
© COPYRIGHTED BY

Simon Stolarczyk

December 2017

DECISION MAKING IN SOCIAL NETWORKS

A Dissertation

Presented to

the Faculty of the Department of Mathematics

University of Houston

In Partial Fulfillment

of the Requirements for the Degree

Doctor of Philosophy

By

Simon Stolarczyk

December 2017

DECISION MAKING IN SOCIAL NETWORKS

Simon Stolarczyk

APPROVED:

Dr. Krešimir Josić (Committee Chair)
Department of Mathematics, University of Houston

Dr. Kevin E. Bassler
Department of Physics, University of Houston

Dr. Zachary Kilpatrick
Department of Appl. Math, University of CO

Dr. Ilya Timofeyev
Department of Mathematics, University of Houston

Dr. Andrew Török
Department of Mathematics, University of Houston

Dean, College of Natural Sciences and Mathematics

Acknowledgements

It is said that work in math and science is built by standing on the shoulders of giants, but to do it, we must also stand on the shoulders of those around us. I would like to thank those people here.

First, I want to thank my advisor, Dr. Krešimir Josić. His guidance has made me a better writer and a better scientist. Without his insight and attention to detail, none of the following would have been possible.

I would also like to give thanks to my committee members Dr. Andrew Török and Dr. Ilya Timofeyev, but especially to Dr. Kevin E. Bassler and Dr. Zachary Kilpatrick who worked with me to produce results on the first and second topics of my research, respectively.

I am grateful to have had many great colleagues, from those who took classes with me my first years to those in my research group in my last years. I am especially glad to have spent time with Adrian, Carlos, and Daniel who all matched my insight with better insight, and sarcasm with *much* better sarcasm.

Next I need to thank all of those friends who stuck around with me even while I was sucked into the hole of graduate school. In particular, without Justin, Kristen, and Regan my years would have been way less fun.

I am hugely grateful to Dave and Valerie Throgmorton, whose bottomless generosity provided me with housing and sustenance and rides to campus. With their support, I most likely would not have made it this far.

I am thankful to my Mom for raising me, and making me hard-headed enough

to do this kind of work. I am also thankful for Nanny, who provided the foundation for my intellectual passion and is the reason I pursued an education. Doing all those puzzles together doubtlessly made me the problem-solver I am today.

Finally, I am thankful to Kaitlin (even thankful enough to not use one her nicknames). Her love and support (including yelling at me to finish my work) were a constant animus.

DECISION MAKING IN SOCIAL NETWORKS

An Abstract of a Dissertation
Presented to
the Faculty of the Department of Mathematics
University of Houston

In Partial Fulfillment
of the Requirements for the Degree
Doctor of Philosophy

By
Simon Stolarczyk
December 2017

Abstract

We investigate how the network topology of social networks impacts decision making. First we look at sequential models of decision making with feedforward network topologies. We see how rational agent incorporate knowledge of the network topology in order to make an optimal estimate of an unknown parameter. We give a condition for making this optimal estimate in terms the row space of a matrix which encodes the network topology. We then show what this condition means for infinitely large networks. Then we extend a model of evidence accumulation for a two-alternative task to general networks where agents are allowed to communicate decisions. We detail the process rational agents must undergo and give detailed computations of the process for basic network structures.

Contents

1	Introduction	1
1.1	Decision Making Agents	5
1.2	Underlying Assumptions and Aumann's Theorem	7
1.3	Herding	8
1.4	Probabilistic Network Herding	10
1.5	Sequential Parameter Estimation in Feedforward Networks	12
1.6	Evidence Accumulation in Networks	14
2	Feedforward Networks	16
2.1	Setup	17
2.2	The Model	20
2.3	Results	26
2.3.1	Graphical Conditions for Ideal Networks	27
2.4	Sufficient Conditions for Ideal Three-Layer Networks	34
2.5	Conclusion	39
3	Asymptotics for Feedforward Networks	42
3.1	Variance and Bias of the Final Estimate	42

