf{x}_\tau, (\sigma_\tau^j(h_\tau))_{j=1}^M\}$ and $\mathbf{x}_\tau = A\mathbf{x}_{\tau-1} + \sum_{j=1}^M \alpha_j \sigma_\tau^j(h_\tau)$ are iteratively defined from h_t . A subgame perfect equilibrium is a profile of strategies $(\sigma_t^m)_{t=0}^\infty$ such that for any h_t , every influencer m 's strategy $(\sigma_\tau^m(h_\tau))_{\tau=t}^\infty$ maximizes (3) given all other influencers' strategies and the opinion updating rule. Given our interest in network influence, we restrict attention to strategies that depend only on the current opinions instead of the entire history. Specifically, we look for Markov Perfect equilibrium (MPE), which is a subgame perfect equilibrium where $\sigma_t^m(h_t) = \sigma_\tau^m(\tilde{h}_\tau)$ for every m, t , and τ such that the two histories have the same state variable $\mathbf{x}_{t-1} = \tilde{\mathbf{x}}_{\tau-1}$.

2.1 Remarks on Model Assumptions

Agents' opinion updating rule. Agents update their opinions using a combination of their neighbors' opinions and messages from influencers according to updating rule (1). This rule differs from the classic DeGroot learning, where the agents' opinions depend solely on their neighbors' opinions according to network A . In our model, influencers send strategic messages in every period. We will show that these messages drive the limit opinions, which depend on the network A and the influencers' payoff parameters. The agents' initial opinions affect the influencers' payoffs, but have no impact on the limit opinions.

Time horizon. The infinite-horizon problem shares many common features with the finite-horizon problem, except that the optimal messages are time dependent in the latter problem.

¹¹We use $\mathbf{1}$ for the vector in which every entry is 1 and $\mathbf{e}_i$ for the basis in which the i -th entry is 1 and all other entries are 0.

To be succinct, we present the results for the infinite-horizon problem while using the finite-horizon problem to explain intuitions due to its simplicity. Specifically, when T is finite, influencer m maximizes her total discounted payoff by choosing messages $\mathbf{r}_t^m$ in each period $t < T$, and no influencer sends messages in the terminal period. The history of the play and the influencer strategies extend easily to the finite-horizon model. But the finite horizon introduces non-stationary strategies which depends on the remaining time periods. Therefore we require the influencers to choose the same strategies for the same current opinions period by period, that is, $\sigma_t^m(h_t) = \sigma_t^m(\tilde{h}_t)$ for every m, t such that the two histories have the same state variable $\mathbf{x}_{t-1} = \tilde{\mathbf{x}}_{t-1}$.

Can influencers reach all agents? The opinion updating rule (1) assumes that each influencer m has the ability to send a message to every agent with the same impact α_m . In reality, however, some agents may be stubborn and discard these messages, while others may be more receptive to like-minded influencers. In Appendix B, we present a more general model that allows messages from influencer m to have different impacts on each agent. When the impact on an agent is zero, it implies that influencer m cannot reach that particular agent, but she may still influence that agent indirectly through his neighbors.

Linear quadratic payoffs with no uncertainty. Each influencer's payoff is decreasing and quadratic in the difference between agents' opinions and her agenda. This payoff structure, along with quadratic costs, leads to the optimal strategies being linear in agents' opinions.¹² Introducing small, independently distributed Gaussian noise into the opinion updating process introduces uncertainty, but it adds only a variance term into the influencers' strategies and does not affect our results qualitatively.

More general model. We use the simplest model possible in the text to highlight the economic insights. In Appendix B, we extend the analysis to a more general multiple-influencer model and establish the necessary conditions for the existence of a MPE. In addition to incorporating the previously mentioned varying impacts, we introduce three generalizations. First, each influencer's payoff assigns a weight to the difference between an agent's opinion and her agenda, and this weight can vary across both influencers and agents. Second, we include varying costs associated with sending messages to different agents. Third, we allow each influencer to have a known initial position, and the cost of her messages depends on their deviation from this position. The spirit of analysis is similar and

¹²Nonlinear models often necessitate linear approximations due to the absence of analytical solutions.

the necessary conditions we derived can be used numerically to solve for the equilibrium.

3 Benchmark: Single Influencer

We begin with a single influencer who tries to persuade the agents, such as the government promoting the effectiveness of vaccines. We normalize the government's agenda to 0; this is without loss because only the distance of each agent's opinion from her agenda matters to her payoff. We also drop the influencer index m and let $\alpha = 1$ for simplicity. We use this model to demonstrate our main solution method—eigen/singular decomposition of matrix A , which simplifies analysis greatly. The same decomposition is used later to show existence and uniqueness of equilibrium in the multiple-influencer model.

Denote the government's optimal value at $t = 0$ as $v(\mathbf{x}_0)$, which is the highest payoff obtained when she chooses messages optimally. The continuation value from period t onward is $v(\mathbf{x}_t)$, and is well-defined and bounded (as proven later in Proposition 1). For all $t \geq 0$, by the principle of optimality, we have

$$v(\mathbf{x}_t) = \max_{\mathbf{r}_t} \{ -\mathbf{x}_t' \mathbf{x}_t - c \mathbf{r}_t' \mathbf{r}_t + \delta v(A \mathbf{x}_t + \mathbf{r}_t) \}.$$

The problem can be solved using standard optimal control method.¹³

3.1 Symmetric Network A

The Transformed Problem: Regaining Independence

Opinion evolution in the network is complicated because the agents' opinions depend on the government's messages, which in turn depends on the agents' opinions. This high degree of interdependence hinders comparative statics study and policy analysis. To gain economic insight, we transform the problem into one with independence across dimensions of opinions by decomposing A . We begin with symmetric network A which can be decomposed as $A = UDU'$, where $UU' = I$ and D is the diagonal matrix of A 's eigenvalues, $\lambda_1 \geq \lambda_2 \dots \geq \lambda_N$. Each column of U is an eigenvector $\mathbf{u}_i$ corresponding to

¹³See [Anderson and Moore \(1989\)](#) and [Bertsekas \(2017\)](#) among many others for details. This method has been used extensively in [Hansen and Sargent \(2013\)](#) to study macroeconomics questions. But there are often no simple analytical solutions. The decomposition method below allows us to perform comparative statics and to do policy analysis much easily.

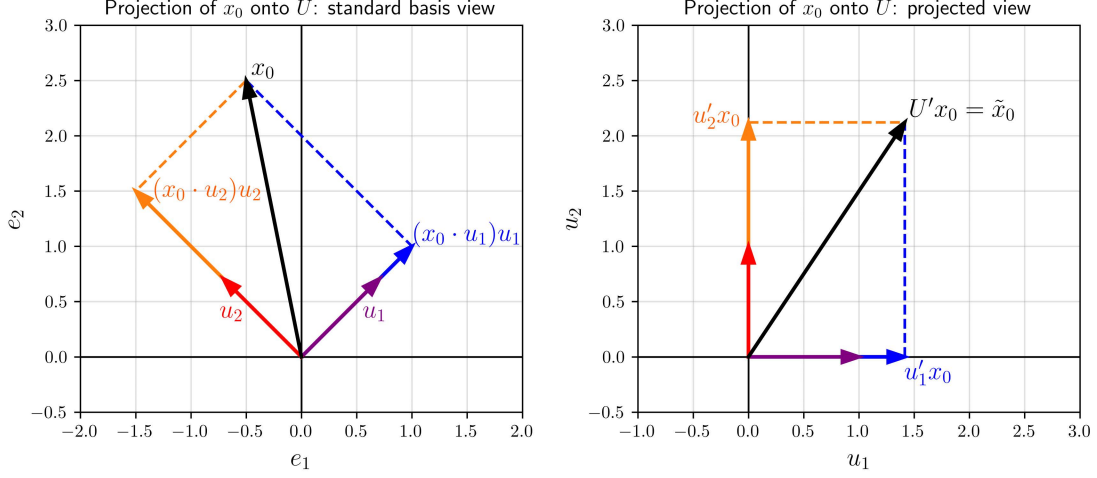


Figure 1: Dimension transformation using eigendecomposition.

eigenvalue λ_i and thus $AU = DU$. Any real symmetric $N \times N$ matrix has N independent eigenvectors (whether the eigenvalues are repeated or not) which form an orthonormal basis, that is, $\|\mathbf{u}_i\| = 1$ and $\mathbf{u}_i \cdot \mathbf{u}_j = 0$ for all $i \neq j$. We then transform the variables under study by letting $\tilde{\mathbf{x}}_t = U' \mathbf{x}_t$ and $\tilde{\mathbf{r}}_t = U' \mathbf{r}_t$.¹⁴ Using the transformed variables, we can rewrite the value function as

$$\tilde{v}(\tilde{\mathbf{x}}_t) = \max_{\tilde{\mathbf{r}}_t} \{ -\tilde{\mathbf{x}}_t' \tilde{\mathbf{x}}_t - c \tilde{\mathbf{r}}_t' \tilde{\mathbf{r}}_t + \delta \tilde{v}(D \tilde{\mathbf{x}}_t + \tilde{\mathbf{r}}_t) \}. \quad (4)$$

We illustrate this decomposition in a two-agent network in Figure 1. Because $\mathbf{u}_1$ and $\mathbf{u}_2$ are orthogonal, in the transformed problem, the projected opinions are located on the two axes. Absent of the messages, the projected opinions evolve according to $\tilde{\mathbf{x}}_{t+1} = D \tilde{\mathbf{x}}_t$, and D is diagonal. Thus, the projected opinions' evolution is independent across dimensions.

Optimal Messages and Comparative Statics

The government's value function in (4) can be written as a quadratic form:

$$\tilde{v}(\tilde{\mathbf{x}}_t) = -\tilde{\mathbf{x}}_t' \tilde{K}^* \tilde{\mathbf{x}}_t,$$

¹⁴Specifically, $\tilde{x}_0^1 = \mathbf{u}_1' \mathbf{x}_0$ is the projection of $\mathbf{x}_0$ on $\mathbf{u}_1$, the first eigenvector of A . Similarly, $\tilde{r}_0^k = \mathbf{u}_k' \mathbf{r}_0$ is the projection of $\mathbf{r}_0$ on eigenvector $\mathbf{u}_k$. To distinguish the transformed problem from the original problem, we denote the variables in the transformed problem by using $\tilde{y}$ if the original variable is y .

where $\tilde{K}^*$ is an endogenous matrix summarizing how her future payoff depends on current opinions.¹⁵ Inserting the above value function form to (4), we have

$$\tilde{v}(\tilde{\mathbf{x}}_t) = -\tilde{\mathbf{x}}_t' \tilde{K}^* \tilde{\mathbf{x}}_t = \max_{\tilde{\mathbf{r}}_t} \left\{ -\tilde{\mathbf{x}}_t' \tilde{\mathbf{x}}_t - c \tilde{\mathbf{r}}_t' \tilde{\mathbf{r}}_t - \delta (D \tilde{\mathbf{x}}_t + \tilde{\mathbf{r}}_t)' \tilde{K}^* (D \tilde{\mathbf{x}}_t + \tilde{\mathbf{r}}_t) \right\}.$$

Because the government's payoff is quadratic and concave in $\tilde{\mathbf{r}}_t$, her optimal message is uniquely derived from the first-order condition and given by (6) in the proposition below. Substituting the optimal $\tilde{\mathbf{r}}_t$ in (6) back into the value function above, we can see the *infinite-horizon Ricatti matrix* $\tilde{K}^*$ is the unique positive definite matrix that solves

$$\tilde{K}^* = I + \delta c (\delta \tilde{K}^* + cI)^{-1} \tilde{K}^* D^2. \quad (5)$$

That is, $\tilde{K}^*$ is the solution to a fixed-point problem. Characterization of the agents' opinions follows immediately.

Proposition 1. *Suppose A is symmetric and is decomposed as $A = UDU'$. The influencer's optimal message is*

$$\tilde{\mathbf{r}}_t = \tilde{L}^* \tilde{\mathbf{x}}_t = \tilde{L}^* (\tilde{L}^* + D)^t \tilde{\mathbf{x}}_0, \quad (6)$$

where $\tilde{L}^* = -(\delta \tilde{K}^* + cI)^{-1} \delta \tilde{K}^* D$ and $\tilde{K}^*$ is the unique positive definite Ricatti matrix defined by (5). All agents agree with the government's agenda in the limit: $\mathbf{x}_\infty = \mathbf{0}$.

The government's optimal strategy is linear in period- t opinions as described in (6),¹⁶ where the slope $\tilde{L}^*$ measures how hard she pushes agents' opinions towards her agenda 0. This slope depends on the Ricatti matrix $\tilde{K}^*$, which is the solution to a nonlinear fixed-point matrix equation and may not have a simple analytical form in general. When A is symmetric, however, we show the unique solution of (5) is a *diagonal* matrix and we can solve each of

¹⁵For intuition about the quadratic form, consider a simple two-period model, where the value function takes the form $\tilde{v}_t(\tilde{\mathbf{x}}_t) = -\tilde{\mathbf{x}}_t' \tilde{K}_t \tilde{\mathbf{x}}_t$, $t = 0, 1$ where the matrix $\tilde{K}_t$ depends on the period t . The terminal value is $\tilde{v}_1(\tilde{\mathbf{x}}_1) = -\tilde{\mathbf{x}}_1' \tilde{\mathbf{x}}_1$ in which $\tilde{K}_1 = I$. In period 0, the influencer chooses $\tilde{\mathbf{r}}_0$ to maximize her total payoff. Substituting the optimal messages $\tilde{\mathbf{r}}_0 = -\delta/(\delta + c) D \tilde{\mathbf{x}}_0$ back into the value function gives us $\tilde{K}_0$ and a quadratic value function $\tilde{v}_0(\tilde{\mathbf{x}}_0)$. The infinite-horizon case is similar except that this matrix $\tilde{K}^*$ is constant.

¹⁶We say the influencer's strategy is *linear* when $\tilde{\mathbf{r}}_t = \tilde{L}^* \tilde{\mathbf{x}}_t$ and *affine* when $\tilde{\mathbf{r}}_t = \tilde{L}^* \tilde{\mathbf{x}}_t + \tilde{\mathbf{l}}^*$, that is, the latter allows for a non-zero constant term.

the N dimensions separately.¹⁷ As a result, the slope $\tilde{L}^*$ is also diagonal and pd, where

$$\tilde{L}_{jj}^* = -\frac{\delta \tilde{K}_{jj}^* \lambda_j}{\delta \tilde{K}_{jj}^* + c}.$$

Moreover, because each diagonal entry of $\tilde{K}_{jj}^*$ is increasing in $|\lambda_j|$, the government pushes hardest against the first dimension of the projected opinions, second hardest against the second dimension, and so on. Note that the projected opinions are related to the agents' opinions in an interesting way.: $\tilde{x}_t^j = \mathbf{u}_j' \mathbf{x}_t$, where each $\mathbf{u}_j$ is the normalized eigenvector associated with the j th largest eigenvalue of A . Because A is stochastic, $\mathbf{u}_1 = (1/\sqrt{N})\mathbf{1}$, and thus the first dimension of the projected opinion is proportional to the average opinion in that period. Intuitively, this eigenvector implies that all agents' opinions are equally important in the long run, and thus the influencer pushes the hardest against the average opinion. The other eigenvectors have descending importance, for example, in a two-agent network, the second eigenvector is proportional to $(1, -1)'$ and thus the second dimension of the projected opinion is the difference in agent's opinions. That is, the influencer also wants to reduce the dispersion of the agents' opinions, which matters to the convergence speed. Moreover, given the optimal message, we have

$$\tilde{\mathbf{x}}_t = D\tilde{\mathbf{x}}_{t-1} + \tilde{\mathbf{r}}_{t-1} = \left(I - (\delta \tilde{K}^* + cI)^{-1} \delta \tilde{K}^* \right) D\tilde{\mathbf{x}}_{t-1} = c(\delta \tilde{K}^* + cI)^{-1} D\tilde{\mathbf{x}}_{t-1}.$$

Iterate and we have $\tilde{\mathbf{x}}_t = c^t (\delta \tilde{K}^* + cI)^{-t} D^t \tilde{\mathbf{x}}_0$. Because $c(\delta \tilde{K}^* + cI)^{-1} D$ is diagonal and each diagonal entry is strictly smaller than 1 (in absolute value), in each dimension, $|\tilde{x}_t^j|$ decreases strictly in t and all opinions converge to the government's agenda 0.

We now show how the government's messages and payoffs depend on several key parameters such as δ , c and the agents' initial opinions.

Proposition 2. *Suppose A is symmetric and consider every transformed dimension j .*

(1) If δ increases or if c decreases, the government sends more extreme messages initially ($|\tilde{r}_0^j|$ increases), but less extreme messages when t is sufficiently large ($|\tilde{r}_t^j|$ decreases). Her messages are more extreme ($|\tilde{r}_t^j|$ increases for all t) if $|\tilde{x}_0^j|$ increases. (2) Her value $\tilde{v}(\tilde{\mathbf{x}}_0)$ decreases in c , δ , and in each $|\tilde{x}_0^j|$.

¹⁷While our proof applies directly to the infinite-horizon setting, the intuition can be seen clearly from the finite-horizon case. Recall that $\tilde{K}_T = I$ is diagonal; when $\tilde{K}_{t+1}$ is diagonal, $\tilde{K}_t = I + \delta c(\delta \tilde{K}_{t+1} + cI)^{-1} \tilde{K}_{t+1} D^2$ is also diagonal. The infinite-horizon Ricatti matrix $\tilde{K}^*$ is the limit of $\tilde{K}_{T-t}$ as $t \rightarrow \infty$ and inherits the properties of $\tilde{K}_t$ being diagonal.

We use the absolute value of the government's messages to measure how extreme these messages are. In each dimension j ,

$$|\tilde{r}_t^j| = |\tilde{L}_{jj}^*| |\tilde{x}_t^j| = \frac{\delta \tilde{K}_{jj}^* |\lambda_j|}{\delta \tilde{K}_{jj}^* + c} \left(\frac{c |\lambda_j|}{\delta \tilde{K}_{jj}^* + c} \right)^t |\tilde{x}_0^j|. \quad (7)$$

Clearly, for given parameters, her message always becomes less extreme as times go on because the opinions converge to 0. But the trajectory of her messages responds differently to the changes in parameters. First, her initial message becomes more extreme and later messages become less extreme if δ increases or if c decreases, reflecting a trade off between short term and long term payoffs. By the Envelope Theorem, it is easy to see that $\tilde{K}_{jj}^*$ increases in δ and c . That is, she is worse off if the future payoff is more important, or if the cost is higher. We can also show that $\tilde{K}_{jj}^*/c$ decreases in c . It follows that the (absolute value of) the slope increases in δ and decreases in c . Therefore the initial message, which depends on the slope of the strategy, becomes more extreme as the future becomes more important or as the cost decreases. Her later messages, however, have the opposite responses. The second exponential term in (7), which measures the evolution of opinions, decreases in δ and increases in c . That is, as δ increases, the slope is steeper and the initial message becomes more extreme; as a result, the opinions converge to 0 faster and later messages are less extreme. Similarly, if c increases, the initial message becomes less extreme, but the opinions converge to 0 slower and she has to send more extreme messages later. Next, the government's optimal value $\tilde{v}(\tilde{x}_0)$ decreases in $\tilde{x}_0$, which represents the initial conditions, because she has to overcome a lot of disagreement. As discussed above, if the average opinion becomes farther away from her agenda, the influencer is worse off. This observation is also true, to a lesser extent, when agents disagree with each other more.

How does the network structure affect the influencer? Let the spectrum of A be $(\lambda_1, \dots, \lambda_N)$. Then we can ask two related questions. First, compare two networks that differ only in one eigenvalue $|\lambda_j|$; second, if a network changes systematically.

Observation 1. *Suppose A is symmetric. The government is worse off if any $|\lambda_j|$ increases.*

For $\varepsilon > 0$, consider perturbations of A : (1) Let $A(\varepsilon) = (1 - \varepsilon)A + \varepsilon J_N/N$, where J_N is an all one $N \times N$ matrix. Then, the eigenvalues of $A(\varepsilon)$ are $(\lambda_1, (1 - \varepsilon)\lambda_2, \dots, (1 - \varepsilon)\lambda_N)$.

(2) Suppose A is positive definite. Consider another symmetric positive definite row stochastic matrix P such that $\lambda_N^P > \lambda_2^A$.¹⁸ Let $A(\varepsilon) = (1 - \varepsilon)A + \varepsilon P$. Then, every eigenvalue of

¹⁸That is, the smallest eigenvalue of P is higher than the second largest eigenvalue of A .

$A(\varepsilon)$ is greater than the corresponding eigenvalue of A .

From (5), we can see each $\tilde{K}_{jj}^*$ decreases in $|\lambda_j|$, and thus the influencer is worse off if $|\lambda_j|$ increases because dimension j 's opinion is more persistent, and convergence of agents' opinions is slower. For instance, in any two-agent networks, the influencer gets a higher payoff from the network with a smaller $|\lambda_2|$. Similarly, any two networks differing in only the magnitude of one single eigenvalue can be payoff ranked. This answer is incomplete, however, because whenever network A changes, the entire set of eigenvalues changes accordingly. To gain insight, we illustrate perturbations of A in which the magnitude of every eigenvalue shifts in the same direction.