CONTENTS

3.2	Inference in random feedforward networks	49
3.2.1	More than 3 Layers	53
3.2.2	Simulation Details	54
3.3	Conclusion	54
4	Evidence Accumulation on Networks	56
4.1	Single-Agent Setup	58
4.2	Two-Agent Setup	62
4.3	Thresholds and Decision Evidence	69
4.3.1	Decision Evidence	69
4.3.2	Non-decision evidence and Symmetry	72
4.4	Discrete Example	76
4.4.1	Setup	76
4.4.2	Non-Decision Evidence	77
4.4.3	Simulations	84
4.5	Continuum Limit	86
4.6	Conclusion	95
5	Bidirectional Coupling	96
5.1	Setup and Symmetric Boundaries	97
5.2	Equilibration Process	101
5.3	Recursive Process in Time	107
6	General Network Accumulation	110
6.1	Accumulation of Evidence on General Networks	111
6.1.1	Terminology and Notation	112
6.1.2	Marginalization	113
6.1.3	Pre-decision	114

CONTENTS

6.2	Three-Agent Networks	117
6.2.1	NS1: The Fully Connected Network	117
6.2.2	NS11: The 3-Agent Unidirectional Line	125
6.2.3	Other Networks	129
6.3	Large Cliques: Simulations and Asymptotics	131
6.3.1	Agreement Information	133
	Bibliography	141

List of Figures

2.1.1	Illustration of the general setup. Agents in the first layer (top layer in the figure) make measurements, x_1, x_2 , and x_3 , of a parameter s . In each layer agents make an estimate of this parameter, and communicate it to agents in the subsequent layer. Arrows indicate the direction in which information is propagated. We show that information about s degrades across layers in the network in panel (a), but not in the network in (b).	19
2.3.1	A W-motif spanning three layers.	28
2.4.1	Example of a two step network reduction. It is difficult to tell whether the network on top is ideal. However, after two steps of reduction, all first-layer agents in each of the five connected components have equal out-degree. The network is therefore ideal.	38
3.1.1	Example of a network with an inconsistent final estimate. The green and blue nodes represent agents in the first and second layer, respectively. Each second-layer agent receives input from the common, central agent and a distinct first-layer agent, and thus $L_2 = L_1 - 1$	43

3.2.1	The probability that a random, three-layer network is ideal for connection probabilities $p = 0.1$ (left), 0.5 (center), and 0.9 (right). In each panel, the different curves correspond to different, but fixed numbers of agents in the first layer. The number of agents in the second layer is varied. There is a sharp transition in the probability that a network is ideal when the number of agents in the second layer exceeds the number in the first. Simulation details can be found in Section 3.2.2.	52
3.2.2	The probability that a random, four-layer network is ideal for connection probabilities $p = 0.1$ (left), 0.5 (center), and 0.9 (right). Each curve corresponds to equal, fixed numbers of agents in the first two layers, with a changing number of agents in the third layer. Simulation details can be found in Section 3.2.2.	54
4.1.1	An example of the evolution of the log-likelihood ratio when the evidence distributions are normal with means ± 0.1 corresponding to $H = \pm 1$ and equal variance 1. The true state is $H = 1$, the thresholds are $\theta_- = -3$ and $\theta_+ = 3$, and the agent eventually accumulates enough evidence to make a correct decision.	62
4.2.1	A pair of agents with unidirectional coupling.	63
4.3.1	An example of evidence accumulation with two agents, and normal evidence distributions with means ± 0.1 corresponding to $H = \pm 1$ and equal variance 1. The true state is $H = 1$. Agent 1 makes the wrong decision, but this does not mislead the second agent.	75
4.4.1	An example lattice for a random walk where $\theta_- = -3\theta$ and $\theta_+ = 5\theta$ are the absorbing boundaries.	78
4.4.2	A simple example lattice with two transient states, two absorbing states, and transition probabilities indicated by arrows.	80
4.4.3	The value of the non-decision evidence that computes for $\theta_- = -1, \theta_+ = 2, p = \frac{e}{5}, q = \frac{1}{5}, s = 1 - p - q$. We see that the evidence quickly converges to $\frac{1}{2}$	83
4.4.4	Non-decision evidence saturates for different ratios of thresholds. We let $p = \frac{e}{5}, q = \frac{1}{5}, s = 1 - q - p$ so that $\log \frac{p}{q} = 1$ and $p + q + s = 1$	85

4.4.5	The amount of evidence obtained by agent 2 after observing a non-decision at time step n where n is the minimal number of steps needed to reach the closer, negative threshold. The parameter p defining the observational distribution is varied in the interval $p \in [0.5, 1)$ to obtain the different curves. For comparison, the amount of evidence obtained from a single A observation is given by the blue curve.	87
4.5.1	A simulation of the conditional belief distributions for $H = -1$ (left) and $H = 1$ (right) using Crank-Nicolson discretization. Each plot shows how the belief is distributed and evolves in time where brightness corresponds to probability. Here we let the boundaries be $\theta_- = -1$ and $\theta_+ = 3$. We used a drift rate $a = 1$, diffusion rate $b = 1$, space step size $dx = 0.01$, time step size $dt = 0.001$, and total time $T = 10$	93
4.5.2	Using the simulations from Figure 4.5.1, we compute the survival probabilities and resulting non-decision evidence and plot them on the same graph.	94
5.1.1	Two bidirectionally coupled agents.	97
6.2.1	We outline in pink the two networks whose behavior we analyze in detail. The other networks can be analyzed using a similar approach, or can be understood from the analysis of the two agent case (e.g. NS8 is essentially a two-agent network with an observer). We will refer to each network by the given label.	118
6.2.2	Agents on a unidirectional Line	125
6.3.1	The first decider, agent 1 is on top. The agreeing agents are on the right in the purple group and the disagreeing agents are on the left in the green group.	132
6.3.2	Simulations for a fully-connected network with 10 agents for $H = 1$. Top left: probability some agent decides (red) by time t and probability a given agent has positive belief (blue). Top right: the probability exactly k agents have positive belief. Bottom left: the expected amount of evidence gained after seeing the distribution of agreeing and disagreeing agents. Bottom right: how much evidence is gained per agreeing agent.	137

LIST OF FIGURES

- 6.3.3 The plots are the same as in Figure 6.3.2, here done for 100 agents.
The only difference is, in the top right, we plot the probability k
agents are correct for all k larger than half the total number of agents.
This gives us an idea of how likely the majority of agents agree. . . 138
- 6.3.4 The plots are the same as in Figure 6.3.2, here done for 1000 agents.
The only difference is, in the top right, we plot the probability k
agents are correct for all k larger than half the total number of agents.
This gives us an idea of how likely the majority of agents agree. . . 139
- 6.3.5 The expected first passage time for a given sized clique. Each point
was done by averaging the first exit times over 100 simulations. . . 140

List of Tables

Introduction

When making decisions in natural or economic situations, agents typically rely on the processing of information to choose whether and how to act. This information may come from observations of the environment, or it may be shared by other agents as they take actions or communicate information. Here we refer to anything capable of using information to decide on an action as an **agent**. This could be an animal in the wild or in a lab, a human deciding on whether to make a particular purchase, or a robot deciding where to move to based on a goal and the locations of its neighbors. We will frequently assume that the agents are rational: There is some measure of success or reward that depends on their actions, and they use all available information to maximize this reward [41]. For instance, agents trying to make an estimate of some real valued parameter, s , might want to make an estimate $\hat{s}$ which maximizes the expected value of

$$u(\hat{s}) = -(s - \hat{s})^2. \tag{1.0.1}$$

In many situations, an agent acting as part of a group will, on average, outperform an agent working alone. Agents in a network can share information, and can thus be expected to do better than any individual agent relying solely on information it has gathered on its own. Typically not all agents can interact with each other. This may be due to limited bandwidth or physical distances between agents. Such constraints are modeled by specifying which pairs of agents interact, and which do not. Formally, we identify agents with a set of vertices, V , of a graph (network). The set of all possible pairwise, directed interactions defines the set of edges, E , in this network [30] which we will also refer to as the **network topology**. We will always consider directed graphs, where edges have an orientation which indicate that one agent is a source of information that is communicated to another agent. In Chapters 4, 5, and 6 we allow return edges, so that two agents can provide information to each other. When two vertices are connected by an edge, we will informally say that the agents are neighbors, but we will formally define neighborhoods when we need to use them.