We call the first perturbation an “equalizing” perturbation, as it shifts each agent’s weight on himself closer to those on his neighbors. Consider a simple two-agent network A . If $A_{ii} > A_{ij}$, the introduction of $J_2/2$ decreases A_{ii} and increases A_{ij} . Then, opinions converge faster because every agent becomes less stubborn and opinions are less persistent. Conversely, if $A_{ii} < A_{ij}$, the introduction of $J_2/2$ increases A_{ii} and decreases A_{ij} , which again accelerates the convergence of opinion. Before the perturbation, the agents listen to their neighbors more than themselves, and thus the opinions fluctuate too much and convergence takes longer. In either case, as A_{ii} gets closer to A_{ij} , the influencer’s optimal payoff increases. We call the second perturbation a “polarizing” perturbation as it mixes matrix A with a more persistent matrix; then, convergence takes longer. To begin with, suppose A is combined with the identity matrix I , in which no agent listens to his neighbors; all the diagonal terms increase and all the eigenvalues increase as well. As a result, the influencer has to intervene more to persuade the agents in the perturbed network and gets a lower payoff. This result also holds for any matrix P that is similar to A .¹⁹ Even for matrices that are not similar to A , we can still find some conditions bounding the eigenvalues under which this result holds, as shown in part (2) of Observation 1.

3.2 Comparing Common Intervention Policies

The government is often constrained in how she can send her messages for institutional or practical reasons. For example, recent work by [Galeotti, Golub, and Goyal \(2020\)](#)

¹⁹If P is similar to A , then it can be diagonalized by the same eigenvector matrix, that is, $A = UDU'$ and $P = UDU'P$. Then if $A > (<)P$ in the positive (negative) definite sense, the absolute value of every eigenvalue of $A(\varepsilon)$ is smaller (greater) than the corresponding eigenvalue of A .

investigates optimal one-shot intervention where influencers can send one message only. In network learning and formation games, [Bala and Goyal \(1998\)](#) and [Watts \(2001\)](#) consider myopic strategies where influencers maximize their next period's payoffs. Myopic strategy is also a good proxy for a government facing reelection pressures. To enhance policy analysis and differentiate our model from existing ones, we define and compare four interventions: optimal dynamic intervention yielding payoff $\tilde{v}^*$, one-shot intervention yielding payoff $\tilde{v}^{os}$, myopic (dynamic) intervention yielding payoff $\tilde{v}^{mp}$, and no intervention yielding payoff $\tilde{v}^\emptyset$. Everything in this subsection follows the same superscript convention.

All these interventions can be analyzed easily in our framework. For myopic intervention, the government's optimal message is $\tilde{\mathbf{r}}_t^{mp} = -\frac{\delta}{\delta+c}D\tilde{\mathbf{x}}_t$, just as in a two-period model because she maximizes only the next period's payoff without considering any long-term effects. Under no intervention, opinions evolve according to $\tilde{\mathbf{x}}_t = D^t\tilde{\mathbf{x}}_0$. If A is strongly connected and aperiodic, then her total discounted payoff is $-\tilde{\mathbf{x}}_0'(I - \delta D^2)^{-1}\tilde{\mathbf{x}}_0$.²⁰ In the long run, the agents' consensus is $\tilde{\mathbf{x}}_0^1$, which is the average initial opinions. In one-shot intervention, the government is restricted to intervene only initially. The opinions then evolve without further intervention, i.e., $\tilde{\mathbf{x}}_1 = D\tilde{\mathbf{x}}_0 + \tilde{\mathbf{r}}_0$, and $\tilde{\mathbf{x}}_t = D\tilde{\mathbf{x}}_{t-1}$ for $t > 1$. Hence, her total discounted payoff resembles the no intervention case:

$$-\tilde{\mathbf{x}}_0'\tilde{\mathbf{x}}_0 - c\tilde{\mathbf{r}}_0'\tilde{\mathbf{r}}_0 - \delta\tilde{\mathbf{x}}_1'(I - \delta D^2)^{-1}\tilde{\mathbf{x}}_1.$$

Solving the FOC, we find that the optimal initial message for each dimension j is $\tilde{r}_{0,j}^{os} = -\frac{\delta\lambda_j}{\delta+c(1-\delta\lambda_j^2)}\tilde{x}_0^j$. We now compare these intervention policies.

Proposition 3. *Suppose A is symmetric, strongly connected and aperiodic. For any $\delta > 0$, $c \in (0, \infty)$ and $\tilde{\mathbf{x}}_0$, we have (1) the payoff ranking: $\tilde{v}^* > \tilde{v}^{mp} > \tilde{v}^{os} > \tilde{v}^\emptyset$;*

(2) the (absolute value of) message ranking: the government sends the most extreme initial message in one-shot intervention, less in optimal dynamic intervention, and the least in myopic intervention. After a cutoff period, the government sends less extreme messages in optimal dynamic intervention than those in myopic intervention.

In extreme cases, active interventions yield the same payoff. For instance, when $c = 0$, all agents agree with the influencer's agenda in one period, and when $c = \infty$ or $\delta = 0$, no intervention occurs. In all other cases, interventions differ in messages and payoffs. Clearly,

²⁰We need A to satisfy the additional assumptions so that opinions converge under no intervention, which is not necessary when there are influencers.

the government pushes the hardest in one-shot intervention, as she has only one chance to influence the agents. The comparison of messages between dynamic intervention and myopic intervention is more subtle. The initial messages are stronger in dynamic intervention, while later messages are stronger in myopic intervention. Recall from Proposition 1 that each message depends on both the slope and the current opinions. We can show that the slope is always higher in dynamic intervention, because the influencer is willing to push harder now to increase future payoffs. In the initial period, the opinions are the same in both interventions, thus only the slope matters and the initial message is stronger in dynamic intervention. As time goes on, the opinions converge faster to zero in dynamic intervention. After a cutoff period, the effect of lower (absolute values of) opinions in dynamic intervention outweighs the higher slope, and then the message becomes less extreme in dynamic intervention.

In terms of payoffs, optimal dynamic intervention clearly leads to the highest payoff, while no intervention leads to the lowest. The comparison between one-shot intervention and myopic intervention is less clear. The former takes long-term effect into account but sends only one message, and the latter focuses only on short-term payoffs but sends messages in every period. Recall that when $c = 0$ or $\delta = 0$, the influencer's payoffs under all interventions are the same. By the Envelope theorem, the influencer's payoff under all interventions decreases in δ and c . Our study reveals a positive and increasing payoff difference between myopic and one-shot intervention as δ rises. Myopic intervention ensures opinion convergence to the government's agenda, unlike one-shot intervention. The higher the δ , the lower is the payoff for opinions not converging to 0. In contrast, the payoff difference is non-monotonic in c , peaking at an intermediate cost where one-shot intervention cannot use extreme messages but myopic intervention can spread the cost over many periods. Moreover, if we compare the cumulative payoffs from *only the early periods*, myopic intervention may lead to a higher payoff than dynamic intervention where the government pushes harder initially to speed up convergence and increase her long-term payoff. As a result, myopic intervention may appeal to a government facing short-term reelection considerations.

3.3 General network A

We extend the analysis beyond symmetric networks to characterize optimal dynamic intervention more generally. As before, the government chooses a sequence of messages

$\{\mathbf{r}_0, \dots, \mathbf{r}_t, \dots\}$ to maximize her total discounted payoff: $-\sum_0^\infty \delta^t (\mathbf{x}'_t \mathbf{x}_t + c \mathbf{r}'_t \mathbf{r}_t)$. The following result extends Proposition 1 to the general network.²¹

Proposition 4. *The government's optimal strategy is $\mathbf{r}_t = L^* \mathbf{x}_t$, in which $L^* = -(\delta K^* + cI)^{-1} \delta K^* A$ and K^* is the unique pd solution of*

$$K^* = I + \delta A' (K^* - \delta K^* (\delta K^* + cI)^{-1} K^*) A. \quad (8)$$

Moreover, $v(\mathbf{x}_0) = -\mathbf{x}'_0 K^* \mathbf{x}_0$. All opinions converge to the government's agenda if δ is sufficiently high.

The optimal messages are still linear in current opinions

$$\mathbf{x}_t = (A + L^*) \mathbf{x}_{t+1} = (I - (\delta K^* + cI)^{-1} \delta K^*) A \mathbf{x}_{t-1} = c(\delta K^* + cI)^{-1} A \mathbf{x}_{t-1}.$$

Moreover, the value function is bounded because even under myopic strategy, $\mathbf{r}_t^{mp} = \frac{\delta}{\delta+c} A$ and $\mathbf{x}_t^{mp} = \frac{c}{\delta+c} A \mathbf{x}_{t-1}^{mp}$. The transition matrix $\frac{c}{\delta+c} A$ is stable since the absolute value of all its eigenvalues are strictly smaller than 1. Therefore the opinions converge under myopic strategy and the value $v(\mathbf{x}_t^{mp})$ converges to zero as time goes on. The government must do better under optimal strategy, and thus her value function is also bounded. When δ is sufficiently high,, a unique solution to the Ricatti equation (5) exists because $c(\delta K^* + cI)^{-1}$ is a stable matrix and thus the opinions converge: $\lim_{t \rightarrow \infty} \mathbf{x}_t \rightarrow \mathbf{0}$. Moreover, the initial message under optimal strategy is still more extreme than under myopic strategy:

$$\|\mathbf{x}_1\| \leq \|c(\delta K^* + cI)^{-1}\| \|A \mathbf{x}_0\| < c/(c + \delta) \|A \mathbf{x}_0\| = \|\mathbf{x}_1^{mp}\|,$$

where the second inequality is true because the spectrum norm of a symmetric matrix $c(\delta K^* + cI)^{-1}$ is equal to its largest eigenvalue, which is smaller than $c/(\delta + c)$.²²

4 Multiple Influencers and Long-Run Consensus

Many influencers may push their (possibly different) agendas at the same time as in Section 2. In this section, we examine conditions under which long-run consensus among agents is

²¹Note that eigendecomposition does not work for asymmetric networks. For this part, we return to the original problem; for example, opinions are $\mathbf{x}_t$ instead of $\tilde{\mathbf{x}}_t$ and the value function is v instead of $\tilde{v}$.

²²Because the dimensions are no longer independent, we cannot compare messages for each dimension as in Section 3.2; instead we use the Euclidean norm to measure the strength of a (message or opinion) vector. We also follow the convention of using spectrum norm for matrices.

sustained, and scenarios with perpetual disagreements are studied in the next section. We first show that a sufficient condition for consensus is equal impact among all influencers, that is, agents attach the same weight to each influencer's messages. Otherwise, when influencers have different impacts, consensus is sustained when the network is symmetric.

4.1 Influencers with Equal Impact

When all influencers have equal impact on the agents, $\alpha_m = \alpha$ in equation (1). We begin with symmetric network A and decompose it by $A = UDU'$ as in Section 3.1. By multiplying U' to all variables of updating equation (1), the opinion updating equation becomes

$$\tilde{\mathbf{x}}_{t+1} = D\tilde{\mathbf{x}}_t + \alpha \sum_{m=1}^M \tilde{\mathbf{r}}_t^m. \quad (9)$$

Influencer m 's projected agenda is $\tilde{\mathbf{b}}^m = U'\mathbf{b}^m = U'b^m\mathbf{1}$ and her stage payoff in period t is

$$\tilde{u}_t^m(\tilde{\mathbf{x}}_t) = -(\tilde{\mathbf{x}}_t - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_t - \tilde{\mathbf{b}}^m) - c(\tilde{\mathbf{r}}_t^m)' \tilde{\mathbf{r}}_t^m.$$

We conjecture (and later prove) that the continuation value for each influencer m in period t has the form

$$\tilde{v}^m(\tilde{\mathbf{x}}_t) = -(\tilde{\mathbf{x}}_t - \tilde{\mathbf{k}}^m)' \tilde{K}^m (\tilde{\mathbf{x}}_t - \tilde{\mathbf{k}}^m) - \tilde{\kappa}^m. \quad (10)$$

This value function contains both quadratic and linear terms of $\tilde{\mathbf{x}}_t$, with new terms $\tilde{\mathbf{k}}^m \in \mathbb{R}^N$ and $\tilde{\kappa}^m \in \mathbb{R}$. All terms $\tilde{K}^m$, $\tilde{\mathbf{k}}^m$ and $\tilde{\kappa}^m$ are determined endogenously. Each influencer chooses $\{\tilde{\mathbf{r}}_t^m, t \geq 0\}$ to maximize her value function, given the strategies of other influencers.

First, we show that all influencers share the same Ricatti matrix $\tilde{K}^m = \tilde{K}^*$ in the value function (10). This result is easily seen from the two-period example. In the terminal period $T = 1$, all $\tilde{K}_1^m = I$ by assumption. In period $t = 0$, influencer m 's first-order condition is simply $c\tilde{\mathbf{r}}_0^m = -\delta\alpha\tilde{K}_1^m(\tilde{\mathbf{x}}_1 - \tilde{\mathbf{b}}^m)$. Because influencers have the same cost c , impact α , and period-1 Ricatti matrix $\tilde{K}_1^m = I$, their optimal messages have the same slope, and thus their value functions have the same quadratic terms as captured by $\tilde{K}_0^m$. When $T = \infty$, we show that there exists a unique diagonal and positive definite Ricatti matrix $\tilde{K}^*$ solving the following fixed point equation for all influencers:

$$\tilde{K}^* = I + c\delta\tilde{K}^*(cI + \delta\alpha^2\tilde{K}^*)(cI + \delta M\alpha^2\tilde{K}^*)^{-2}D^2, \quad (11)$$

which becomes the Ricatti equation (5) in the single-influencer model when $M = 1$.²³ The following result provides an equilibrium characterization.

Proposition 5. *Suppose A is symmetric and invertible, and influencers' impacts are $\alpha_m = \alpha$. (1) There exists a unique MPE. In this MPE, each influencer's strategy is affine in the opinions:*

$$\tilde{\mathbf{r}}_t^m = -\frac{\delta\alpha}{c}\tilde{K}^* \left(cHD\tilde{\mathbf{x}}_t + \delta H\alpha^2\tilde{K}^* \sum_{l=1}^M \tilde{\mathbf{b}}^l - \tilde{\mathbf{k}}^m \right), \quad (12)$$

where $\tilde{K}^*$ is the Ricatti matrix defined by (11) and $H = (cI + \delta M\alpha^2\tilde{K}^*)^{-1}$.

(2) Agents form consensus in the limit: $\mathbf{x}_\infty = \sum_1^M \mathbf{b}^m / M$.

First, conditional on every other influencer plays a Markov strategy, each influencer's message is unique and affine in current opinions due to the quadratic stage payoff. Moreover, from equation (12), every influencer's message has the same slope because they have the same Ricatti matrix. What sets each influencer apart is the term $\tilde{\mathbf{k}}^m$, which captures the effect of their distinct agendas, leading to different constant terms in each influencer's message. We can think of $\tilde{\mathbf{k}}^m$ as her *dynamic agenda*. As seen from value function (10), each influencer's payoff decreases in the distance between $\tilde{\mathbf{x}}_t$ and $\tilde{\mathbf{k}}^m$, instead of the distance between $\tilde{\mathbf{x}}_t$ and $\tilde{\mathbf{b}}^m$ as in her stage payoff. In equilibrium, $\tilde{\mathbf{k}}^m$ can be expressed as:

$$\left(\tilde{K}^* - \delta DH\tilde{K}^*(\delta\alpha^2\tilde{K}^* + cI) \right) \tilde{\mathbf{k}}^m = \tilde{\mathbf{b}}^m - (\delta DH\tilde{K}^*(\delta\alpha^2\tilde{K}^* + cI))H\delta\alpha^2\tilde{K}^* \sum_1^M \tilde{\mathbf{b}}^m.$$

The projected agenda has a particular form noted in the following observation.

Observation 2. *When A is symmetric, for influencer m , $\tilde{\mathbf{b}}^m = U'b^m\mathbf{1} = (b^m/\sqrt{N}, 0, \dots, 0)'$.*

This observation follows from the fact that the first column of U is $\mathbf{u}_1 = (1/\sqrt{N})\mathbf{1}'$, and all other columns are orthogonal to $\mathbf{u}_1$. By Observation 2 and the formula of $\tilde{\mathbf{k}}^m$, the j th element of $\tilde{\mathbf{k}}^m$ is zero unless $j = 1$, so we focus on the first dimension $\tilde{k}_1^m$ below. Note that the dynamic agenda $\tilde{\mathbf{k}}^m$ grows linearly in influencer m 's agenda $\tilde{\mathbf{b}}^m$, implying that influencer m has the highest dynamic agenda if she has the highest $\tilde{\mathbf{b}}^m$. Moreover, when the average agenda of other influencers is exactly $\tilde{\mathbf{b}}^m$, then $\tilde{\mathbf{k}}^m = \tilde{\mathbf{b}}^m$. Otherwise, the average agenda of the other influencers and $\tilde{\mathbf{k}}^m$ lie on the opposite sides of $\tilde{\mathbf{b}}^m$. For instance, suppose

²³The solution $\tilde{K}^*$ is diagonal because of the independence among dimensions in the transformed problem. Moreover, there exists a unique solution in each dimension by the Intermediate Value Theorem.

$\tilde{\mathbf{b}}^m = \mathbf{0}$, then if the average agenda of the other influencers is positive, then $\tilde{k}_1^m < 0$. As the average agenda of the others increases, $\tilde{k}_1^m$ further decreases. Intuitively, as other influencers push for higher opinions, influencer m needs to pay a higher cost to keep opinions closer to her agenda.

Second, we show that surprisingly, the agents still believe in the average agenda in the limit even though influencers may use asymmetric strategies (due to $\tilde{\mathbf{k}}^m$) in the equilibrium. This result arises from the identical slopes of the optimal strategies, indicating that each influencer's agenda carries equal weight in the limit opinion. It is worth noting that convergence of opinions does not imply convergence of individual influencers' messages to zero. Instead, messages from the influencers counterbalance each other, resulting in the aggregate message converging to zero: $\sum_1^M \tilde{\mathbf{r}}_\infty^m = \mathbf{0}$. That is, to maintain this consensus requires costly continuous intervention by every influencer.

We now examine the strategic interactions among influencers. In every period, the best response of influencer m to the others' messages is:

$$(cI + \delta\alpha^2 \tilde{K}^*) \tilde{\mathbf{r}}_{t-1}^m = -\delta\alpha \tilde{K}^* \left(D\tilde{\mathbf{x}}_{t-1} + \alpha \sum_{l \neq m} \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}^m \right). \quad (13)$$

Clearly, each influencer's message is decreasing in the sum of all other influencers' messages. Because each message may be positive or negative, we study how extreme each influencer's message is (its absolute value) in response to another influencer's message.²⁴ If both messages have the same sign, their messages are strategic substitutes in that if one's message becomes more extreme, the others becomes more moderate. Conversely, if the messages have opposite signs, then the messages become strategic complements: if one becomes more extreme, so is the other.

When all influencers' agendas are identical and normalized to 0,²⁵ according to (13), their messages are strategic substitutes. Intuitively, persuading the agents is a public good when they share the same agenda. How significant is this strategic substitute effect? One way is to study how each influencer's payoff changes with M , the total number of influencers. Another way is to compare their payoffs with that of a representative influencer who chooses M messages in each period to maximize the sum of all influencers' payoffs.

²⁴We use the terms strategic substitutes and complements intuitively to talk about how influencers' messages in each period respond to each other even though they are often used in static games.

²⁵If all influencers have agenda $b^m = b$, then we can have a change of variable $\mathbf{z}_t = \mathbf{x}_t - b\mathbf{1}$, such that $\mathbf{z}_t = A\mathbf{z}_{t-1} + \alpha \sum_1^M \mathbf{r}_t^m$. That is, it is without loss to assume that all influencers have agenda 0.

Corollary 1. *Suppose A is symmetric, $\alpha_m = \alpha$ and $b^m = 0$ for all m .*

(1) *All influencers get the same payoff which increases in M . Moreover, every influencer sends less extreme messages and the opinions converge faster as M increases.*

(2) *Compared to a representative influencer, the case with multiple influencers exhibits a lower total payoff and slower opinion convergence.*

Despite the free-riding effect, every influencer is better off as the number of influencers M increases. The payoff gain comes from two sources: lower message cost and faster opinion convergence. First, everyone pushes less against the current opinions as M increases:

$$\tilde{\mathbf{r}}_t^m = -(\delta M \tilde{K}^*(M) + cI)^{-1} \delta \tilde{K}^*(M) D \tilde{\mathbf{x}}_t. \quad (14)$$

The absolute value of the slope in (14) decreases in M , because we show in the proof that $\tilde{K}^*(M)$ decreases in M while $M \tilde{K}^*(M)$ increases in M . So fix $\tilde{\mathbf{x}}_t$, influencer m chooses a less extreme message $\tilde{r}_t^m$ in period t and pays a lower cost. In addition, the opinions become closer to their agenda 0, period by period, as M increases. Intuitively, while each influencer pushes less, since there are more of them, the aggregate message is more extreme and opinions converge faster to 0. The free-riding effect among influencers, however, implies that the messages are too moderate and payoffs are too low comparing to those of a representative influencer.

What if the influencers have different agendas? While the slope of the strategy is still identical because all influencers have the same $\tilde{K}^*$, the dynamic agenda $\tilde{\mathbf{k}}^m$ varies. By Observation 2, we focus on the first dimension, the only dimension in which influencers are asymmetric. Substituting the limit opinion into equation (12), we have influencer m 's message in the first dimension,

$$\tilde{r}_{1,\infty}^m = \frac{\delta \alpha}{c} \tilde{K}_{11}^* \left(\tilde{b}_1^m - \sum_1^M \tilde{b}_1^m / M \right).$$

In the limit, the messages of influencers with agendas above and below the average agenda are strategic complements, and messages of influencers with agendas all above (or all below) the average agenda are strategic substitutes. Suppose the government is influencer 1 with agenda 0. If another influencer 2's agenda becomes more negative (vaccines more dangerous), then we can see in the limit, the government has to send a more positive message; influencer 2 has to send a more negative message, and the consensus is more negative. In

addition, payoffs are higher for influencers whose agendas are closer to the average agenda. For instance, a government trying to convince a network of agents about the efficacy of vaccines is unlikely to succeed despite paying a high cost if other influencers have extreme agendas biased in the other direction. In contrast, the government will do better if the network faces a more diverse and well-balanced group of influencers.

Lastly, the spirit of Proposition 5 extends to the case of any asymmetric network A . We focus on a *semi-symmetric equilibrium* in which all influencers use affine strategies with the same slope but possibly different constant terms, because the slope is determined by the same Ricatti matrix K^* while $\mathbf{k}^m$ differs due to their different agendas. We show that there exists a semi-symmetric MPE with the same consensus among agents as before.

Proposition 6. *Suppose A is invertible and $\alpha_m = \alpha$ for all m . If there exists a positive definite solution to the Ricatti equation*

$$K^* = I + c\delta A' H' K^* (\delta \alpha^2 K^* + cI) H A, \quad (15)$$

where $H = (cI + \delta M \alpha^2 K^*)^{-1}$. Then there exists a semi-symmetric MPE in which all influencers have the same Ricatti matrix K^* . In this MPE, agents form consensus in the limit: $\mathbf{x}_\infty = \sum_1^M \mathbf{b}^m / M$.

When the network is asymmetric, the agents' weights on their neighbors' opinions are not uniform; for instance, some agents are listened to with higher weights than others. But the same impact α means that all influencers can affect each agent in the same way. In this equilibrium (if it exists), they all send more extreme messages to more central agents and less extreme messages to less central agents. In the limit, messages eventually offset each other and the consensus remains the average agenda. Furthermore, this MPE is the unique semi-symmetric MPE if matrix $K^* - (K^* - I)A^{-1}$ is non-singular.²⁶ Because K^* is the solution of a nonlinear fixed point equation of matrices, we cannot directly show such a positive definite solution exists. But for any finite T , there exists a unique semi-symmetric MPE with similar properties.

²⁶We do not rule out the possibility of a MPE in which influencers use strategies that differ both in slope and in constant term, that is, each influencer has her own K^m and $\mathbf{k}^m$, similar to the results in Proposition 7. The existence of solutions to the coupled Ricatti equations is still an open question for a general network A , so we focus on the semi-symmetric equilibrium.

4.2 Influencers with Different Impacts

As before, we begin with symmetric network A and decompose it into $A = UDU'$. In the transformed problem, the opinions evolve according to (9) and influencer m 's value function is given by (10). Notice that distinct values of α_m indicate different marginal impacts of influencers on agents' opinions, resulting in distinct Ricatti matrices denoted by $\tilde{K}^m$. To gain tractability, we focus on the case of two asymmetric influencers with different α_m and leave the general analysis to Appendix B. Without loss, let $\alpha_1 > \alpha_2$, meaning that every message of influencer 1 has a bigger impact on the next period's opinion. The next result illustrates that consensus is still sustained when the network is symmetric.

Proposition 7. *Suppose A is symmetric and $\alpha_1 > \alpha_2$. (1) The Ricatti matrices satisfy $\tilde{K}^1 > \tilde{K}^2$. In the unique MPE, the slope of influencer 1's optimal strategy is steeper than that of influencer 2. In the limit, influencer 1 intervenes no more than influencer 2: $|\tilde{\mathbf{r}}_\infty^1| \leq |\tilde{\mathbf{r}}_\infty^2|$ with inequality holds in the first dimension when $b^1 \neq b^2$.*

(2) Agents form consensus in the limit, which is a weighted average of the influencers' agendas and is closer to b^1 than to b^2 .

Recall that the optimal message for influencer m is affine in agents' opinions, taking the form $\tilde{\mathbf{r}}_t^m = \tilde{L}^m \tilde{\mathbf{x}}_t + \tilde{\mathbf{I}}^m$ where $\tilde{L}^m$ is a diagonal matrix. Unlike Proposition 5, the slope differs between the influencers, and part (1) shows that $|\tilde{L}_{jj}^1| > |\tilde{L}_{jj}^2|$ for each dimension j . Intuitively, as influencer 1 has a bigger impact on the agents' opinions, her optimal message features a steeper slope in every dimension which fastens convergence. Next, part (1) compares the influencers' limit messages, which depends on whether their agendas are the same or not. If $b^1 = b^2$, then their messages converge to zero as the agents' opinions converge to their agenda. Otherwise, the stronger influencer 1 sends a less extreme message than the weaker influencer 2 in the limit. Notice that the messages offset each other when there is consensus (shown in part 2), which means that $\alpha_1 \tilde{\mathbf{r}}_t^1 + \alpha_2 \tilde{\mathbf{r}}_t^2 \rightarrow 0$ as $t \rightarrow \infty$. Because $\alpha_1 > \alpha_2$, we have $|\tilde{\mathbf{r}}_t^1| < |\tilde{\mathbf{r}}_t^2|$ dimension by dimension in the limit when they are not both converging to zero.²⁷ In terms of payoffs, when $b^1 = b^2$, influencer 1 sends more extreme messages and gets a lower payoff than influencer 2 because of the free-riding effect discussed in Corollary 1. When $b^1 \neq b^2$, the payoff comparison is ambiguous. Influencer 1's early

²⁷This result also shows the property of the constant term $\tilde{\mathbf{I}}^m$ in influencer m 's optimal message. Recall that the slope of influencer 1's optimal message is steeper, so the effect from the constant term must dominate the slope effect so that influencer 1 sends a less extreme message in the limit than influencer 2.

messages may be more extreme and costly due to steeper slopes, but her later messages are less extreme due the property of the limit messages.

Part (2) of Proposition 7 shows in the limit, agents forms consensus, which is closer to b^1 than to b^2 . In particular, the weight on each agenda is a nonlinear and endogenous function of their impacts and the Ricatti equations. The first dimension of the limit projected opinion $\tilde{\mathbf{x}}_\infty^1$, the only non-zero dimension by Observation 2, is proportional to

$$\alpha_1^2 \left(1 - \frac{\delta(\delta\alpha_2^2 \tilde{K}_{22}^2 + c)}{c + \delta\alpha_1^2 \tilde{K}_{11}^2 + \delta\alpha_2^2 \tilde{K}_{11}^2} \right) \tilde{b}^1 + \alpha_2^2 \left(1 - \frac{\delta(\delta\alpha_1^2 \tilde{K}_{11}^1 + c)}{c + \delta\alpha_1^2 \tilde{K}_{11}^2 + \delta\alpha_2^2 \tilde{K}_{11}^2} \right) \tilde{b}^2,$$

where $\tilde{K}_{11}^m$ is the first diagonal entry of $\tilde{K}^m$.²⁸ The denominator $c + \delta\alpha_1^2 \tilde{K}_{11}^2 + \delta\alpha_2^2 \tilde{K}_{11}^2$ reflects the marginal effect of the two messages, both the cost today and the total change in future payoffs. Since $\alpha_1 > \alpha_2$, and $\tilde{K}_{11}^1 > \tilde{K}_{11}^2$, the limit opinion puts a higher weight on influencer 1's agenda. But why do agents reach a consensus when the influencers have different agendas? Hypothetically, it is possible for each influencer to focus on different agents, such as influencer 1 sends stronger messages to one group of agents and influencer 2 sends stronger messages to the rest. But in a symmetric network A , each agent's eigenvector centrality is the same, that is, they have identical weights in affecting the limit opinions. Thus, influencers do not find it optimal to treat agents differently, and they form consensus. In the next section, we will show that consensus is unlikely when the network is asymmetric.

5 Two Applications with Perpetual Disagreements

In real-life social networks, disagreements among agents are prevalent rather than rare occurrences. We apply our model to shed light on the reasons behind agents' failure to reach consensus. Essentially, their disagreements stem from the differential targeting employed by competing influencers, either as a consequence of strategic choices or due to the influence of spam bots sending constant messages.

5.1 Uniform Opinion Leader Network

Suppose the government in our running example is influencer 1, who knows the presence of an adversarial influencer pushing the agenda b^2 that vaccines are dangerous and should not

²⁸The projected limit opinions are $\tilde{\mathbf{x}}_\infty = \{\tilde{x}_\infty^1, 0, \dots, 0\}$, and the unprojected limit opinions have identical entries: $\mathbf{x}_\infty = U' \tilde{\mathbf{x}}_\infty = \tilde{\mathbf{x}}_\infty^1 \mathbf{1}$.

be used. In the network, there are opinion leaders such as medical officials, news reporters, and researchers, whose opinions about vaccination hold greater significance to the general population. To capture this crucial aspect of real-life networks while keeping the analysis tractable, we focus on the following type of stylized networks.

Definition 1. *In a uniform opinion leader network, each column j has identical entries $A_{ij} = a_j$ for all $i \in \mathcal{N}$, and $a_1 > a_2 > \dots > a_N$.*

Agent 1 is the clear opinion leader in this network as all agents put the highest weight on his opinions. Any network with opinion leaders is asymmetric, and thus we apply our general network model in Section 3.3. Moreover, we use the singular value decomposition (SVD) method to transform the problem into one with independent dimensions, which simplifies the analysis in networks with low rank such as the uniform opinion leader network.

Specifically, let $A = USV'$, where $S = \text{diag}\{\sigma_1, \dots, \sigma_N\}$ is the singular value matrix and each σ_j^2 is an eigenvalue of $A'A$.²⁹ Let $\sigma_j^2 \geq \sigma_{j+1}^2$ for all $1 \leq j < N$. The columns of U and V are orthonormal eigenvectors of AA' and $A'A$ respectively: $U'U = V'V = I$. Each column of U is denoted as $\mathbf{u}_j$ and each column of V is denoted as $\mathbf{v}_j$. Because $A = \sum_j \sigma_j \mathbf{u}_j \mathbf{v}_j'$ and σ_j is ordered by magnitude, dimensions associated with larger singular values are more important for the influencer. We transform the problem by using projected variables like before. Let $\tilde{\mathbf{x}}_t = V'\mathbf{x}_t$ and $\tilde{\mathbf{r}}_t = V'\mathbf{r}_t$. In the transformed problem, we have $\tilde{\mathbf{x}}_{t+1} = V'US\tilde{\mathbf{x}}_t + \sum \alpha_m \tilde{\mathbf{r}}_t^m$.

SVD greatly simplifies the analysis of the uniform opinion leader network, a rank one matrix, as it transforms the influencers' problems in all but the first dimension into myopic problems. This result is due to the feature of the singular values in the uniform opinion leader network: $\sigma_1 = \sqrt{N(a_1^2 + \dots + a_N^2)} > 1$ and $\sigma_j = 0$ for all $j \neq 1$. As the opinions updating matrix is $A = \sum_j \sigma_j \mathbf{u}_j \mathbf{v}_j'$, when $\sigma_j = 0$, dimension- j opinion in period $t + 1$ does not have any effect on the future opinions. Thus, when choosing the optimal dimension- j message in period t , the influencer considers only her period- $(t + 1)$ payoff. This myopic problem in each dimension $j \neq 1$ simplifies the Ricatti matrix, which has no closed-form solutions when the network is asymmetric in general.

Observation 3. *In a uniform opinion leader network, the Ricatti matrix $\tilde{K}^m$ is diagonal with $\tilde{K}_{11}^m > 1$ and $\tilde{K}_{jj}^m = 1$ for all $j \neq 1$.*

²⁹SVD generalizes the eigendecomposition used in Section 3.1 when A is symmetric. Because $A'A$ is symmetric and positive semi-definite for any A , it has real, non-negative eigenvalues.

This result is a direct implication of the myopic problem in dimension $j \neq 1$. Recall that the stage payoff depends on the opinions in the form of $-(\tilde{\mathbf{x}}_t - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_t - \tilde{\mathbf{b}}^m)$. When opinions have no future effect on the influencer's payoff other than the stage payoff in the next period, the quadratic term of the opinions is weighted by 1, which is the j th diagonal entry of $\tilde{K}^m$.

Next, we show the government and the adversarial influencer 2 *endogenously* choose to persuade different agents with different intensity, leading to perpetual disagreement.

Proposition 8. *Suppose A is an uniform opinion leader network and $\alpha_1 > \alpha_2$. In the unique MPE, the agents' opinions converge in the limit and every agent's limit opinion is closer to b^1 than to b^2 . When $b^1 \neq b^2$, they disagree in the limit and the opinion leaders' limit opinions are closer to b^1 than those of opinions followers: $|x_\infty^1 - b^1| < \dots < |x_\infty^N - b^1|$.*

From Observation 3, the influencers' messages in the first dimension are chosen to maximize their total discounted payoffs, but in all other dimensions are myopic best response.³⁰ Several parts of Proposition 7 still hold in uniform opinion leader networks. Specifically, $\tilde{K}_{11}^1 > \tilde{K}_{11}^2$ and the slope of influencer 1's optimal strategy is steeper than that of influencer 2. Also, the agents' limit opinions are closer to influencer 1's agenda than that of influencer 2. More interestingly, whenever $b^1 \neq b^2$, the two influencers focus on different subsets of agents, who then do not reach consensus in the limit. This differential targeting is reflected in the limit aggregate message $\alpha_1 \mathbf{r}_\infty^1 + \alpha_2 \mathbf{r}_\infty^2$. Suppose $b^1 > b^2$, the limit aggregate message is a decreasing vector with the first entry being the highest. There exists a threshold agent k , such that the limit aggregate message $\alpha_1 \mathbf{r}_\infty^1 + \alpha_2 \mathbf{r}_\infty^2$ is positive in its first k entries and non-positive in all other entries. It implies that influencer 1 with a higher impact α_1 focuses on the first k agents, who are the (relative) opinion leaders. Finding it not optimal to compete with influencer 1 on these leaders, influencer 2 focuses more on agents $k + 1$ through N who are the opinion followers. Therefore, the agents disagree permanently, with the leaders' opinions closer to b^1 than those of the followers.

While the uniform opinion leader network is stylized, the decomposition method used here can be extended to cases when A has a low rank, or when A has many small singular values. In general networks with asymmetric influencers, analytical results beyond the necessary conditions provided in Appendix B are harder to obtain. We use a three-agent asymmetric

³⁰Unlike Observation 2 in the symmetric A case, $\tilde{\mathbf{b}}^m = V' \mathbf{b}^m$ can have many dimensions with none zero and different entries. As a result, in all dimension $j \neq 1$, the optimal strategy to maximize next-period payoff differs in both the slope due to different α_m and the constant term.

network to show the essence of differential targeting remains, leading to perpetual disagreement. Let the initial opinions be $\mathbf{x}_0 = (7, -2, 5)'$. Let $c = 0.1$ and $\delta = 0.9$. Two influencers have opposing agendas $b^1 = 10$ and $b^2 = -10$, and different levels of influence: $\alpha_1 = 0.6$ and $\alpha_2 = 0.5$. The network is represented by

$$A = \begin{bmatrix} 0.6 & 0.3 & 0.1 \\ 0.4 & 0.1 & 0.5 \\ 0.5 & 0.2 & 0.3 \end{bmatrix}.$$

Agent 1 is the opinion leader of the network because everyone listens to him with a relatively high weight. Agent 2 is the least important one receiving a total weight of 0.6. Figure 2 depicts the opinion evolution in this network.

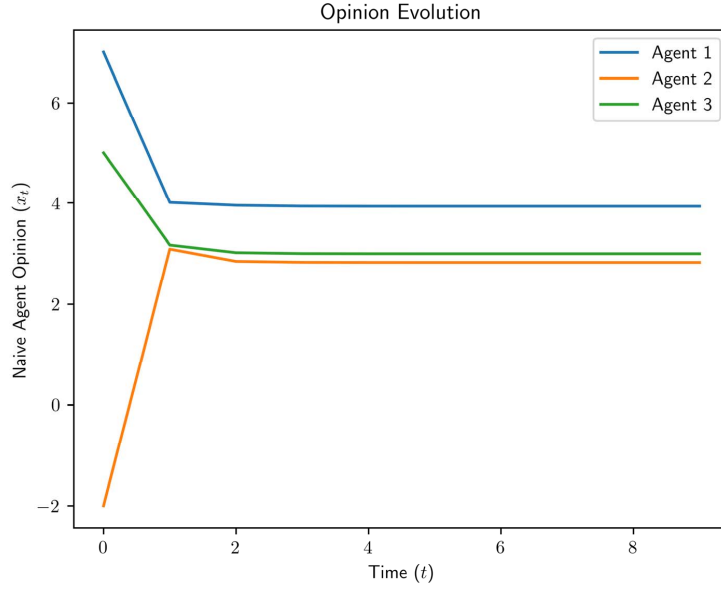


Figure 2: Opinions over time in an asymmetric network under asymmetric influencers.

Clearly, the agents fail to reach consensus. The limit opinion vector is $(3.93, 2.83, 3)$, with every agent's limit opinion being closer to the agenda of influencer 1. The disagreement is due to the fact that the influencers focus on different agents, such that the limit aggregate message $\alpha_1 \mathbf{r}_\infty^1 + \alpha_2 \mathbf{r}_\infty^2$ is $(0.42, -0.53, -0.43)$. In net, influencer 1 pushes harder on agent 1 and influencer 2 pushes harder on the other two agents. Notably, as agent 1 is the opinion leader, both influencers send the most extreme messages to him compared to messages to the other agents. Influencer 1 with the higher impact benefits from focusing on agent 1 and

keep his opinions closer to her agenda $b^1 = 10$. Then, influencer 2 focuses more on the other two agents to keep their opinions relatively closer to agenda $b^2 = -10$. In the limit, influencer 1 sends a less costly message and gets a higher payoff than influencer 2.

5.2 Confronting Bots with Constant Messages

We can also apply the general network model in Section 3.3 and Proposition 4 in particular to study how the government in our running example optimally counters spam bots who mimic genuine network users and send repetitive messages at little to no cost. For instance, [Vosoughi, Roy, and Aral \(2018\)](#) report that in 2017, Twitter, Facebook, and Instagram were home to 23, 140 and 27 million bots respectively, representing around 8.5%, 5.5%, and 8.2% of all accounts on these platforms.³¹ In addition, [Azzimonti and Fernandes \(2023\)](#) show that significant levels of polarization are possible even though only 15% of agents believe in fake news from the bots. Using our model, we show that the government chooses not to achieve her agenda when facing bots, that is, it is optimal for her to counter different messages of the bots differently, resulting in generic disagreement and polarization.³²

Suppose there are M bots.³³ Each bot m sends a message $\mathbf{z}^m$ in each period, and the vector of their collective messages is $\mathbf{z} \in \mathbb{R}^M$. The influencer sends messages $\mathbf{r}_t \in \mathbb{R}^N$, and the agents update their opinions in each period after listening to the bots and the influencer:

$$\mathbf{x}_{t+1} = A^n \mathbf{x}_t + A^z \mathbf{z} + \mathbf{r}_t,$$

where A^n and A^z represent the weights with which the agents' listen to each other's opinions and to the bots respectively. Each agent's total weight on the agents and the bots is 1, that is,

³¹According to the 2021 research report titled "Bot Attacks: Top Threats and Trends" from security firm Barracuda, more than two-thirds of internet traffic is generated by bots. In addition, 67% of bad bot traffic originates from public data centers in North America.

³²The mechanism underlying our polarization result differs from the existing literature. It is well known that long-run disagreement can arise due to stubborn agents, homophily, and a lack of interaction among some clusters of agents. [Acemoglu, Como, Fagnani, and Ozdaglar \(2013\)](#) find disagreements in a network with stubborn agents and regular agents who update opinions by exchanging information with a neighbor chosen randomly at any given time. [Della Lena \(2022\)](#) shows that when agents receive informative signals and listen to sources with fixed opinions, the limit opinion does not feature consensus.

³³In the main model, we use M to refer to the number of strategic influencers. In this section with a single influencer, we abuse the notation to let M be the number of bots and m be a generic bot.

$\sum_j a_{ij}^n + \sum_m a_{im}^z = 1$ for all i , and all weights are non-negative.³⁴ We further assume that A^n is strictly substochastic (i.e., each row sum is strictly less than 1) and diagonalizable, and thus $(I - A^n)^{-1}$ exists. We can rewrite the problem which then fits nicely into our general network model as follows:

$$\chi_{t+1} = \begin{bmatrix} \mathbf{x}_{t+1} \\ \mathbf{z} \end{bmatrix} = \begin{bmatrix} A^n & A^z \\ \mathbf{0} & I \end{bmatrix} \begin{bmatrix} \mathbf{x}_t \\ \mathbf{z} \end{bmatrix} + \begin{bmatrix} I \\ \mathbf{0} \end{bmatrix} \mathbf{r}_t = A\chi_t + W\mathbf{r}_t.$$

Observe that the new opinion variable is χ_t , an $(N + M) \times 1$ vector of the agents' opinions and the bots' agendas. The new opinion variable evolves according to an expanded matrix A , which is row stochastic by definition, accounting for the fact that bots never change their messages. Also, the influencer's messages $\mathbf{r}_t$ are multiplied by the $(N + M) \times N$ matrix W because the influencer never sends costly messages to the bots. The government's total discounted payoff can be expressed as:

$$- \sum_{t=0}^{\infty} \delta^t (\chi_t' \chi_t + c \mathbf{r}_t' \mathbf{r}_t - \mathbf{z}' \mathbf{z}).$$

As she cares only about persuading the agents, we remove the bots' impact on her payoff by subtracting $\mathbf{z}' \mathbf{z}$ from $\chi_t' \chi_t$. The next result shows that while the government pushes back against the bots in every period, she does not fully achieve her agenda.

Corollary 2. *The government's optimal message is $\mathbf{r}_t = L^* \chi_t$, in which $L^* = -(\delta W' K^* W + cI)^{-1} \delta W' K^* A$ and K^* is the unique positive definite solution of*

$$K^* = I + \delta A' (K^* - \delta K^* W (\delta W' K^* W + cI)^{-1} W' K^*) A.$$

Opinions and messages converge: $\lim_{t \rightarrow \infty} \mathbf{x}_t = \mathbf{x}_\infty$ and $\mathbf{r}_\infty = \lim_{t \rightarrow \infty} \mathbf{r}_t$. The agents disagree with each other and the government generically: $\mathbf{x}_\infty = (I - A^n)^{-1} (A^z \mathbf{z} + \mathbf{r}_\infty)$.

The limit opinions are a weighted average of the bots' messages and the government's limit message, where the weights measure their cumulative influences. In particular, the influence from the agents' initial opinions is equal to $\lim_{t \rightarrow \infty} (A^n)^t = \mathbf{0}$ because A^n is substochastic. The cumulative weight of the government's message is $I + A^n + \dots + (A^n)^\infty = (I - A^n)^{-1}$.³⁵

³⁴We modify our assumption to $A^n + A^z$, instead of A , being row stochastic to ensure the limit opinions are well defined. To highlight the change, we denote the network matrix by A^n .

³⁵This weight includes direct and indirect influence. For example, after one period, the agents place weight $(I + A^n)$ on the influencer's message $\mathbf{r}_\infty$ (walk of length 1 and 2 from the influencer to her), and so on. This measure is closely related to the centrality of [Katz \(1953\)](#).

Similarly, the agents put a total weight of $(I - A^n)^{-1}A^z$ on the bots. With bots, the government sends non-vanishing messages in the limit ($\mathbf{r}_\infty \neq \mathbf{0}$), and the agents disagree with her generically ($\mathbf{x}_\infty \neq \mathbf{0}$). To understand this claim, for the first part, suppose $\mathbf{r}_\infty = \mathbf{0}$, then from Corollary 2, the limit opinion is a linear function of the bots' agendas and is not zero in general, $\mathbf{x}_\infty = (I - A^n)^{-1}A^z\mathbf{z}$. The influencer should push the opinions by sending non-zero messages according to the optimal strategy, and thus $\mathbf{r}_\infty \neq \mathbf{0}$. Second, suppose $\mathbf{x}_\infty = \mathbf{0}$, then from the corollary, the limit message need to completely offset the bots' influence: $\mathbf{r}_\infty = -A^z\mathbf{z}$, but this is not the optimal message.

It is more subtle to show that the agents do not have consensus in general. Since $\mathbf{r}_\infty$ is endogenous, we need to have a direct solution of $\mathbf{x}_\infty$. Because A is upper triangular and $W = (I, 0)'$, K^* has an interesting feature: it is a block matrix of the form³⁶

$$K^* = \begin{bmatrix} K_1 & K_2 \\ K_2' & K_3 \end{bmatrix}.$$

Here, K_1 is a function of itself which depends only on the model parameters and network structure A^n , while K_2 depends on both K_1 and the bots' access to the network A^z (see equation (50) and (51) in the proof of Corollary 2). The limit opinion becomes:

$$\mathbf{x}_\infty = (\delta K_1 + cI - cA^n)^{-1}(cA^z - \delta K_2)\mathbf{z}.$$

We can express the limit opinion as the product of a positive definite matrix independent of $\mathbf{z}$ and $A^z\mathbf{z}$. For any positive definite matrix, there can only be one $\mathbf{z}$ such that the agents reach consensus. Therefore, agents don't agree with each other generically and the influencer is worse off as $\|\mathbf{z}\|$ increases.

6 Conclusion

Our model is portable and can be used to study how strategic influencers competing over networks, such as elections, marketing, and adversarial nations' disinformation campaigns. By decomposing the model appropriately, we obtain tractable results, including comparative statics and insights into limit opinions. Our model can investigate network intervention

³⁶Submatrix K_1 summarizes the future disutility from agents' opinion deviation from the influencer, and K_2 summarizes her future disutility because the agents are influenced by the bots' agendas. Submatrix K_3 summarizes the disutility due to the bots' deviation from her agenda, which does not matter to her strategy.

questions, such as determining which social network platforms (e.g., Facebook, Twitter, Instagram) influencers should target.³⁷ Timing, for instance, plays a crucial role: our model suggests that early intervention in a network is key. Once a network accumulates numerous influencers with extreme agendas, it becomes increasingly difficult to persuade others.

One promising research direction is a model of coarse targeting, where influencers can send messages to only some agents in the network. Preliminary findings suggest that strategic influencers face a trade-off between short-term and long-term payoffs in their target selection. Convergence speed plays a vital role in long-term outcomes, making opinion leaders and stubborn agents potential candidates for targeting. Moreover, it is essential to target agents with extreme initial opinions, as their opinions affect short-term payoffs significantly. The question of optimal dynamic targeting holds relevance and presents theoretical interest.

Our findings suggest that in asymmetric networks with asymmetric influencers, disagreement in the limit is the norm rather than the exception. We consider a general model that accommodates these asymmetries and provide the influencers’ optimal strategies in Appendix B, which are necessary conditions for possible MPE. Relatedly, an active area of research in computer science involves investigating the existence and properties of equilibria in general-sum multiple-player linear quadratic (LQ) games.³⁸ In addition, researchers are exploring when standard tools such as the policy gradient algorithm lead to an equilibrium (Mazumdar, Ratliff, Jordan, and Sastry 2019).

References

- ACEMOGLU, D., G. COMO, F. FAGNANI, AND A. OZDAGLAR (2013): “Opinion Fluctuations and Disagreement in Social Networks,” *Mathematics of Operations Research*, 38(1), 1–27.
- ALATAS, V., A. BANERJEE, A. G. CHANDRASEKHAR, R. HANNA, AND B. A. OLKEN (2016): “Network Structure and the Aggregation of Information: Theory and Evidence from Indonesia,” *American Economic Review*, 106(7), 1663–1704.

³⁷A 2019 Oxford report shows that Facebook is the No. 1 platform for governments and political parties seeking to manipulate public opinion. See “The Global Disinformation Order: 2019 Global Inventory of Organised Social Media Manipulation” by Samantha Bradshaw and Philip N. Howard, University of Oxford.

³⁸While sufficiency results exist for zero-sum LQ games and two-player LQ games, conditions for the existence of Ricatti matrices for the infinite Ricatti equations are open questions.

- ALEVY, J. E., M. S. HAIGH, AND J. A. LIST (2007): “Information Cascades: Evidence from a Field Experiment with Financial Market Professionals,” *Journal of Finance*, 62, 151–180.
- ANDERSON, B. D., AND J. B. MOORE (1989): *Optimal Control: Linear Quadratic Methods*. Prentice-Hall International, Inc.
- ANDERSON, L. R., AND C. A. HOLT (1997): “Information Cascades in the Laboratory,” *American Economic Review*, 87, 847–862.
- AZZIMONTI, M., AND M. FERNANDES (2023): “Social media networks, fake news, and polarization,” *European Journal of Political Economy*, 76.
- BAI, J., M. GOLOSOV, N. QIAN, AND Y. KAI (2015): “Understanding the Influence of Government Controlled Media: Evidence from Air Pollution in China,” *working paper*.
- BALA, V., AND S. GOYAL (1998): “Learning from Neighbours,” *The Review of Economic Studies*, 65(3), 595–621.
- BALLESTER, C., A. CALVO-ARMENGOL, AND Y. ZENOU (2006): “Who’s Who in Networks. Wanted: The Key Player,” *econometrica*, 74(5), 1403–1417.
- BERTSEKAS, D. P. (2017): *Dynamic programming and optimal control*, vol. 1. Athena Scientific, 4th edition edn.
- BLOCH, F., AND S. SHABAYEK (2023): “Targeting in social networks with anonymized information,” *Games and Economic Behavior*.
- CAI, H., Y. CHEN, AND H. FANG (2009): “Observational Learning: Evidence from a Randomized Natural Field Experiment,” *American Economic Review*, 99(3), 864–882.
- CELEN, B., AND S. KARIV (2004): “Distinguishing Informational Cascades from Herd Behavior in the Laborator,” *American Economic Review*, 94, 484–498.
- CHANDRASEKHAR, A. G., H. LARREGUY, AND J. P. XANDRI (2020): “Testing Models of Social Learning on Networks: Evidence from Two Experiments,” *Econometrica*, 88(1), 1–32.
- DEGROOT, M. H. (1974): “Reaching a Consensus,” *Journal of the American Statistical Association*, 69(345), 118–121.

- DELLA LENA, S. (2022): “The Spread of Misinformation in Networks with Individual and Social Learning,” Available at: <http://dx.doi.org/10.2139/ssrn.3511080>.
- DEMARZO, P. M., D. VAYANOS, AND J. ZWIEBEL (2003): “Persuasion Bias, Social Influence, and Uni-Dimensional Opinions,” *Quarterly Journal of Economics*, 118(3), 909–968.
- ENGWERDA, J. (1998): “On Scalar Feedback Nash Equilibria in the Infinite Horizon LQ-Game,” in *IFAC Proceedings Volumes*, vol. 31, pp. 193–198.
- ENKE, B., AND F. ZIMMERMANN (2019): “Correlation neglect in belief formation,” *The Review of Economic Studies*, 86(1), 313–332.
- FREILING, G., G. JANK, AND H. ABOU-KANDIL (1996): “On global existence of solutions to coupled matrix Riccati equations in closed-loop Nash games,” *IEEE Transactions on Automatic Control*, 41(2), 264–269.
- FRIEDKIN, N. E., AND E. C. JOHNSEN (1990): “Social Influence and opinions,” *Journal of Mathematical Sociology*, 15, 193–206.
- GALE, D., AND S. KARIV (2003): “Bayesian learning in social networks,” *Games and Economic Behavior*, 45(2), 329–346.
- GALEOTTI, A., B. GOLUB, AND S. GOYAL (2020): “Targeting interventions in networks,” *Econometrica*, 88(6), 2445–2471.
- GALEOTTI, A., AND S. GOYAL (2009): “Influencing the Influencers: A Theory of Strategic Diffusion,” *The Rand Journal of Economics*, 40(3), 509–532.
- GHADERI, J., AND R. SRIKANT (2014): “Opinion dynamics in social networks with stubborn agents: Equilibrium and convergence rate,” *Automatica*, 50(12), 3209–3215.
- GOLUB, B., AND M. O. JACKSON (2010): “Naive Learning in Social Networks and the Wisdom of Crowds,” *American Economic Journal: Microeconomics*, 2(1), 112–49.
- GRABISCH, M., A. MANDEL, A. RUSINOWSKA, AND E. TANIMURA (2018): “Strategic Influence in Social Networks,” *Mathematics of Operations Research*, 43(11111), 29–50.
- GRIMM, V., AND F. MENGEL (2020): “Experiments on belief formation in networks,” *Journal of the European Economic Association*, 18(1), 49–82.

- HANSEN, L. P., AND T. J. SARGENT (2013): *Recursive Models of Dynamic Linear Economies*. Princeton University Press.
- HESPANHA, J. P. (2018): *Linear systems theory*. Princeton university press.
- JACKSON, M. O. (2008): *Social and Economic Networks*. Princeton University Press.
- JEONG, D., AND E. SHIN (2022): “Optimal Influence Design in Networks,” KAIST College of Business Working Paper Series No. 2022-001.
- KATZ, L. (1953): “A new status index derived from sociometric analysis,” *Psychometrika*, 18, 39–43.
- LEVY, G., AND R. RAZIN (2018): “Information diffusion in networks with the Bayesian Peer Influence heuristic,” *Games and Economic Behavior*, 109, 262–270.
- LI, W., AND X. TAN (2020): “Locally Bayesian Learning in Networks,” *Theoretical Economics*, 15, 239–278.
- (2021): “Cognitively-Constrained Learning from Neighbors,” *Games and Economic Behavior*, 129, 32–54.
- MAZUMDAR, E., L. J. RATLIFF, M. I. JORDAN, AND S. S. SASTRY (2019): “Policy-Gradient Algorithms Have No Guarantees of Convergence in Linear Quadratic Games,” <https://arxiv.org/abs/1907.03712>.
- MOBIUS, M., T. PHAN, AND A. SZEIDL (2015): “Treasure Hunt: Social Learning in the Field,” *NBER Working Paper No. 21014*.
- MOLAVI, P., A. TAHBAZ-SALEHI, AND A. JADBABAIE (2018): “A Theory of Non-Bayesian Social Learning,” *Econometrica*, 86(2), 445–490.
- MOSSEL, E., A. SLY, AND O. TAMUZ (2015): “Strategic Learning and the Topology of Social Networks,” *Econometrica*, 83(5), 1755–1794.
- MUELLER-FRANK, M., AND C. NERI (2019): “A General Analysis of Boundedly Rational Learning in Social Networks,” working paper, SSRN 2933411.
- PAPAVASSILOPOULOS, G. P., AND G. J. OLSDER (1984): “On the linear-quadratic, closed-loop, no-memory Nash game,” *Journal of optimization theory and applications*, 42(4), 551–560.

SADLER, E. (forthcoming): “Influence Campaigns,” *American Economic Journal: Microeconomics*.

VOHRA, A. (2023): “Strategic Influencers and the Shaping of Beliefs,” mimeo.

VOSOUGHI, S., D. ROY, AND S. ARAL (2018): “The spread of true and false news online,” *Science*, 359(6380), 1146–1151.

WATTS, A. (2001): “A Dynamic Model of Network Formation,” *Games and Economic Behavior*, 34(2), 331–341.

A Proofs

Proof of Proposition 1: We characterize both the finite-horizon and the infinite-horizon optimal messages and payoffs in details in this proposition. In later results, we focus on the infinite-horizon case, but we may use the finite-horizon results as part of the proof. First, we present the finite-horizon result formally.

Proposition 1 (cont’d). *Suppose A is symmetric and is decomposed as $A = UDU'$. If T is finite, then for all $t < T$, the influencer’s optimal message is*

$$\tilde{\mathbf{r}}_t = \tilde{L}_t \tilde{\mathbf{x}}_t = \tilde{L}_t \prod_{\tau=0}^{t-1} (\tilde{L}_\tau + D) \tilde{\mathbf{x}}_0,$$

in which $\tilde{L}_t = -(\delta \tilde{K}_{t+1} + cI)^{-1} \delta \tilde{K}_{t+1} D$, and $\tilde{K}_t$ is the Ricatti matrix defined by (16).

Finite-horizon proof: we use $\tilde{v}_t(\tilde{\mathbf{x}}_t)$ for the continuation value of the influencer at time t . In the transformed problem, $\tilde{v}_T(\tilde{\mathbf{x}}_T) = -\tilde{\mathbf{x}}_T' \tilde{\mathbf{x}}_T$, and

$$\tilde{v}_{T-1}(\tilde{\mathbf{x}}_{T-1}) = \max_{\tilde{\mathbf{r}}_{T-1}} - \left\{ \tilde{\mathbf{x}}_{T-1}' \tilde{\mathbf{x}}_{T-1} + c \tilde{\mathbf{r}}_{T-1}' \tilde{\mathbf{r}}_{T-1} + \delta (D \tilde{\mathbf{x}}_{T-1} + \tilde{\mathbf{r}}_{T-1})' (D \tilde{\mathbf{x}}_{T-1} + \tilde{\mathbf{r}}_{T-1}) \right\}.$$

Differentiate it with respect to $\tilde{\mathbf{r}}_{T-1}$ and set the derivative to zero we obtain

$$-(\delta + c) \tilde{\mathbf{r}}_{T-1} = \delta D \tilde{\mathbf{x}}_{T-1}, \text{ then } \tilde{\mathbf{r}}_{T-1} = -\frac{\delta}{\delta + c} D \tilde{\mathbf{x}}_{T-1}.$$

Since the objective function is concave, and the opinion evolution is linear, the FOC suffices. Substituting the optimal $\tilde{\mathbf{r}}_{T-1}$ into the value function at time $T - 1$, we have

$$\begin{aligned} \tilde{v}_{T-1}(\tilde{\mathbf{x}}_{T-1}) &= -\tilde{\mathbf{x}}_{T-1}' \tilde{\mathbf{x}}_{T-1} - c \tilde{\mathbf{r}}_{T-1}' \tilde{\mathbf{r}}_{T-1} - \delta \tilde{\mathbf{x}}_T' \tilde{\mathbf{x}}_T \\ &= -\tilde{\mathbf{x}}_{T-1}' \left(I + c \frac{\delta^2}{(\delta + c)^2} D^2 + \delta \frac{c^2}{(\delta + c)^2} D^2 \right) \tilde{\mathbf{x}}_{T-1}. \end{aligned}$$

This value function can be written as

$$\tilde{v}_{T-1}(\tilde{\mathbf{x}}_{T-1}) = -\tilde{\mathbf{x}}'_{T-1}\tilde{K}_{T-1}\tilde{\mathbf{x}}_{T-1}, \text{ where } \tilde{K}_{T-1} = I + \frac{\delta c}{\delta + c}D^2.$$

Note that the value function is quadratic in $\tilde{K}_{T-1}$. Clearly, $\tilde{K}_{T-1}$ is diagonal and positive definite (pd) with each j -th diagonal entry increasing in c and δ , and in λ_j^2 . Also, the optimal $\tilde{\mathbf{r}}_{T-1}$ characterizes the influencer's myopic strategy when she maximizes only the next period's payoff, which we use in the proof of Proposition 3. We now show if $\tilde{K}_t$ is diagonal and pd, then $\tilde{K}_{t-1}$ is also diagonal and pd. From the value function at $t-1$, we can show the optimal message

$$\tilde{\mathbf{r}}_{t-1} = -(\delta\tilde{K}_t + cI)^{-1}\delta\tilde{K}_tD\tilde{\mathbf{x}}_{t-1}.$$

Every matrix above is diagonal and $(\delta\tilde{K}_t + cI)^{-1}$ exists because it is pd. Therefore we can express the value function as:

$$\begin{aligned} \tilde{v}_{t-1}(\tilde{\mathbf{x}}_{t-1}) &= -\tilde{\mathbf{x}}'_{t-1}\tilde{\mathbf{x}}_{t-1} - c\tilde{\mathbf{r}}'_{t-1}\tilde{\mathbf{r}}_{t-1} + \delta\tilde{v}_t(\tilde{\mathbf{x}}_t) \\ &= -\tilde{\mathbf{x}}'_{t-1}\tilde{\mathbf{x}}_{t-1} - c\tilde{\mathbf{r}}'_{t-1}\tilde{\mathbf{r}}_{t-1} - \delta\tilde{\mathbf{x}}'_t\tilde{K}_t\tilde{\mathbf{x}}_t \\ &= -\tilde{\mathbf{x}}'_{t-1}\left(I + c((\delta\tilde{K}_t + cI)^{-1}\delta\tilde{K}_t)^2D^2 + \delta(I - (\delta\tilde{K}_t + cI)^{-1}\delta\tilde{K}_t)^2\tilde{K}_tD^2\right)\tilde{\mathbf{x}}_{t-1}. \end{aligned}$$

Simplifying above, we have:

$$\tilde{K}_{t-1} = I + \delta c(\delta\tilde{K}_t + cI)^{-1}\tilde{K}_tD^2. \quad (16)$$

Clearly, $\tilde{K}_{t-1}$ is diagonal and pd. Moreover, $\tilde{K}_t$ decreases in t by induction. First, $\tilde{K}_{T-1} > \tilde{K}_T = I$. Second, suppose this claim is true for all $\tilde{K}_t, \dots, \tilde{K}_{T-1}$. From equation (16) above, for every dimension j , $\tilde{K}_{t-1,jj} - \tilde{K}_{t,jj} = f(\tilde{K}_{t,jj}) - f(\tilde{K}_{t+1,jj})$, where $f(x) = \frac{\delta c \lambda_j^2 x}{\delta x + c}$. Since $f(\cdot)$ is increasing and $\tilde{K}_{t,jj} > \tilde{K}_{t+1,jj}$, we have $\tilde{K}_{t-1,jj} - \tilde{K}_{t,jj} > 0$ for all j .

Infinite-horizon proof: there are two ways to analyze the infinite-horizon case. First, we show that the finite-horizon Ricatti matrix has a well defined limit, which is the infinite-horizon Ricatti matrix. Second, we directly solve the fixed point problem (5) and study the properties of its solution $\tilde{K}^*$. In later proofs, we may use either method depending on convenience.

First, we show that the finite-horizon Ricatti matrix has a well defined limit. Recall that $\tilde{\mathbf{x}}_t = D^t\tilde{\mathbf{x}}_0$ under no invention, that is, under $\tilde{\mathbf{r}}_t = 0$ in every period. The value denoted as $\tilde{v}^\emptyset(\tilde{\mathbf{x}}_0) = -\tilde{\mathbf{x}}'_0(I - \delta D)^{-1}\tilde{\mathbf{x}}_0$ is bounded, where $\emptyset$ refers to the no intervention case. The

optimal strategy must do better, that is, $\tilde{v}(\tilde{\mathbf{x}}_0) > \tilde{v}^\emptyset(\tilde{\mathbf{x}}_0)$. Thus, $\tilde{K}_{T-t}$ is bounded and recall from above, it increases in t . A bounded and monotonically increasing sequence has a limit. Let $\lim_{t \rightarrow \infty} \tilde{K}_{T-t} = \tilde{K}^*$. Taking the limit of both sides of the finite-horizon Ricatti equation (16), we have the infinite-horizon Ricatti equation (5). Moreover, $\tilde{K}^*$ has all the properties of the finite-horizon Ricatti matrix, namely, diagonal, positive definite, and unique.

Second, we can directly study the infinite-horizon Ricatti equation. Note from the above that the value of this problem is bounded and thus a solution exists. We claim $\tilde{v}(\tilde{\mathbf{x}}) = -\tilde{\mathbf{x}}' K^* \tilde{\mathbf{x}}$, and verify this claim indeed holds. In period t , the influencer chooses $\tilde{\mathbf{r}}_t$ to maximize

$$\tilde{v}(\tilde{\mathbf{x}}_t) = -\tilde{\mathbf{x}}_t' \tilde{\mathbf{x}}_t - c \tilde{\mathbf{r}}_t' \tilde{\mathbf{r}}_t + \delta \tilde{v}(\tilde{\mathbf{x}}_{t+1}),$$

where $\tilde{\mathbf{x}}_{t+1} = D\tilde{\mathbf{x}}_t + \tilde{\mathbf{r}}_t$. Differentiate it with respect to $\tilde{\mathbf{r}}_t$, we can easily see the the unique solution (if exists) of infinite-horizon Ricatti equation is a necessary condition for optimality:

$$\tilde{K}^* = I + \delta c(\delta \tilde{K}^* + cI)^{-1} \tilde{K}^* D^2.$$

Next, we show that there exists a diagonal and pd matrix $\tilde{K}^*$ that solves this equation. If $\tilde{K}^*$ is diagonal, then the problem can be solved in each dimension j . At $\tilde{K}_{jj}^* = 1$, the left hand side (LHS) is smaller than the right hand side (RHS), but at $\tilde{K}_{jj}^* = \infty$, the LHS is bigger than the RHS. Moreover, the slope of the RHS with respect to $\tilde{K}_{jj}^*$ is positive but smaller than 1, which is the slope of the LHS. Thus, in each dimension, there is one intersection, which clearly defines $\tilde{K}^*$. Then, we have to show under the optimal strategy $\tilde{\mathbf{r}}_t = -(\delta \tilde{K}^* + cI)^{-1} \delta \tilde{K}^* D \tilde{\mathbf{x}}_t$, the opinions converge and thus the solution is also sufficient. It is easy to see that $\tilde{\mathbf{x}}_t = c^t (\delta \tilde{K}^* + cI)^{-t} D^t \tilde{\mathbf{x}}_0$. The matrix $c(\delta \tilde{K}^* + cI)^{-1} D$ is strictly substochastic with all eigenvalues strictly less than 1 in absolute value, and thus the system is asymptotically stable and the limit opinion is 0. Asymptotic stability is a necessary and sufficient condition for the existence of a unique Ricatti matrix that solves the infinite-horizon Ricatti equation, therefore the solution $\tilde{K}^*$ is unique. Putting all together, we verify the claim that the value function is $\tilde{v}(\tilde{\mathbf{x}}) = -\tilde{\mathbf{x}}' K^* \tilde{\mathbf{x}}$ for any opinion $\tilde{\mathbf{x}}$. $\parallel$

Proof of Proposition 2: Recall from expression (7) that the influencer's message in each period is a linear function of the current opinions. In dimension j ,

$$|\tilde{r}_t^j| = |\tilde{L}_{jj}^*| |\tilde{x}_t^j| = \frac{\delta \tilde{K}_{jj}^* |\lambda_j|}{\delta \tilde{K}_{jj}^* + c} \left(\frac{c |\lambda_j|}{\delta \tilde{K}_{jj}^* + c} \right)^t |\tilde{x}_0^j|.$$

The optimal message depends on the slope $|\tilde{L}_{jj}^*|$ as well as the second (exponential) term, capturing opinion convergence. Both terms depend on the endogenous Ricatti matrix. Since the value function is continuous, by Envelope Theorem, for any $\tilde{\mathbf{x}}_0$, the influencer is worse off as the cost of intervention c increases, or as δ increases since the future losses matter more. This observation immediately implies that $\tilde{K}^*$ increases in c and δ . Thus, as δ increases, $|\tilde{L}_{jj}^*|$ increases. Moreover, since for all j , $\tilde{K}_{jj}^*$ is the solution to $\delta(\tilde{K}_{jj}^*)^2 + (c - \delta - \delta c \lambda_j^2) \tilde{K}_{jj}^* - c = 0$, it is easy to see that $\tilde{K}_{jj}^*$ increases in λ_j^2 and $\tilde{K}_{jj}^*/c$ decreases in c by the Implicit Function Theorem. Thus, as c increases, $|\tilde{L}_{jj}^*|$ decreases.

The initial message $|\tilde{r}_0^j|$ depends only on the slope. Clearly, it increases in δ and decreases in c . The exponential term, however, decreases in δ and increases in c . As t increases, the exponential term dominates, and thus the later messages are less extreme when δ increases or c decreases. In terms of payoff, since the value function is $\tilde{v}(\tilde{\mathbf{x}}_0) = -\tilde{\mathbf{x}}_0' \tilde{K}^* \tilde{\mathbf{x}}_0$, it is easy to see that it decreases in δ , c , $|\lambda_j|$, and $\tilde{\mathbf{x}}_0$. ||

Proof of Observation 1: The first statement is shown in the previous proof, and we prove the ensuing two parts in turn. Part (1): the matrix J_N/N has a unique spectrum: $(1, 0, \dots, 0)$. Let D_J be a diagonal matrix with this spectrum in the diagonal. Recall that $A = UDU'$, it is easy to verify the $J_N = UD_JU'$. That is, they can be diagonalized by the same eigenvectors. Therefore, the spectrum of $A(\varepsilon)$ is simply $(1, (1 - \varepsilon)\lambda_2, \dots, (1 - \varepsilon)\lambda_N)$.

Part (2): because A and P are both symmetric and positive definite matrices, they have all positive eigenvalues. Let the eigenvalues of $A(\varepsilon)$ be $\{\lambda_1^\varepsilon, \dots, \lambda_N^\varepsilon\}$. Recall that λ_1^A is the largest eigenvalue of A and λ_N^A is the smallest eigenvalue; we label the eigenvalues of P similarly. By Weyl's inequality, since A, P are both symmetric, for all i

$$(1 - \varepsilon)\lambda_i^A + \varepsilon\lambda_N^P \leq \lambda_i^\varepsilon \leq (1 - \varepsilon)\lambda_i^A + \varepsilon\lambda_1^P.$$

If $\lambda_N^P > \lambda_2^A$, or the smallest eigenvalue of P is higher than the second largest eigenvalue of A , then $\lambda_i^\varepsilon > \lambda_i^A$ for all $i > 2$. ||

Proof of Proposition 3: The optimal strategy for each intervention is given in the text. Here we add the appropriate superscript to the opinions to be clear. The opinion evolution is given by $\tilde{\mathbf{x}}_t^{os} = D^{t-1} \tilde{\mathbf{x}}_1$ and $\tilde{\mathbf{x}}_t^\emptyset = D^t \tilde{\mathbf{x}}_0$ for one-shot and no intervention. Moreover, $\tilde{\mathbf{x}}_t^{mp} = \left(\frac{c}{\delta+c}\right)^t D^t \tilde{\mathbf{x}}_0$ for myopic intervention, and $\tilde{\mathbf{x}}_t^* = (cI + \delta \tilde{K}^*)^{-t} c^t D^t \tilde{\mathbf{x}}_0$ for optimal dynamic intervention. We now compare initial message intensity of active interventions.

From the text, in each dimension j ,

$$|\tilde{r}_{0,j}^{mp}| = \frac{\delta}{\delta + c} |\lambda_j \tilde{x}_0^j| < \frac{\delta}{\delta + c/\tilde{K}_{jj}^*} |\lambda_j \tilde{x}_0^j| = |\tilde{r}_{0,j}^*|.$$

Thus, dynamic intervention sends a stronger initial message in each dimension than that in myopic intervention. Next, as $|\tilde{r}_{0,j}^{os}| = \frac{\delta}{\delta + c(1 - \delta\lambda_j^2)} |\lambda_j \tilde{x}_0^j|$, we need to compare $1 - \delta\lambda_j^2$ versus $1/\tilde{K}_{jj}^*$. By (5),

$$\tilde{K}_{jj}^* = 1 + \frac{\delta c \tilde{K}_{jj}^* \lambda_j^2}{\delta \tilde{K}_{jj}^* + c} \Rightarrow \tilde{K}_{jj}^* - 1 = \frac{\delta \tilde{K}_{jj}^* \lambda_j^2}{\delta \tilde{K}_{jj}^* / c + 1} \Rightarrow 1 - 1/\tilde{K}_{jj}^* = \frac{\delta \lambda_j^2}{\delta \tilde{K}_{jj}^* / c + 1} < \delta \lambda_j^2.$$

Therefore, $|\tilde{r}_{0,j}^{os}| > |\tilde{r}_{0,j}^*|$, that is, one-shot intervention sends the strongest initial message. In terms of later messages, notice that $|\tilde{\mathbf{x}}_t^{mp}| > |\tilde{\mathbf{x}}_t^*|$ in every dimension because $\tilde{K}^* > I$ in the pd sense. When t is large enough, the lower magnitude of current opinions dominates the steeper slope, and thus the later messages are less extreme under optimal intervention than that under myopic intervention.

We now compare the payoffs from one-shot and myopic intervention. Substituting the optimal strategy $\tilde{\mathbf{r}}_0^{os}$ in the one-shot intervention, we get the one-shot influencer's payoff is:

$$\tilde{v}^{os}(\tilde{\mathbf{x}}_0) = -\tilde{\mathbf{x}}_0' \tilde{\mathbf{x}}_0 - c(\tilde{\mathbf{r}}_0^{os})' \tilde{\mathbf{r}}_0^{os} - \sum_1^\infty \delta^t (\tilde{\mathbf{x}}_1^{os})' D^{2(t-1)} \tilde{\mathbf{x}}_1^{os} \quad (17)$$

$$= -\tilde{\mathbf{x}}_0' \tilde{\mathbf{x}}_0 - \tilde{\mathbf{x}}_0' c \delta (c(I - \delta D^2) + \delta I)^{-1} D^2 \tilde{\mathbf{x}}_0. \quad (18)$$

As in the text, the myopic strategy is $\tilde{\mathbf{r}}_t^{mp} = -\frac{\delta}{\delta + c} D \tilde{\mathbf{x}}_t^{mp}$, and $\tilde{\mathbf{x}}_t^{mp} = \left(\frac{c}{\delta + c}\right)^t D^t \tilde{\mathbf{x}}_0$. Together, we have the total discounted payoff in myopic intervention is:

$$-\tilde{\mathbf{x}}_0' \left(I - \frac{\delta c^2}{(\delta + c)^2} D^2 \right)^{-1} \left(I + \frac{c \delta^2}{(\delta + c)^2} D^2 \right) \tilde{\mathbf{x}}_0. \quad (19)$$

Use dimension wise comparison here: since $\tilde{\mathbf{x}}_0$ is fixed, for dimension j , (19)-(18) is:

$$\begin{aligned} & -\delta c \lambda_j^2 \left(\frac{\delta + c}{(\delta + c)^2 - \delta c^2 \lambda_j^2} - \frac{1}{\delta + c(1 - \delta \lambda_j^2)} \right) \\ & = -\delta c \lambda_j^2 \left(\frac{\delta + c}{(\delta + c)^2 - \delta c^2 \lambda_j^2} - \frac{\delta + c}{(\delta + c)^2 - c \delta (\delta + c) \lambda_j^2} \right) > 0. \end{aligned}$$

Thus, the myopic payoff is always higher than that of the one-shot payoff for all δ, c and symmetric A . ||

Proof of Proposition 4: By standard result such as Proposition 3.1.1 of Bertsekas (2017), when opinions evolve according to $x_t = Ax_{t-1} + r_{t-1}$ and there is no discounting, if the pair (A, I) is controllable and (A', I) are controllable, the system given optimal strategy is stable. A pair (A, I) is controllable when the matrix $[I, A, \dots, A^{N-1}]$ has full rank N . Moreover, the Ricatti equation in the case of $\delta = 1$ is

$$K^* = I + A' (K^* - K^* (K^* + cI)^{-1} K^*) A$$

has a unique solution. Let $\hat{A} = \sqrt{\delta}A$, then we can rewrite our Ricatti equation given in the proposition as

$$K^* = I + (\sqrt{\delta}A)' \left(K^* - K^* \left(K^* + \frac{c}{\delta} I \right)^{-1} K^* \right) \sqrt{\delta}A.$$

After rewriting, the opinion system evolves according to $\hat{\mathbf{x}}_t = \hat{A}\hat{\mathbf{x}}_{t-1} + \hat{\mathbf{r}}_t$, where $\hat{\mathbf{r}}_t = \sqrt{\delta}\mathbf{r}_t$, and the cost of the $\hat{\mathbf{r}}_t$ is c/δ . Since $(\sqrt{\delta}A, I)$ and $((\sqrt{\delta}A)', I)$ satisfy the same controllability conditions, this system with discounting is also stable and $\hat{\mathbf{x}}_t \rightarrow 0$. Clearly, $\lim_{t \rightarrow \infty} v(\mathbf{x}_t) = -\mathbf{x}_t' K^* \mathbf{x}_t = 0$. Substitute in the optimal strategy, we have $\hat{\mathbf{x}}_t = (\sqrt{\delta})^t \mathbf{x}_t$ where $\hat{\mathbf{x}}_0 = \mathbf{x}_0$. Because at $\delta = 1$, $c(\delta K^* + cI)^{-1}A$ is stable with (absolute value) of all its eigenvalues strictly smaller than 1. Moreover, K^* is continuous in δ and the largest eigenvalue, as a root of a polynomial function is also continuous in δ . Therefore $c(\delta K^* + cI)^{-1}A$ is also stable for δ sufficiently high. $\parallel$

Proof of Proposition 5: we prove the result in a few steps.

Step 1: When T is finite, we claim that the Ricatti matrix $\tilde{K}_t^m$ is diagonal, positive definite, and $\tilde{K}_t^m = \tilde{K}_t$ for all t and m . To begin with, at the terminal period T , since $\tilde{v}^m(\tilde{\mathbf{x}}_T) = -(\tilde{\mathbf{x}}_T - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_T - \tilde{\mathbf{b}}^m)$, $\tilde{K}_T^m = I$, $\tilde{\mathbf{k}}_T^m = \tilde{\mathbf{b}}^m$, and $\tilde{\kappa}_T^m = 0$. Suppose the claim holds at t , we consider the value function at $t - 1$.

$$\begin{aligned} \tilde{v}_{t-1}^m(\tilde{\mathbf{x}}_{t-1}) &= -(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m) - c(\tilde{\mathbf{r}}_{t-1}^m)'(\tilde{\mathbf{r}}_{t-1}^m) \\ &\quad - \delta \left(D\tilde{\mathbf{x}}_{t-1} + \alpha \sum_{l=1}^M \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}_t^m \right)' \tilde{K}_t \left(D\tilde{\mathbf{x}}_{t-1} + \alpha \sum_{l=1}^M \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}_t^m \right) - \delta \tilde{\kappa}_t^m. \end{aligned}$$

From the FOC for agent m :

$$c\tilde{\mathbf{r}}_{t-1}^m = -\delta\alpha\tilde{K}_t \left(D\tilde{\mathbf{x}}_{t-1} + \alpha \sum_{l=1}^M \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}_t^m \right). \quad (20)$$

Summing up all the individual FOCs, we have

$$\alpha \sum \tilde{\mathbf{r}}_{t-1}^m = -(cI + \delta M \alpha^2 \tilde{K}_t)^{-1} \left(\delta M \alpha^2 \tilde{K}_t D \tilde{\mathbf{x}}_{t-1} - \delta \tilde{K}_t \alpha^2 \sum \tilde{\mathbf{k}}_t^m \right).$$

Note that $cI + \delta M \alpha^2 \tilde{K}_t$ is clearly positive definite and thus an inverse exists. For simplicity, let $H_t = (cI + \delta M \alpha^2 \tilde{K}_t)^{-1}$, which is diagonal and positive definite. Then the above equation becomes

$$\alpha \sum \tilde{\mathbf{r}}_{t-1}^m = -H_t \delta \alpha^2 \tilde{K}_t \left(M D \tilde{\mathbf{x}}_{t-1} - \sum \tilde{\mathbf{k}}_t^m \right).$$

It means that

$$\tilde{\mathbf{x}}_t = D \tilde{\mathbf{x}}_{t-1} + \alpha \sum \tilde{\mathbf{r}}_{t-1}^m = (I - H_t \delta M \alpha^2 \tilde{K}_t) D \tilde{\mathbf{x}}_{t-1} + H_t \delta \alpha^2 \tilde{K}_t \sum \tilde{\mathbf{k}}_t^m. \quad (21)$$

Substitute back into the individual FOC, and we have:

$$\begin{aligned} \tilde{\mathbf{r}}_{t-1}^m &= -\frac{\delta}{c} \alpha \tilde{K}_t \left((I - H_t \delta M \alpha^2 \tilde{K}_t) D \tilde{\mathbf{x}}_{t-1} + H_t \delta \alpha^2 \tilde{K}_t \sum \tilde{\mathbf{k}}_t^m - \tilde{\mathbf{k}}_t^m \right), \\ &= -\frac{\delta}{c} \alpha \tilde{K}_t \left(c H_t D \tilde{\mathbf{x}}_{t-1} + H_t \delta \alpha^2 \tilde{K}_t \sum \tilde{\mathbf{k}}_t^m - \tilde{\mathbf{k}}_t^m \right). \end{aligned}$$

The last equation is due to the fact that

$$I - H_t \delta M \alpha^2 \tilde{K}_t = I - (cI + \delta M \alpha^2 \tilde{K}_t)^{-1} \delta M \alpha^2 \tilde{K}_t = (cI + \delta M \alpha^2 \tilde{K}_t)^{-1} cI = cH_t.$$

We now solve for $\tilde{K}_{t-1}^m$ by substituting $\tilde{\mathbf{r}}_{t-1}^m$ into the value function at $t-1$:

$$\begin{aligned} \tilde{K}_{t-1}^m &= I + c \delta^2 \alpha^2 D' H_t' (\tilde{K}_t)^2 H_t D + c^2 \delta D' H_t' \tilde{K}_t H_t D \\ &= I + c \delta D' H_t' \tilde{K}_t (\delta \alpha^2 \tilde{K}_t + cI) H_t D. \end{aligned} \quad (22)$$

Clearly, since $\tilde{K}_t^m = \tilde{K}_t$, we know $\tilde{K}_{t-1}^m = \tilde{K}_{t-1}$ for all m . Moreover, because the right hand side of (22) is a positive definite diagonal matrix, $\tilde{K}_{t-1}$ is also such a matrix.

Step 2: The finite-horizon Ricatti matrix $\tilde{K}_t$ has a limit. Recall that $\tilde{K}_{t-1}$ is a function of $\tilde{K}_t$ according to (22). When A is symmetric, then for dimension j , we have

$$\tilde{K}_{t-1,jj} = 1 + \frac{c \delta (\delta \alpha^2 \tilde{K}_{t,jj} + c) \tilde{K}_{t,jj} \lambda_j^2}{(\delta \alpha^2 M \tilde{K}_{t,jj} + c)^2}.$$

Consider function $f(x) = 1 + \frac{c \delta (\delta \alpha^2 x + c) x \lambda_j^2}{(\delta \alpha^2 M x + c)^2}$, then clearly $\tilde{K}_{t-1,jj} = f(\tilde{K}_{t,jj})$. Notice that $f(\cdot)$ is differentiable, and

$$\frac{\partial f(x)}{\partial x} = \frac{c^2 \delta \lambda_j^2 (\delta \alpha^2 x (2 - M) + c)}{(\delta \alpha^2 M x + c)^3}, \quad (23)$$

Since $c^2 < (\delta\alpha^2 Mx + c)^2$ and $|\delta\alpha^2 x(2 - M) + c| < \delta\alpha^2 Mx + c$ for all $M > 1$, we have $|\partial f(x)/\partial x| < \delta$. Then by the Mean Value Theorem, there exists some point $g \geq 0$ such that

$$|\tilde{K}_{t-1,jj} - \tilde{K}_{t,jj}| = |f'(g)(\tilde{K}_{t,jj} - \tilde{K}_{t+1,jj})| < \delta |\tilde{K}_{t,jj} - \tilde{K}_{t+1,jj}|.$$

Thus, for any fixed T , there exists a limit of $\tilde{K}_{T-t,jj}$ when $t \rightarrow \infty$ because the distance $|\tilde{K}_{t-1,jj} - \tilde{K}_{t,jj}|$ is decreasing. In addition, since the terminal value $\tilde{K}_{T,jj} = 1$ and that $f(x) > 1$ for all $x > 0$, the limit $\tilde{K}_{jj}^* \geq 1$.

Step 3: There exists a MPE when $T = \infty$. First, we show that $\tilde{K}^*$ is the solution to the infinite-horizon Ricatti equation. Take the limit of both sides of equation (22), which exists by the argument above. Let $H = \lim_{t \rightarrow \infty} H_t$, which is also diagonal and pd. We have equation (11), reproduced here:

$$\tilde{K}^* = I + c\delta DH\tilde{K}^*(\delta\alpha^2\tilde{K}^* + cI)HD.$$

It is easy to see equation (11) is a necessary condition for influencer m to maximize his payoff given the other influencers' strategies. Next, we show the dynamic system of opinion evolution given the influencers' strategies is stable and thus the value function above is well defined and there exists a unique solution to the Ricatti equation. To do so, we focus on the special case when $b^m = 0$, that is, all influencers have the same agenda. Note that this assumption does not affect $\tilde{K}^*$ since the agendas only affect the constant terms in the best response equation (20). Then, the value function contains only the quadratic term with $\tilde{K}^*$, and the projected opinions evolve according to

$$\tilde{\mathbf{x}}_t = (I - H\delta M\alpha^2\tilde{K}^*)D\tilde{\mathbf{x}}_{t-1},$$

which is the limit of equation (21). Note that $(I - H\delta M\alpha^2\tilde{K}^*)D$ is a diagonal matrix with every diagonal entry strictly smaller than 1, so $\tilde{\mathbf{x}}_t = (I - H\delta M\alpha^2\tilde{K}^*)^t D^t \tilde{\mathbf{x}}_0$ goes to $\mathbf{0}$ as $t \rightarrow \infty$. Therefore, the solution to equation (11) characterizes each influencer's equilibrium strategy. From the standard result on Lyapunov equations, since $I > 0$, $c > 0$, and the matrix $(I - H\delta M\alpha^2\tilde{K}^*)D$ has all eigenvalues with absolute value between $(0, 1)$, there exists a unique $\tilde{K}^*$ that solves equation (11). Together, the optimal strategy given in Proposition 5, equation (11) and (24) below define the unique MPE of the infinite horizon game.

Step 4: We now show that in the unique MPE, agents reach consensus. We have characterized $\tilde{K}^*$ above. We now consider the evolution of the linear term $\tilde{\mathbf{k}}^m$ which depends on the

influencers' agendas. Using equation (11), we have

$$\begin{aligned}
\tilde{K}^* \tilde{\mathbf{k}}^m &= \tilde{\mathbf{b}}^m - \delta^2 \alpha^2 DH (\tilde{K}^*)^2 (H \delta \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m - \tilde{\mathbf{k}}^m) - c \delta DH \tilde{K}^* (H \delta \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m - \tilde{\mathbf{k}}^m) \\
&= \tilde{\mathbf{b}}^m - \left(\delta^2 \alpha^2 DH (\tilde{K}^*)^2 + c \delta DH \tilde{K}^* \right) (H \delta \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m - \tilde{\mathbf{k}}^m) \\
&= \tilde{\mathbf{b}}^m - (\delta DH \tilde{K}^* (\delta \alpha^2 \tilde{K}^* + cI)) (H \delta \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m - \tilde{\mathbf{k}}^m).
\end{aligned}$$

Summing it up for m , we have (note we need D to be invertible here).

$$\begin{aligned}
\tilde{K}^* \sum \tilde{\mathbf{k}}^m &= \sum \tilde{\mathbf{b}}^m + (\tilde{K}^* - I) D^{-1} (H)^{-1} (H \delta M \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m - \sum \tilde{\mathbf{k}}^m) \\
&= \sum \tilde{\mathbf{b}}^m + (\tilde{K}^* - I) D^{-1} \sum \tilde{\mathbf{k}}^m.
\end{aligned} \tag{24}$$

The last equality holds by substituting equation (11) into the right hand side of the expression. Recall that $\tilde{K}^*$ is diagonal. Moreover, we know that $\tilde{\mathbf{b}}^m = U' b^m \mathbf{1} = (\sqrt{N} b^m, 0, \dots, 0)'$. From Step 3, the value function has no linear terms if all influencers' agendas are zero, that is, $\tilde{\mathbf{k}}_j^m = 0$ for $j \neq 1$. As $\tilde{K}^* - (\tilde{K}^* - I) D^{-1}$ is a diagonal matrix with the first diagonal entry being 1, in the first dimension we have $\sum \tilde{\mathbf{k}}_1^m = \sum \tilde{\mathbf{b}}_1^m = \sqrt{N} \sum b^m$; and all other dimensions of $\sum \tilde{\mathbf{k}}^m$ are 0. As a result, $\sum \tilde{\mathbf{k}}^m$ is uniquely defined.

Next, given the influencer's strategies, the network opinions evolve according to

$$\tilde{\mathbf{x}}_t = D \tilde{\mathbf{x}}_{t-1} + \alpha \sum \tilde{\mathbf{r}}_{t-1}^m = (I - H \delta M \alpha^2 \tilde{K}^*) D \tilde{\mathbf{x}}_{t-1} + H \delta \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m. \tag{25}$$

Recall that $(I - H \delta M \alpha^2 \tilde{K}^*) = cH$, and we can rewrite the above as

$$(\tilde{\mathbf{x}}_t - \beta) = cHD(\tilde{\mathbf{x}}_{t-1} - \beta),$$

where $\beta = (I - cHD)^{-1} H \delta \alpha^2 \tilde{K}^* \sum \tilde{\mathbf{k}}^m$. Clearly, (absolute values) of all eigenvalues of cHD are smaller than 1, and the above system converges to $\lim_{t \rightarrow \infty} \tilde{\mathbf{x}}_t - \beta = 0$, or $\tilde{\mathbf{x}}_\infty = \beta$ is the limit belief. Since $\sum \tilde{\mathbf{k}}_j^m = 0$ for all $j \neq 1$ and $\sum \tilde{k}_1^m = \sqrt{N} \sum b^m$, it is easy to see that $\beta = \left(\sqrt{N} \sum b^m / M, 0, \dots, 0 \right)$ and thus the (unprojected) limit opinions are $\mathbf{x}_\infty = (\sum b^m / M, \dots, \sum b^m / M)'$ in the unique MPE of this game. ||

Proof of Corollary 1: From the proof of Proposition 5, influencer m 's strategy is linear, not affine, in the projected opinions when $b^m = 0$ for all m . Therefore, the payoff depends only on $\tilde{K}^*$, which is diagonal, pd, and identical for all m . Rewrite the infinite-horizon Ricatti equation (11) as a function of M :

$$\tilde{K}^*(M) = I + c \delta DH(M) \tilde{K}^*(M) (\delta \alpha^2 \tilde{K}^*(M) + cI) H(M) D, \tag{26}$$

where $H(M) = (cI + \delta M \alpha^2 \tilde{K}^*(M))^{-1}$. We now show that each diagonal entry of $\tilde{K}^*(M)$ decreases in M . In dimension j ,

$$\tilde{K}_{jj}^*(M) = 1 + \frac{c\delta K_{jj}^*(M)(\delta\alpha^2 K_{jj}^*(M) + c)\lambda_j^2}{(\delta M \alpha^2 K_{jj}^*(M) + c)^2}.$$

Let $x = \tilde{K}_{jj}^*(M)$. The above equation can be rewritten as

$$g(x, M) \equiv 1 + \frac{c\delta x(\delta\alpha^2 x + c)\lambda_j^2}{(\delta M \alpha^2 x + c)^2} - x = 0. \quad (27)$$

It is easy to see that $\partial g / \partial M < 0$. Next,

$$\frac{\partial g}{\partial x} = \frac{c^2 \delta \lambda_j^2 (\delta \alpha^2 x (2 - M) + c)}{(M \delta \alpha^2 x + c)^3} - 1.$$

Note that $c^2 < (M \delta \alpha^2 x + c)^2$. Also, $(\delta \alpha^2 x (2 - M) + c) < M \delta \alpha^2 x + c$ for all $M > 1$. Therefore, the ratio above is smaller than 1, which implies $\frac{\partial g}{\partial x} < 0$. By implicit function theorem, $\tilde{K}_j^*(M)$ strictly decreases in M . Recall that each influencer's payoff is $-\tilde{\mathbf{x}}_0' \tilde{K}^*(M) \tilde{\mathbf{x}}_0$. Thus, each influencer's payoff increases in M .

Next, we examine the convergence speed and intervention intensity. Recall that

$$\tilde{\mathbf{x}}_t = D\tilde{\mathbf{x}}_{t-1} + M\alpha\tilde{\mathbf{r}}_{t-1} = \left(I - (\delta M \alpha^2 \tilde{K}^*(M) + cI)^{-1} \delta M \alpha^2 \tilde{K}^*(M)\right)^t D^t \tilde{\mathbf{x}}_0.$$

Using the fact that $\tilde{K}^*(M)$ is diagonal, we can rewrite it as:

$$\tilde{\mathbf{x}}_t = \left(\frac{c}{\delta M \alpha^2 \tilde{K}_{11}^*(M) + c}\right)^t \lambda_1^t \tilde{x}_0^1 + \dots + \left(\frac{c}{\delta M \alpha^2 \tilde{K}_{NN}^*(M) + c}\right)^t \lambda_N^t \tilde{x}_0^N.$$

We now show that $M\tilde{K}_{jj}^*(M)$ increases in M for every dimension j , and thus the convergence speed strictly increases in M . Let $y = Mx$ (recall $x = \tilde{K}_{jj}^*(M)$), and let $x' = \frac{dx}{dM}$, $y' = \frac{dy}{dM}$. Differentiating $g(x, M)$ in (27) with respect to M , we have:

$$-x' \left(1 - \frac{c\delta(2\delta\alpha^2 x + c)\lambda_j^2}{(\delta M \alpha^2 x + c)^2}\right) = \frac{2c\delta^2 \alpha^2 x(\delta\alpha^2 x + c)\lambda_j^2}{(\delta M \alpha^2 x + c)^3} y'.$$

We know that $x' < 0$ from above. Moreover, the second term on the left hand side is positive because $c < \delta M \alpha^2 x + c$ and $2 \leq M$. Therefore $y' > 0$, that is, $M\tilde{K}_{jj}^*(M)$ increases in M for all j . Therefore, the convergence speed increases in M and the absolute value of (14) decreases in M for every influencer.

We now compare the total payoff of the M influencers with that of a single representative influencer. Recall the representative influencer's stage payoff is $-M\tilde{\mathbf{x}}'_t\tilde{\mathbf{x}}_t - c\sum_1^M(\tilde{\mathbf{r}}_t^m)'\tilde{\mathbf{r}}_t^m$. Since the cost is quadratic, for any aggregate message $\tilde{\mathbf{r}}_t = \sum_1^M \tilde{\mathbf{r}}_t^m$, it is optimal to send M identical messages. Her value function in period 0 is $\tilde{v}_0^s(M) = -M\tilde{\mathbf{x}}'_0\tilde{K}^s(M)\tilde{\mathbf{x}}_0$, where $\tilde{K}^s(M)$ is the Ricatti matrix of the influencer who sends message $\tilde{\mathbf{r}}_t/M$. In this way, we can compare $\tilde{K}^s(M)$ directly with $\tilde{K}^m(M)$ in the multiple-influencer model. We have

$$\tilde{v}_t^s(M) = -M\tilde{\mathbf{x}}'_t\tilde{\mathbf{x}}_t - \frac{c}{M}\tilde{\mathbf{r}}'_t\tilde{\mathbf{r}}_t - \delta M(D\tilde{\mathbf{x}}'_t + \alpha\tilde{\mathbf{r}}_t)'\tilde{K}^s(M)(D\tilde{\mathbf{x}}'_t + \alpha\tilde{\mathbf{r}}_t).$$

Her optimal strategy is

$$\left(\delta M\alpha^2\tilde{K}^s(M) + \frac{c}{M}I\right)\tilde{\mathbf{r}}_t = -\delta M\alpha\tilde{K}^s(M)D\tilde{\mathbf{x}}_t.$$

The Ricatti equation becomes:

$$\tilde{K}^s(M) = I + \delta c(\delta M^2\alpha^2\tilde{K}^s(M) + cI)^{-1}\tilde{K}^s(M)D^2. \quad (28)$$

It is easy to verify that $\tilde{K}^s(M)$ decreases in M and $M\tilde{K}^s(M)$ increases in M similar to the multiple-influencer case above.

The Ricatti matrix in (26) and (28) are both diagonal, so we can compare their respective diagonal entry. In dimension j , we have

$$g^s(x) = 1 + \frac{\delta c}{\delta M^2\alpha^2x + c}x\lambda_j^2, \text{ and } g^*(x) = 1 + \frac{\delta c(\delta\alpha^2x + c)}{(\delta M\alpha^2x + c)^2}x\lambda_j^2.$$

We can verify that $g^s(x) \leq g^*(x)$ and the equality holds if and only if $M = 1$. So for $M \geq 2$, $K_{jj}^s(M) < K_{jj}^*(M)$, which implies that the single representative influencer can do better than the multiple influencers. For convergence speed, we need to compare $M^2K_{jj}^s(M)$ and $MK_{jj}^*(M)$, or equivalently to compare $MK_{jj}^s(M)$ and $K_{jj}^*(M)$. Notice that at $M = 1$, $MK_{jj}^s(M) = K_{jj}^*(M)$. Since the LHS increases in M and the RHS decreases in M , $MK_{jj}^s(M) > K_{jj}^*(M)$ for all $M > 1$. Thus, the opinions converge faster under the single representative influencer than under M influencers. ||

Proof of Proposition 6: Since A is asymmetric, we do not decompose A and thus all variables are in their original forms. In the infinite-horizon model, we focus on the semi-symmetric equilibrium in which all influencers share the same Ricatti matrix $K^m = K^*$ if such an equilibrium exists. The finite-horizon model is similar with the exception that the

unique MPE must be semi-symmetric in which all influencers share the same Ricatti matrix K_t^* for $0 \leq t \leq T$ by backward induction.

First, we characterize the infinite-horizon Ricatti equation (15), which is a necessary condition of the equilibrium. We assume $K^m = K^*$ is a positive definite matrix for all m and derive the value function of each influencer.

$$\begin{aligned} v^m(\mathbf{x}_{t-1}) = & -(\mathbf{x}_{t-1} - \mathbf{b}^m)'(\mathbf{x}_{t-1} - \mathbf{b}^m) - c(\mathbf{r}_{t-1}^m)'(\mathbf{r}_{t-1}^m) \\ & - \delta \left(A\mathbf{x}_{t-1} + \alpha \sum_{l=1}^M \mathbf{r}_{t-1}^l - \mathbf{k}^m \right)' K^* \left(A\mathbf{x}_{t-1} + \alpha \sum_{l=1}^M \mathbf{r}_{t-1}^l - \mathbf{k}^m \right) - \delta \kappa^m. \end{aligned}$$

From the FOC for influencer m :

$$c\mathbf{r}_{t-1}^m = -\delta\alpha K^* \left(A\mathbf{x}_{t-1} + \alpha \sum \mathbf{r}_{t-1}^m - \mathbf{k}^m \right). \quad (29)$$

Summing up all the individual FOCs, we have

$$\alpha \sum \mathbf{r}_{t-1}^m = -(cI + \delta M \alpha^2 K^*)^{-1} \left(\delta M \alpha^2 K^* A\mathbf{x}_{t-1} - \delta K^* \alpha^2 \sum \mathbf{k}^m \right).$$

Note that $cI + \delta M \alpha^2 K^*$ is clearly positive definite and thus an inverse exists. For simplicity, let $H = (cI + \delta M \alpha^2 K^*)^{-1}$, which is positive definite. Then the above equation becomes

$$\alpha \sum_m \mathbf{r}_{t-1}^m = -H \delta \alpha^2 K^* \left(M A\mathbf{x}_{t-1} - \sum \mathbf{k}^m \right). \quad (30)$$

Similar to the calculations in Step 1 of the proof of Proposition 5, we get

$$\begin{aligned} K^* &= I + c\delta^2 \alpha^2 A' H' (K^*)^2 H A + c^2 \delta A' H' K^* H A \\ &= I + c\delta A' H' K^* (\delta \alpha^2 K^* + cI) H A. \end{aligned}$$

Moreover, we can use the Ricatti equation to obtain the evolution of $\mathbf{k}^m$:

$$\begin{aligned} K^* \mathbf{k}^m &= \mathbf{b}^m - (\delta A H K^* (\delta \alpha^2 K^* + cI)) (H \delta \alpha^2 K^* \sum \mathbf{k}^m - \mathbf{k}^m), \\ K^* \sum \mathbf{k}^m &= \sum \mathbf{b}^m + (K^* - I) A^{-1} \sum \mathbf{k}^m. \end{aligned} \quad (31)$$

Note that the Ricatti equation depends only on the quadratic terms of $\mathbf{x}_t$, and thus the Ricatti equation (15) is the same as one in which all $\mathbf{k}^m = \mathbf{0}$ and $\kappa^m = 0$. We can rewrite the Ricatti equation (15) as

$$(cHA)' K^* cHA - K^* + I + c\delta^2 \alpha^2 A' H' (K^*)' I K^* H A = 0.$$

Because $I > 0$ and $cI > 0$, $I + c\delta^2\alpha^2 A'H'(K^*)'IK^*HA > 0$, this is a Lyapunov equation. If a positive definite solution K^* exists, then by the known theorem on Lyapunov inequality (see for example, Theorem 8.4 from [Hespanha \(2018\)](#)), the transition matrix cHA is stable and has all eigenvalues strictly smaller than 1 (in absolute value).

Next, the network opinions evolve according to

$$\mathbf{x}_t = A\mathbf{x}_{t-1} + \alpha \sum \mathbf{r}_{t-1}^m = cHA\mathbf{x}_{t-1} + H\delta\alpha^2 K^* \sum \mathbf{k}^m.$$

Again, this is an affine dynamic system and we can rewrite the above as

$$(\mathbf{x}_t - \boldsymbol{\beta}) = cHA(\mathbf{x}_{t-1} - \boldsymbol{\beta}),$$

where $\boldsymbol{\beta} = (I - cHA)^{-1}H\delta\alpha^2 K^* \sum \mathbf{k}^m$. Note that because cHA is asymptotically stable, the magnitude of all eigenvalues is strictly smaller than 1, $(I - cHA)$ is pd and thus invertible. The above system converges to $\lim_{t \rightarrow \infty} \mathbf{x}_t - \boldsymbol{\beta} = 0$, that is, the network opinions converge, and $\mathbf{x}_\infty = \boldsymbol{\beta}$ is the limit opinion.

Next, we know from equation (31) that

$$(K^* - (K^* - I)A^{-1}) \sum \mathbf{k}^m = \sum \mathbf{b}^m.$$

If $K^* - (K^* - I)A^{-1}$ is non-singular, the solution $\sum \mathbf{k}^m$ is unique. We now show in the limit, consensus is always an equilibrium outcome. In this case, $\mathbf{x}_\infty$ is a constant vector and thus $\mathbf{x}_\infty = A\mathbf{x}_\infty$. Since $\mathbf{x}_\infty = A\mathbf{x}_\infty + \alpha \sum \mathbf{r}_\infty^m$, the sum of long run messages $\sum_1^M \mathbf{r}_\infty^m = \mathbf{0}$. Taking (30) to the limit, we have $A\mathbf{x}_\infty = \mathbf{x}_\infty = \sum \mathbf{k}^m / M$. So, $\sum \mathbf{k}^m$ is a constant vector. Substitute into equation (31) and use the fact $M\mathbf{x}_\infty = \sum \mathbf{k}^m = A^{-1} \sum \mathbf{k}^m$, we have

$$K^* M\mathbf{x}_\infty = \sum \mathbf{b}^m + (K^* - I)M\mathbf{x}_\infty.$$

That is, $\mathbf{x}_\infty = \sum \mathbf{b}^m / M$. Therefore there always exists a semi-symmetric MPE in which all influencers have the same K^* and the network opinions converge to consensus $\mathbf{x}_\infty$. This MPE is the unique semi-symmetric MPE if $K^* - (K^* - I)A^{-1}$ is non-singular. $\parallel$

Proof of Proposition 7: Since A is symmetric, we decompose it as $A = UDU'$ and transform the problem into one using the projected variables.

Step 1: There exist a unique pair of diagonal and pd Ricatti matrices $(\tilde{K}^1, \tilde{K}^2)$ that solve the influencers' infinite-horizon Ricatti equations. Moreover, $\tilde{K}^1 > \tilde{K}^2$ in the pd sense.

We assume that the value function of influencer $m = 1, 2$ exists when $T = \infty$ and takes the following quadratic form, and we will verify this assumption later at the end of Step 2.

$$\begin{aligned} \tilde{v}_{t-1}^m(\tilde{\mathbf{x}}_{t-1}) = & -(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m) - c(\tilde{\mathbf{r}}_{t-1}^m)'(\tilde{\mathbf{r}}_{t-1}^m) \\ & - \delta \left(D\tilde{\mathbf{x}}_{t-1} + \sum_{l=1}^M \alpha_l \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}^m \right)' \tilde{K}^m \left(D\tilde{\mathbf{x}}_{t-1} + \sum_{l=1}^M \alpha_l \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}^m \right) - \delta \tilde{\kappa}^m. \end{aligned}$$

We first derive the coupled Ricatti equations under the assumption that $\tilde{K}^m$ is diagonal and pd and will verify later that there exist such Ricatti matrices that solve the Ricatti equations. From the FOC for influencer m :

$$c\tilde{\mathbf{r}}_{t-1}^m = -\delta\alpha_m\tilde{K}^m \left(D\tilde{\mathbf{x}}_{t-1} + \sum \alpha_m \tilde{\mathbf{r}}_{t-1}^m - \tilde{\mathbf{k}}^m \right). \quad (32)$$

Rearrange and we have:

$$\left(cI + \delta\alpha_1^2\tilde{K}^1 + \delta\alpha_2^2\tilde{K}^2 \right) \tilde{\mathbf{r}}_{t-1}^1 = -\delta\alpha_1\tilde{K}^1 \left(D\tilde{\mathbf{x}}_{t-1} - \frac{\delta}{c}\alpha_2^2\tilde{K}^2 \left(\tilde{\mathbf{k}}^1 - \tilde{\mathbf{k}}^2 \right) - \tilde{\mathbf{k}}^1 \right); \quad (33)$$

$$\left(cI + \delta\alpha_1^2\tilde{K}^1 + \delta\alpha_2^2\tilde{K}^2 \right) \tilde{\mathbf{r}}_{t-1}^2 = -\delta\alpha_2\tilde{K}^2 \left(D\tilde{\mathbf{x}}_{t-1} - \frac{\delta}{c}\alpha_1^2\tilde{K}^1 \left(\tilde{\mathbf{k}}^2 - \tilde{\mathbf{k}}^1 \right) - \tilde{\mathbf{k}}^2 \right). \quad (34)$$

We can focus on the evolution of $\tilde{K}^m$ because it is independent of $\tilde{\mathbf{k}}^m$ and b^m . Let $\Gamma = (cI + \delta\alpha_1^2\tilde{K}^1 + \delta\alpha_2^2\tilde{K}^2)^{-1}$. Substitute equation (33) and (34) into the evolution of opinions:

$$\tilde{\mathbf{x}}_t = D\tilde{\mathbf{x}}_{t-1} - \delta\alpha_1^2\Gamma\tilde{K}^1D\tilde{\mathbf{x}}_{t-1} - \delta\alpha_2^2\Gamma\tilde{K}^2D\tilde{\mathbf{x}}_{t-1} = (I - \delta\alpha_1^2\Gamma\tilde{K}^1 - \delta\alpha_2^2\Gamma\tilde{K}^2)D\tilde{\mathbf{x}}_{t-1}.$$

Note that $(I - \delta\alpha_1^2\Gamma\tilde{K}^1 - \delta\alpha_2^2\Gamma\tilde{K}^2) = c\Gamma$, and thus we have $\tilde{\mathbf{x}}_t = c\Gamma D\tilde{\mathbf{x}}_{t-1}$. Combined with $\tilde{\mathbf{r}}_{t-1}^m = -\delta\alpha_m\Gamma\tilde{K}^mD\tilde{\mathbf{x}}_{t-1}$, we can derive the coupled Ricatti equations. Note that influencer m 's Ricatti matrix depends on the other's Ricatti matrix through Γ .

$$\tilde{K}^1 = I + c\delta D\Gamma\tilde{K}^1(\delta\alpha_1^2\tilde{K}^1 + cI)\Gamma D; \quad (35)$$

$$\tilde{K}^2 = I + c\delta D\Gamma\tilde{K}^2(\delta\alpha_2^2\tilde{K}^2 + cI)\Gamma D. \quad (36)$$

The above coupled Ricatti equations are necessary conditions for the infinite-horizon game to have a MPE. Next, we show that there exist positive definite and diagonal solutions $(\tilde{K}^1, \tilde{K}^2)$ to the coupled Ricatti equations. In each dimension j , let $x = \tilde{K}_{jj}^1$ and $y = \tilde{K}_{jj}^2$, and the coupled Ricatti equations become

$$x = 1 + \frac{c\delta\lambda_j^2x(\delta\alpha_1^2x + c)}{(c + \delta\alpha_1^2x + \delta\alpha_2^2y)^2}; \quad (37)$$

$$y = 1 + \frac{c\delta\lambda_j^2y(\delta\alpha_2^2y + c)}{(c + \delta\alpha_1^2x + \delta\alpha_2^2y)^2}. \quad (38)$$

Because $(\tilde{K}^1, \tilde{K}^2)$ need to be pd and diagonal, we search for solutions in the range $x \geq 1, y \geq 1$. Subtracting RHS from LHS of equation (37) defines an implicit function $g_1(x, y) = 0$, and we can show $\frac{\partial g_1(x, y)}{\partial x} > \frac{\partial g_1(x, y)}{\partial y} > 0$. Given $g_1(x, y) = 0$, x strictly decreases in y and $|\partial y / \partial x| > 1$. Similarly, subtracting RHS from LHS of equation (38) defines an implicit function $g_2(x, y) = 0$, and we can show $\frac{\partial g_2(x, y)}{\partial y} > \frac{\partial g_2(x, y)}{\partial x} > 0$. Given $g_2(x, y) = 0$, we can show that y strictly decreases in x and $|\partial y / \partial x| < 1$. Thus, the two curves $g_1(x, y) = 0$ and $g_2(x, y) = 0$ intersect at most once. Moreover, at $y = 1, x > 1$ and at $y \rightarrow \infty, x \rightarrow 1$, and vice versa. So there exists a unique solution $\tilde{K}_{jj}^1 > 1, \tilde{K}_{jj}^2 > 1$ for each dimension j .

Then, we show $\tilde{K}^1 > \tilde{K}^2$. Notice that whenever $\alpha_1 \neq \alpha_2, x \neq y$. Also, $\alpha_1^2 x \neq \alpha_2^2 y$, because otherwise $x = y$ is a solution. Recall $\alpha_1 > \alpha_2$ by assumption. Multiply equation (37) by α_1^2 and equation (38) by α_2^2 and take their difference. Let $x' = \alpha_1^2 x, y' = \alpha_2^2 y$, we have:

$$x' - y' = \alpha_1^2 - \alpha_2^2 + \frac{c\delta\lambda_j^2}{(c + \delta x' + \delta y')^2}(\delta((x')^2 - (y')^2) + c(x' - y')).$$

Since $\alpha_1^2 x \neq \alpha_2^2 y$ from above, divide $x' - y'$ from both sides, and we have:

$$1 = \frac{\alpha_1^2 - \alpha_2^2}{x' - y'} + \frac{c\delta\lambda_j^2}{c + \delta x' + \delta y'}.$$

The second term is strictly between 0 and 1, and thus $0 < \frac{\alpha_1^2 - \alpha_2^2}{x' - y'} < 1$, or $x' > y'$. Together with the fact that $\frac{x-1}{y-1} = \frac{x}{y} \frac{\delta x' + c}{\delta y' + c}$, we have $\frac{x-1}{y-1} < \frac{x}{y}$ and thus $x > y$. Because $\tilde{K}_{jj}^1 > \tilde{K}_{jj}^2$ for all j , $\tilde{K}^1 > \tilde{K}^2$ in the positive definite sense.

Step 2: The network opinions converge: $\lim_{t \rightarrow \infty} \tilde{\mathbf{x}}_t = \tilde{\mathbf{x}}_\infty$. Since $b^m \neq 0$ in general, we also need to derive the evolution of $\tilde{\mathbf{k}}^m$. Summing up the weighted messages, we have:

$$\sum \alpha_m \tilde{\mathbf{r}}_t^m = -\delta \Gamma \left((\alpha_1^2 \tilde{K}^1 + \alpha_2^2 \tilde{K}^2) D \tilde{\mathbf{x}}_t - (\alpha_1^2 \tilde{K}^1 \tilde{\mathbf{k}}^1 + \alpha_2^2 \tilde{K}^2 \tilde{\mathbf{k}}^2) \right). \quad (39)$$

Let $\mathbf{w} = \alpha_1^2 \tilde{K}^1 \tilde{\mathbf{k}}^1 + \alpha_2^2 \tilde{K}^2 \tilde{\mathbf{k}}^2$. From above, we have

$$\tilde{\mathbf{x}}_t = c\Gamma D \tilde{\mathbf{x}}_{t-1} + \delta \Gamma \mathbf{w}; \quad (40)$$

$$\tilde{\mathbf{r}}_{t-1}^m = -\frac{\delta}{c} \alpha_m \tilde{K}^m (c\Gamma D \tilde{\mathbf{x}}_{t-1} + \delta \Gamma \mathbf{w} - \tilde{\mathbf{k}}^m). \quad (41)$$

Together, we have the evolution of $\tilde{\mathbf{k}}^m$, which only depends on $\tilde{K}^m$ and thus is well-defined:

$$\tilde{K}^1 \tilde{\mathbf{k}}^1 = \tilde{\mathbf{b}}^1 - \delta \Gamma D (\delta \alpha_1^2 \tilde{K}^1 + cI) (\delta \tilde{K}^1 \Gamma \mathbf{w} - \tilde{K}^1 \tilde{\mathbf{k}}^1); \quad (42)$$

$$\tilde{K}^2 \tilde{\mathbf{k}}^2 = \tilde{\mathbf{b}}^2 - \delta \Gamma D (\delta \alpha_2^2 \tilde{K}^2 + cI) (\delta \tilde{K}^2 \Gamma \mathbf{w} - \tilde{K}^2 \tilde{\mathbf{k}}^2). \quad (43)$$

We can then solve for $\tilde{K}^m \tilde{\mathbf{k}}^m$ as a function of $\tilde{K}^m$ and other parameters. To simplify the derivation, let

$$Y = I - \delta \Gamma^2 D(\delta \alpha_1^2 \tilde{K}^1 + cI)(\delta \alpha_2^2 \tilde{K}^2 + cI),$$

$$Z_1 = \delta^2 \alpha_2^2 \Gamma D(\delta \alpha_1^2 \tilde{K}^1 + cI) \tilde{K}^1 \Gamma, \text{ and } Z_2 = \delta^2 \alpha_1^2 \Gamma D(\delta \alpha_2^2 \tilde{K}^2 + cI) \tilde{K}^2 \Gamma.$$

Then straightforward calculations can show that

$$(Y^2 - Z_1 Z_2) \tilde{K}^1 \tilde{\mathbf{k}}^1 = Y \tilde{\mathbf{b}}^1 - Z_1 \tilde{\mathbf{b}}^2; \quad (44)$$

$$(Y^2 - Z_1 Z_2) \tilde{K}^2 \tilde{\mathbf{k}}^2 = Y \tilde{\mathbf{b}}^2 - Z_2 \tilde{\mathbf{b}}^1; \quad (45)$$

Moreover, it is easy to show $Y^2 - Z_1 Z_2 > 0$. Clearly, $\tilde{K}^1 \tilde{\mathbf{k}}^1$ is strictly increasing in $\tilde{\mathbf{b}}^1$ and decreasing in $\tilde{\mathbf{b}}^2$, and vice versa for $\tilde{K}^2 \tilde{\mathbf{k}}^2$. Moreover, because $\tilde{K}^1$ and $\tilde{K}^2$ are unique, there exist a unique pair of $\tilde{\mathbf{k}}^1$ and $\tilde{\mathbf{k}}^2$. Given the opinion evolution process (40), since $c\Gamma D$ is substochastic, the dynamic system is stable and the opinions converge to

$$\tilde{\mathbf{x}}_\infty = (I - c\Gamma D)^{-1} \delta \Gamma \mathbf{w} = \left(\alpha_1^2 \tilde{K}^1 + \alpha_2^2 \tilde{K}^2 \right)^{-1} \mathbf{w}.$$

Moreover, because $(\tilde{K}^1, \tilde{K}^2)$ and $(\tilde{\mathbf{k}}^1, \tilde{\mathbf{k}}^2)$ exist and are uniquely defined, and that the opinions converge under the affine strategies given by (41), the value functions are well defined, and there exists a unique MPE of this game.

Step 3: The limit opinion $\tilde{\mathbf{x}}_\infty$ features consensus in the unique MPE of this game. Substitute expressions (44) and (45) into the expression for $\tilde{\mathbf{x}}_\infty$, and we have:

$$\begin{aligned} & (\alpha_1^2 \tilde{K}^1 + \alpha_2^2 \tilde{K}^2)(Y^2 - Z_1 Z_2) \tilde{\mathbf{x}}_\infty \\ &= \alpha_1^2 \left(I - \delta \Gamma D(\delta \alpha_2^2 \tilde{K}^2 + cI) \right) \tilde{\mathbf{b}}^1 + \alpha_2^2 \left(I - \delta \Gamma D(\delta \alpha_1^2 \tilde{K}^1 + cI) \right) \tilde{\mathbf{b}}^2. \end{aligned} \quad (46)$$

Clearly, in the limit, $\tilde{\mathbf{x}}_\infty$ is zero in all dimensions $j \neq 1$ because $\tilde{\mathbf{b}}_j^1 = 0$ and $\tilde{\mathbf{b}}_j^2 = 0$. The first dimension puts more weight on $\tilde{\mathbf{b}}^1$ than $\tilde{\mathbf{b}}^2$ because $\alpha_1^2 > \alpha_2^2$ and $\alpha_1^2 \tilde{K}_{11}^1 > \alpha_2^2 \tilde{K}_{11}^2$. Finally, as $\tilde{\mathbf{x}}_\infty$ is 0 in all dimensions $j \neq 1$, the unprojected opinions $\mathbf{x}_\infty = U \tilde{\mathbf{x}}_\infty$ have all equal entries. Because the MPE is unique, it must feature consensus in the long run. ||

Proof of Proposition 8: We consider two normalizations of the model, both of which are without loss. First, we assume $b^1 > b^2$ and $\alpha_1^2 b^1 + \alpha_2^2 b^2 = 0$. Notice that we can add a constant to all opinions and to both influencers' agendas and the problem remains the same, because the influencers care only about the difference between opinions and their agendas. Thus, we can assume $\alpha_1^2 b^1 + \alpha_2^2 b^2 = 0$ by choosing the proper constant. Then, $b^1 > 0$ and

$b^2 < 0$, and $b^1 < |b^2|$ because $\alpha_1 > \alpha_2$. Second, we assume $\mathbf{v}_j \cdot \mathbf{1} \geq 0$ for all j where $\mathbf{v}_j$ is the j th column of V . Otherwise, we can simply replace $\mathbf{v}_j$ with $-\mathbf{v}_j$ to make it satisfy the assumption, since both are the unit eigenvectors of the j th eigenvalue.

We first characterize the limit projected opinions in dimensions $j \neq 1$ in Step 1 because they are solutions to a myopic problem. We then study the projected opinion evolution in the first dimension and prove disagreement in the limit in Step 2.

Step 1: Recall that $A = USV'$, where S is the singular value matrix with diagonal entries $\sigma_1 = \sqrt{N(a_1^2 + a_2^2 + \dots + a_N^2)}$ and $\sigma_j = 0$ for all $j \neq 1$. We use $\mathbf{u}_j$ and $\mathbf{v}_j$ for the j th column of U and V ; similarly, we use $[\Gamma]_j$ for the j th column of any matrix Γ . In the singular decomposition, $\mathbf{u}_1 = \frac{1}{\sqrt{N}}\mathbf{1}'$ and $\mathbf{v}_1 = (a_1, a_2, \dots, a_N)' / \sqrt{a_1^2 + a_2^2 + \dots + a_N^2}$. In the transformed problem, $\tilde{\mathbf{x}}_{t+1} = V'US\tilde{\mathbf{x}}_t + \sum \alpha^m \tilde{\mathbf{r}}_t^m$, where $\tilde{\mathbf{x}}_t = V'\mathbf{x}_t$ and $\tilde{\mathbf{r}}_t = V'\mathbf{r}_t$. Notice that $[V'US]_j = \mathbf{0}$ for all $j \neq 1$, moreover, $[V'US]_{11} = 1$. Therefore only $\tilde{x}_t^1$ enters into the opinion $\tilde{\mathbf{x}}_{t+1}$, that is, the opinion evolution depends only on $\tilde{x}_t^1$ for all t . As a result, in any dimension $j \neq 1$, each influencer chooses $\tilde{r}_t^{m,j}$ to maximize her myopic payoff. Specifically,

$$c\tilde{r}_t^{m,j} = -\delta\alpha_m \left(\tilde{x}_t^j - \tilde{b}_j^m \right).$$

Solve for these messages together, we have

$$(c + \delta\alpha_1^2 + \delta\alpha_2^2)\tilde{r}_t^{m,j} = -\delta\alpha_m \left(([SU'V]_j)' \tilde{\mathbf{x}}_t - \delta\alpha_m^2 (\tilde{b}_j^m - \tilde{b}_j^{m'})/c - \tilde{b}_j^m \right).$$

Let $\tau = (c + \delta\alpha_1^2 + \delta\alpha_2^2)^{-1}$. Assume the limit opinion $\tilde{x}_\infty^1$ exists (we will show it holds in Step 2), then in each dimension $j \neq 1$, the limit opinion becomes:

$$\tilde{x}_\infty^j = c\tau ([SU'V]_j)' \tilde{\mathbf{x}}_\infty + \delta\tau (\alpha_1^2 \mathbf{v}_j' b^1 \mathbf{1} + \alpha_2^2 \mathbf{v}_j' b^2 \mathbf{1})$$

Thus, the limit opinion in the original problem is $\mathbf{x}_\infty = V\tilde{\mathbf{x}}_\infty$, which is equal to

$$V \left(c\tau V'US\tilde{x}_\infty^1 \mathbf{1} + \delta\tau (\alpha_1^2 V' b^1 \mathbf{1} + \alpha_2^2 V' b^2 \mathbf{1}) \right) + V \begin{bmatrix} \tilde{x}_\infty^1 - c\tau \tilde{x}_\infty^1 - \frac{\delta\tau(\alpha_1^2 b^1 + \alpha_2^2 b^2)}{\sqrt{a_1^2 + a_2^2 + \dots + a_N^2}} \\ 0 \\ \vdots \\ 0 \end{bmatrix}. \quad (47)$$

As $V'V = I$, the first part is $c\tau US\tilde{x}_\infty^1 \mathbf{1} + \delta\tau (\alpha_1^2 b^1 \mathbf{1} + \alpha_2^2 b^2 \mathbf{1}) = c\tau US\tilde{x}_\infty^1 \mathbf{1}$ given our first normalization. As $\mathbf{u}_1 = \frac{1}{\sqrt{N}}\mathbf{1}'$ and $S_{ij} = 0$ except when $i = j = 1$, we can show the first

part is a constant vector. So to make the limit opinion a consensus, we need the second part to be a constant vector. Recall that $\mathbf{v}_1$ has different entries, for consensus, we need

$$\tilde{x}_\infty^1 - c\tau\tilde{x}_\infty^1 - \frac{\delta\tau(\alpha_1^2b^1 + \alpha_2^2b^2)}{\sqrt{a_1^2 + a_2^2 + \dots + a_N^2}} = 0, \text{ or } (\alpha_1^2 + \alpha_2^2)\tilde{x}_\infty^1 = \frac{\alpha_1^2b^1 + \alpha_2^2b^2}{\sqrt{a_1^2 + a_2^2 + \dots + a_N^2}};$$

the last term is zero due to the first normalization. We will show in Step 2 that $\tilde{x}_\infty^1 > 0$, and thus consensus is impossible.

Step 2: Characterize limit consensus $\mathbf{x}_\infty$. Recall the value function for influencer m is

$$\begin{aligned} \tilde{v}_{t-1}^m(\tilde{\mathbf{x}}_{t-1}) &= -(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m) - c(\tilde{\mathbf{r}}_{t-1}^m)'(\tilde{\mathbf{r}}_{t-1}^m) \\ &\quad - \delta \left(V'US\tilde{\mathbf{x}}_{t-1} + \sum_{l=1}^2 \alpha_l \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}^m \right)' \tilde{K}^m \left(V'US\tilde{\mathbf{x}}_{t-1} + \sum_{l=1}^2 \alpha_l \tilde{\mathbf{r}}_{t-1}^l - \tilde{\mathbf{k}}^m \right) - \delta \tilde{\kappa}^m. \end{aligned}$$

We first examine the form of Ricatti matrix $\tilde{K}^m$. Since $V'US$ is a matrix with only the first column being non-zero, $V'US\tilde{\mathbf{x}}_{t-1}$ contains only $\tilde{x}_{t-1}^1$. As a result, $\tilde{x}_{t-1}^j$ ($j \neq 1$) enters the value function only in its first component $(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m)'(\tilde{\mathbf{x}}_{t-1} - \tilde{\mathbf{b}}^m)$. Thus, the Ricatti matrix $\tilde{K}^m$ is diagonal and all diagonal entries are 1 except for $\tilde{K}_{11}^m$. We can then apply similar analysis as in Proposition 7 where A is symmetric (and the Ricatti matrix is diagonal):

$$\begin{aligned} \tilde{K}^1 &= I + c\delta(V'US)' \Gamma \tilde{K}^1 (\delta\alpha_1^2 \tilde{K}^1 + cI) \Gamma V'US; \\ \tilde{K}^2 &= I + c\delta(V'US)' \Gamma \tilde{K}^2 (\delta\alpha_2^2 \tilde{K}^2 + cI) \Gamma V'US. \end{aligned}$$

The only difference from the Ricatti equations in Proposition 7 is that $V'US$ is not a diagonal matrix. Instead, we have

$$\tilde{K}_{11}^m = 1 + \frac{c\delta(\delta\alpha_m^2 + c)}{(c + \delta\alpha_1^2 + \delta\alpha_2^2)^2} \sum_2^N ((V'US)_{1j})^2 + \frac{c\delta\tilde{K}_{11}^m(\delta\alpha_m^2 \tilde{K}_{11}^m + c)}{(c + \delta\alpha_1^2 \tilde{K}_{11}^1 + \delta\alpha_2^2 \tilde{K}_{11}^2)^2}.$$

The second term on the right hand side is new but it is independent of $\tilde{K}_{11}^m$. We can follow similar steps as those in Step 1 of Proposition 7 to show solutions exist and $\tilde{K}_{11}^1 > \tilde{K}_{11}^2$.

Next, let $\mathbf{w} = \alpha_1^2 \tilde{K}^1 \tilde{\mathbf{k}}^1 + \alpha_2^2 \tilde{K}^2 \tilde{\mathbf{k}}^2$, we have the evolution of $\tilde{\mathbf{k}}^m$, which only depends on $\tilde{K}^m$ and thus is well-defined:

$$\tilde{K}^1 \tilde{\mathbf{k}}^1 = \tilde{\mathbf{b}}^1 - \delta(V'US)' \Gamma \tilde{K}^1 (\delta\alpha_1^2 \tilde{K}^1 + cI) (\delta\Gamma \mathbf{w} - \tilde{\mathbf{k}}^1); \quad (48)$$

$$\tilde{K}^2 \tilde{\mathbf{k}}^2 = \tilde{\mathbf{b}}^2 - \delta(V'US)' \Gamma \tilde{K}^2 (\delta\alpha_2^2 \tilde{K}^2 + cI) (\delta\Gamma \mathbf{w} - \tilde{\mathbf{k}}^2). \quad (49)$$

Unlike Observation 2 in the symmetric A , $\tilde{\mathbf{b}}^m$ is not necessarily zero in dimension $j \neq 1$ because $\mathbf{v}_1$ is not collinear with $\mathbf{b}^m$. Because $(V'US)'$ is a matrix with only the first row being non-zero; in dimension $j \neq 1$ of (48) and (49), it becomes $\tilde{k}_j^m = \tilde{b}_j^m$ (recall $\tilde{K}_{jj}^m = 1$). Thus, $w_j = 0$ for all $j \neq 1$ under the first normalization.

We claim that $w_1 \neq 0$. Suppose $w_1 = 0$, we want to find a contradiction. The second normalization implies that every entry in the first row of $(V'US)'$ must be non-negative because $\mathbf{u}_1$ is proportional to $\mathbf{1}$. Recall that $b^1 > 0$ and $b^2 < 0$; the second normalization also implies that $\tilde{\mathbf{b}}^1 \geq 0$ and $\tilde{\mathbf{b}}^2 \leq 0$ with strict inequality holds in the first entry. Then, $\tilde{k}_j^1 \geq 0$ and $\tilde{k}_j^2 \leq 0$ for all $j \neq 1$. Inserting the hypothesis $w_1 = 0$ into (48) and (49), we can show $\tilde{k}_1^1 > 0$ and $\tilde{k}_1^2 < 0$. Next, we want to show $(\alpha_1^2 \tilde{K}^1) \alpha_1^2 \tilde{K}^1 \tilde{\mathbf{k}}^1 + (\alpha_2^2 \tilde{K}^2) \alpha_2^2 \tilde{K}^2 \tilde{\mathbf{k}}^2 \geq 0$ with strict inequality holds in the first entry. In the first entry,

$$\begin{aligned} & (\alpha_1^2 \tilde{K}_{11}^1) \alpha_1^2 \tilde{K}_{11}^1 \tilde{k}_1^1 + (\alpha_2^2 \tilde{K}_{11}^2) \alpha_2^2 \tilde{K}_{11}^2 \tilde{k}_1^2 \\ &= (\alpha_1^2 \tilde{K}_{11}^1) \alpha_1^2 \tilde{K}_{11}^1 \tilde{k}_1^1 + (\alpha_2^2 \tilde{K}_{11}^2) (-\alpha_1^2 \tilde{K}_{11}^1 \tilde{k}_1^1) \\ &= (\alpha_1^2 \tilde{K}_{11}^1 - \alpha_2^2 \tilde{K}_{11}^2) \alpha_1^2 \tilde{K}_{11}^1 \tilde{k}_1^1 > 0 \end{aligned}$$

The first equality uses the fact that $w_1 = \alpha_1^2 \tilde{K}_{11}^1 \tilde{k}_1^1 + \alpha_2^2 \tilde{K}_{11}^2 \tilde{k}_1^2 = 0$, and the last inequality is due to the fact that $\alpha_1 > \alpha_2$, $\tilde{K}_{11}^1 > \tilde{K}_{11}^2 > 1$ and $\tilde{k}_1^1 > 0$. Results in all other entries can be proved in a similar way. Finally, multiply (48) with α_1^2 and multiply (49) with α_2^2 and sum them up, we have

$$(V'US)' \left((\alpha_1^2 \tilde{K}^1) \alpha_1^2 \tilde{K}^1 \tilde{\mathbf{k}}^1 + (\alpha_2^2 \tilde{K}^2) \alpha_2^2 \tilde{K}^2 \tilde{\mathbf{k}}^2 \right) = \mathbf{0}.$$

This equation is a contradiction because we have shown that every entry in the first row of $(V'US)'$ must be non-negative with $(V'US)_{11} = 1$ and $(\alpha_1^2 \tilde{K}^1) \alpha_1^2 \tilde{K}^1 \tilde{\mathbf{k}}^1 + (\alpha_2^2 \tilde{K}^2) \alpha_2^2 \tilde{K}^2 \tilde{\mathbf{k}}^2 \geq 0$ with strict inequality holds in the first entry. Thus, $w_1 \neq 0$. Lastly, we show that $w_1 > 0$. Notice that $w_1 > 0$ when $a_1 = \dots = a_N$ from Proposition 7 and w_1 is a continuous function of $(a_1, \dots, a_N)$. If for some $(a_1, \dots, a_N)$, $w_1 < 0$, then by the intermediate value theorem, $w_1 = 0$ for some other $(a'_1, \dots, a'_N)$, which contradicts our earlier claim. Thus, $w_1 > 0$.

Again, using the proof of Proposition 7, we have

$$\tilde{\mathbf{x}}_t = c\Gamma V'US \tilde{\mathbf{x}}_{t-1} + \delta\Gamma \mathbf{w}, \text{ and in the limit, } (I - c\Gamma V'US) \tilde{\mathbf{x}}_\infty = \delta\Gamma \mathbf{w}.$$

Thus, $\tilde{x}_\infty^1 > 0$ since it is proportional to w_1 . Using the formula (47) of the limit opinions $\mathbf{x}_\infty$, the first part is $c\tau US \tilde{x}_\infty^1 \mathbf{1} > 0$. The first element in the vector of the second part becomes

$(1 - c\tau)\tilde{x}_\infty^1 > 0$. Thus, all limit opinions are positive. Moreover, the opinion leaders' limit opinions are closer to influencer 1's agenda and there is total disagreement in the limit: $x_\infty^1 > \dots > x_\infty^N$ as $a_1 > \dots > a_N$. $\parallel$

Proof of Corollary 2: This model is very similar to that of Proposition 4. Thus, we sketch out only the parts that are different. The derivation of optimal strategies and the Ricatti equation are the same (except for the new term W). The Ricatti equation here is

$$K^* = I + \delta A' (K^* - \delta K^* W (\delta W' K^* W + cI)^{-1} W' K^*) A.$$

A similar inductive argument can show that $K^* = \lim_{t \rightarrow \infty} K_{T-t}$, which is pd and decreasing. We now show the opinions converge under the optimal strategy.

Because A is upper triangular and W has a special form $W = (I, 0)'$, K^* has an interesting feature: it is a block matrix of the form

$$K^* = \begin{bmatrix} K_1 & K_2 \\ K_2' & K_3 \end{bmatrix}.$$

Submatrix K_1 summarizes the future disutility from agents' opinion deviation from the influencer, and K_2 summarizes her future disutility because the agents are influenced by the bots' agendas. Submatrix K_3 summarizes the disutility due to the bots' deviation from her agenda, which is a constant term and does not matter to her strategy. We now show that K^* can be decomposed in that K_1 is a function of itself, while K_2 depends on both K_1 and K_2 . Simple calculations can show

$$\begin{aligned} I - \delta W (\delta W' K^* W + cI)^{-1} W' K^* &= I - \delta W (\delta K_1 + cI)^{-1} (K_1, K_2) \\ &= I - \delta \begin{bmatrix} (\delta K_1 + cI)^{-1} K_1 & (\delta K_1 + cI)^{-1} K_2 \\ 0 & 0 \end{bmatrix} \\ &= \begin{bmatrix} I - (\delta K_1 + cI)^{-1} \delta K_1 & -(\delta K_1 + cI)^{-1} \delta K_2 \\ 0 & I \end{bmatrix} \end{aligned}$$

Substitute into the Ricatti equation and we have:

$$K_1 = I + c\delta(A^n)' K_1 (\delta K_1 + cI)^{-1} A^n. \quad (50)$$

$$K_2 = c\delta(A^n)' (\delta K_1 + cI)^{-1} (K_1 A^z + K_2). \quad (51)$$

From equation (50) we can solve for K_1 , and then K_2 can be solved. The influencer's optimal strategy $\mathbf{r}_t = -(\delta W' K^* W + cI)^{-1} \delta W' K^* A \chi_t$. Therefore opinions evolve according to

$$\chi_t = ((I - \delta W (\delta W' K^* W + cI)^{-1} W' K^*) A)^t \chi_0.$$

We can use the expression above and get:

$$\lim_{t \rightarrow \infty} \mathbf{x}_t = \lim_{t \rightarrow \infty} (c(\delta K_1 + cI)^{-1} A^n)^t \mathbf{x}_0 + (I - c(\delta K_1 + cI)^{-1} A^n)^{-1} (\delta K_1 + cI)^{-1} (cA^z - \delta K_2) \mathbf{z}.$$

It is easy to see that $(c(\delta K_1 + cI)^{-1} A^n)^t \rightarrow \mathbf{0}$ as $t \rightarrow \infty$. Therefore

$$\mathbf{x}_\infty = (\delta K_1 + cI - cA^n)^{-1} (cA^z - \delta K_2) \mathbf{z}.$$

Note that K_1 and K_2 do not depend on $\mathbf{z}$, $(\delta K_1 + cI - cA^n)$ is positive definite. Moreover, using expression (51), we can show $\mathbf{x}_\infty = \Gamma A^z \mathbf{z}$, where Γ is a positive definite matrix and thus of full rank, and Γ is independent of $\mathbf{z}$. If $\Gamma A^z \mathbf{z} = \gamma \mathbf{1}$ for all $\mathbf{z}$, where γ is a scalar, Γ must feature linearly dependent rows. But this is impossible since Γ is positive definite. Thus for a generic $\mathbf{z}$, $\mathbf{x}_\infty$ is not proportional to $\mathbf{1}$, that is, there is no consensus. ||

B Multiple Strategic Influencers: General Model

We consider a general model of multiple strategic influencers in which we allow each strategic influencer to send a varying number of messages each period. They also have more general access to the agents and face a more general cost of sending messages. As in the text, suppose that there are M strategic influencers with agenda $b^1, \dots, b^M$ respectively, with m being a generic influencer. Also, let m 's message be $\mathbf{r}_t^m \in \mathbb{R}^{l_m}$. That is, m sends $l_m \geq 1$ messages to the agents in each period, which can be considered as her channels of communication. Influencer m pays the following cost of sending messages: let ρ^m be m 's zero-cost message, any departure from ρ^m entails a quadratic cost. If $\rho^m = 0$ as in the text, it can be thought of as a financial cost of sending messages. If $\rho^m \neq 0$, it can represent a reputational cost, that is, influencer m is penalized if she sends a message that is different from her known position ρ^m .

We look at the finite horizon problem before commenting on the infinite horizon problem. Each influencer m maximizes her total discounted payoff by choosing messages $\mathbf{r}_t^m$, for $t = 0, \dots, T - 1$. Let $\mathbf{b}^m = b^m \mathbf{1}_N$ and $\boldsymbol{\rho}^m = \rho^m \mathbf{1}_{l_m}$. Specifically, influencer m 's stage t payoff is

$$u_t^m(\mathbf{x}_t) = -(\mathbf{x}_t - \mathbf{b}^m)' Q^m (\mathbf{x}_t - \mathbf{b}^m) - (\mathbf{r}_t^m - \boldsymbol{\rho}^m)' R^m (\mathbf{r}_t^m - \boldsymbol{\rho}^m),$$

where Q^m captures how much weights influencer m assigns to the agents and R^m captures the cost if m 's message differs from her known position. We assume both Q^m and R^m are symmetric and positive definite.

Each agent listens to the influencers' messages according to W^m , a $N \times l_m$ matrix. That is, he can listen from multiple channels of influencer m 's message. Agents update their opinions according to:

$$\mathbf{x}_{t+1} = A\mathbf{x}_t + \sum_1^M W^m \mathbf{r}_t^m.$$

As before, for each agent, the weights he assigns to all other agents' opinions are summarized by A . We allow for both possibilities for generality. First, A is stochastic and the influencers are outside the network. Second, every row sum of A and $\sum_1^M W^m$ is 1, and thus the influencers are part of the network.

Because each influencer cares about how far the agents' opinions are from her agenda, we can ease exposition by letting $\chi_t^m = \mathbf{x}_t - \mathbf{b}^m \in \mathbb{R}^N$, which measures the distance between $\mathbf{x}_t$ and m 's agenda. Similarly, let $\gamma_t^m = \mathbf{r}_t^m - \boldsymbol{\rho}^m \in \mathbb{R}^{l_m}$, which measures the distance between $\mathbf{r}_t^m$ and m 's known position. In particular, χ_t^m is the new state variable for influencer m , who chooses γ_t^m instead of $\mathbf{r}_t^m$ in the following analysis. The opinion updating process above then becomes, for all $t \leq T - 1$ and m ,

$$\chi_{t+1}^m = A\chi_t^m + \sum_h W^h \gamma_t^h + \mathbf{c}^m, \quad (52)$$

where the time independent constant $\mathbf{c}^m = \sum_m W^m \boldsymbol{\rho}^m + (A - I_N)\mathbf{b}^m \in \mathbb{R}^N$.

We assume that all influencers know the parameters of this game, namely, matrix A and W^m for all m , the payoff parameters Q^m and R^m and the horizon T . Thus this is a dynamic game of complete and imperfect information.³⁹ We also assume that $\delta < 1$ for all influencers so that the value function is bounded when $T = \infty$.

As before, we start with influencer m 's problem at $T - 1$. The value function has both quadratic and linear terms. We still use K_t^m and $\mathbf{k}_t^m, \kappa_t^m$ to represent these terms, but the use of $\mathbf{k}_t^m$ is slightly different from the text because this notation is simpler in the more general model. At T , the value function for m is

$$v_T^m(\chi_T^m) = -(\chi_T^m)' K_T^m (\chi_T^m) - (\mathbf{k}_T^m)' (\chi_T^m) - \kappa_T^m,$$

with terminal values $K_T^m = Q^m$, $\mathbf{k}_T^m = \mathbf{0} \in \mathbb{R}^N$ and $\kappa_T^m = 0$. We proceed iteratively and

³⁹As in the benchmark case, adding small amount of normally distributed iid noise does not affect the results qualitatively.

show if the value function at any $1 \leq t \leq T$ is of the form

$$v_t^m(\mathbf{x}_t^m) = -(\mathbf{x}_t^m)' K_t^m \mathbf{x}_t^m - (\mathbf{k}_t^m)' \mathbf{x}_t^m - \kappa_t^m,$$

the value function at $t - 1$ also has this form. We now derive the optimal strategy of the influencers. To begin with, the value function at $t - 1$ is:

$$v_{t-1}^m(\mathbf{x}_{t-1}^m) = \max_{\gamma_{t-1}^m} -(\mathbf{x}_{t-1}^m)' Q^m \mathbf{x}_{t-1}^m - (\gamma_{t-1}^m)' R^m \gamma_{t-1}^m + \delta v_t^m(\mathbf{x}_t^m).$$

Use expression (52) and differentiate with respect to γ_{t-1}^m , and we have the first order condition for each m :

$$\begin{aligned} & - \left((W^m)' K_t^m W^m + \frac{R^m}{\delta} \right) \gamma_{t-1}^m \\ & = (W^m)' K_t^m A \mathbf{x}_{t-1}^m + (W^m)' K_t^m \left(\sum_{h \neq m} W^h \gamma_{t-1}^h \right) + (W^m)' K_t^m \mathbf{c}^m + \frac{1}{2} (W^m)' \mathbf{k}_t^m. \end{aligned} \tag{53}$$

Influencer m 's strategy in $t - 1$ is an affine function of the current opinions, messages of the other influencers, and a constant term depending on her agenda. Her strategy (net of her known position) goes in the opposite direction of $\mathbf{x}_{t-1}^m$: if the agents' opinions are higher than m 's agenda, she aims to reduce them by sending a more negative message (relative to her agenda), and vice versa. It also decreases in the weighted sum of the other influencers messages (net of their agendas).

Next, we can find the stage NE for period $t - 1$ by solving the M first order conditions together. Reorganize the best response functions, and we have for all m :

$$\begin{aligned} & \left((W^m)' K_t^m W^m + \frac{R^m}{\delta} \right) \gamma_{t-1}^m + (W^m)' K_t^m \left(\sum_{h \neq m} W^h \gamma_{t-1}^h \right) \\ & = - (W^m)' K_t^m A \mathbf{x}_{t-1}^m - (W^m)' K_t^m \mathbf{c}^m - \frac{1}{2} (W^m)' \mathbf{k}_t^m. \end{aligned}$$

The left hand side is the sum of all messages weighted by each m 's influence on the network W^m , her cost, and her Ricatti matrix K_t^m . Let the coefficient matrix of the influencers' messages on the left hand side above be:

$$M_{t-1} = \begin{bmatrix} (W^1)' K_t^1 W^1 + R^1/\delta & (W^1)' K_t^1 W^2 & \dots & (W^1)' K_t^1 W^M \\ (W^2)' K_t^2 W^1 & (W^2)' K_t^2 W^2 + R^2/\delta & \dots & (W^2)' K_t^2 W^M \\ \vdots & \vdots & \ddots & \vdots \\ (W^M)' K_t^M W^1 & (W^M)' K_t^M W^2 & \dots & (W^M)' K_t^M W^M + R^M/\delta \end{bmatrix}.$$

Here M_{t-1} is an $(\sum_m l_m \times \sum_m l_m)$ matrix, where the first $(1_m \times \sum_m 1_m)$ submatrix concerns influencer 1's equilibrium strategy at $t-1$, and the m -th submatrix concerns agent m . One necessary condition for M_{t-1} to be invertible is that each R^m is positive definite. But because in general the Ricatti matrix K_t^m differs by m and it depends on all other K_t^h for $h \neq m$, it is difficult to give a sufficient condition for M_{t-1} to be non-singular. We proceed by assuming M_{t-1} is invertible.

Note that $\chi_t^h = \chi_{t-1}^m + \mathbf{b}^m - \mathbf{b}^h$ by definition. Then the influencers' equilibrium messages at $t-1$ are:

$$\gamma_{t-1}^m = -E_{t-1}^m \chi_{t-1}^m - F_{t-1}^m, \text{ where } E_{t-1}^m = [M_{t-1}^{-1} H_{t-1}]^m, \quad F_{t-1}^m = [M_{t-1}^{-1} C_{t-1}]^m. \quad (54)$$

The coefficient matrix H_{t-1} on the opinions χ_{t-1}^m and the residual term C_{t-1}^m are respectively:

$$H_{t-1} = \begin{bmatrix} (W^1)' K_t^1 A \\ \vdots \\ (W^M)' K_t^M A \end{bmatrix}, \text{ and } C_{t-1}^m = \begin{bmatrix} (W^1)' K_t^1 A (\mathbf{b}^m - \mathbf{b}^1) + (W^1)' K_t^1 \mathbf{c}^1 + \frac{1}{2} (W^1)' \mathbf{k}_t^1 \\ \vdots \\ (W^M)' K_t^M A (\mathbf{b}^m - \mathbf{b}^M) + (W^M)' K_t^M \mathbf{c}^M + \frac{1}{2} (W^M)' \mathbf{k}_t^M \end{bmatrix}.$$

Similarly, H_{t-1} and C_{t-1}^m are $(\sum_m l_m \times N)$ and $(\sum_m l_m \times 1)$ matrix respectively. In particular, $[M_{t-1}^{-1} H_{t-1}]^m$ refers to the m -th submatrix of the product, and $[M_{t-1}^{-1} C_{t-1}^l]^l$ refers to the l -th submatrix of the product.

Next, derive $v_{t-1}^m(\chi_{t-1}^m)$ by substituting in the influencers' equilibrium strategies at $t-1$. Using $\chi_t^h = \chi_{t-1}^m + \mathbf{b}^m - \mathbf{b}^h$ again, the opinions evolve from agent m 's perspective as follows:

$$\chi_t^m = \left(A - \sum_{h=1}^M W^h E_{t-1}^h \right) \chi_{t-1}^m - \sum_{h=1}^M W^h (E_{t-1}^h (\mathbf{b}^m - \mathbf{b}^h) + F_{t-1}^h) + \mathbf{c}^m.$$

To simplify notations, we can let $G_{t-1} = A - \sum_{h=1}^M W^h E_{t-1}^h$, which determines the period-to-period evolution of the agent's opinions, and is common to all influencers. Also, we let $\mathbf{g}_{t-1}^m = -\sum_{h=1}^M W^h (E_{t-1}^h (\mathbf{b}^m - \mathbf{b}^h) + F_{t-1}^h) + \mathbf{c}^m$. Then the opinions evolve given the influencers' optimal strategy as:

$$\chi_t^m = G_{t-1} \chi_{t-1}^m + \mathbf{g}_{t-1}^m.$$

We can now expand v_{t-1}^m and equate the quadratic terms and get:

$$K_{t-1}^m = Q^m + (E_{t-1}^m)' R^m E_{t-1}^m + \delta G_{t-1}' K_t^m G_{t-1}. \quad (55)$$

Note that if M_{t-1} is invertible for all t , then from the above expression that each individual K_t^m is symmetric and positive definite by an induction argument similar to that in the main model. But in general, the M iterative solutions for K_{t-1}^m depend on all K_t^m in a nonlinear way through M_{t-1}^{-1} and H_{t-1} . So these M equations need to be solved simultaneously. But they only depend on K_t^m , not $\mathbf{k}_t^m$.

Next, we equate the linear term and find:

$$(\mathbf{k}_{t-1}^m)' = 2\delta(\mathbf{g}_{t-1}^m)'K_t^m G_{t-1} + \delta(\mathbf{k}_t^m)'G_{t-1} + 2(F_{t-1}^m)'R^m E_{t-1}^m. \quad (56)$$

The constant term is

$$\kappa_{t-1}^m = \delta((\mathbf{g}_{t-1}^m)'K_t^m \mathbf{g}_{t-1}^m + (\mathbf{k}_t^m)' \mathbf{g}_{t-1}^m - \mathbf{k}_t^m) + (F_{t-1}^m)'R^m F_{t-1}^m. \quad (57)$$

By backward induction and equations (55), (56) and (57), we can iteratively find the optimal strategy $\gamma^m = \{\gamma_0^m, \dots, \gamma_{T-1}^m\}$ with a value function $v_0^m(\mathbf{x}_0^m)$ for all m .

If each K_t^m has a limit $\lim_{t \rightarrow \infty} K_{T-t}^m = K^m$, then taking limit of equations (55), (56) and (57) on both sides, we obtain the infinite horizon solution. This is because all the matrices defined above are continuous in K_t^m .