In the following we assume that agents can acquire two types of information: private and social. We refer to the information an agent receives from its neighbors as **social information**, and the information that is available through other sources, such as observations or measurements of environmental variables, as **private information**. Private information can be shared, and thus become social. An agent shares social information only with its neighbors, that is agents at the head of edges pointing from the agent. Thus social information transmitted by one agent

is not necessarily available to all agents in the network. However, social information can propagate through a network, and reach all agents eventually.

As each agent only has access to the social information communicated by its neighbors, different agents will typically have different information at any point in time, and may reach different decisions. Moreover, it is not clear that agents will make use of all information that is available to the totality of agents in the network, especially when they are not connected to all of the agents. To evaluate how well the agents are working together, one must compare their performance to that of a **global agent**: an idealized actor who has access to all information in the network and uses it to maximize its reward, or the probability of being correct. This global agent has access to the private information of all individuals and can use all this information unobstructed by the constraints of the network.

As an example, suppose that every agent has a private, i.i.d. signal about which of two choices is better. Assume that none of the signals are completely reliable, but are independent. In this case, it would be better to poll all agents about which choice they think is better and follow the majority than to select some agent at random and go with its choice. Indeed, an early version of the Law of Large numbers known as Condorcet's Jury Theorem states that the majority will choose the better option with probability 1, as the number of agents grows to infinity, if all agents are equivalent and every agent has a better than even chance of choosing the better option (see [17]). An agent who is able to poll all others would thus be able to make the best use of the information from this collective. In general, we will make no assumption that such an agent is part of the network,

but will use a fictitious global agent for comparison.

Thus the social information available to an agent is typically weaker than the information available to a global agent. Social information can be the result of an inference or a decision. In the case of an inference, the agents infer the state of some variable in the environment, and can communicate the inferred value (this can be a vector). Thus if an agent receives the inferred value it can learn something about the variable it does not directly observe. However, the agent may not learn as much as if it directly received all the information the communicating agent used to make the inference. In the case of a decision, the agents will make a choice between different options based on the information that they have received, and this choice is what is communicated to their neighbors. Here agents learn something about the belief of other agents, but may not know how certain the neighbors are that their choice is correct. In both cases, the information communicated (inferred value or decision) is not the same as the information that is available to the communicating agent and is less informative than what a global agent has access to.

The unifying theme of the work that follows is to identify structural obstacles to the performance of rational agents in a network. We will do so by comparing the performance of a global agent making use of all information in the network to that of agents making the best use of their private information and the social information communicated by their neighbors. How does the network topology (“who talks to whom”) affect the computations and performance of rational agents, and the group as a whole?

We will assume that each agent makes some or all of its private information public (social) by communicating it to its neighbors on the network. Interestingly, in some cases rational agents will discard their private information and only base decisions on social information. This can happen even when agents optimally use all available information, as we discuss below. As a result, the performance of agents in the network can, on average, be much worse than if agents work independently. Thus it is not always beneficial for a group of rational agents to exchange information: The information generated by a small subset of agents can dominate the decisions of the collective. In general the decisions based on a fraction of the information available to the collective will not be as good as that based on all the information.

In the remainder of this chapter, we give an overview of some of the main theoretical results and concepts in the literature. After going over the main assumptions of our models, we focus on sequential models of decision-making agents where actions (choices which communicate some of an agent's belief about correct decision to its neighbors) are done in a predefined order and next examine models where agents are continuously integrating information. We conclude this chapter by previewing the rest of the themes in this work.

1.1 Decision Making Agents

In the following, a set of agents is attempting to determine a true state, s , of the world in order to choose an action. Initially we only model the process of inferring

this hidden state of the world. In practice, this state could be some aspect of the environment that can be inferred from a stimulus (like which direction a sound is emanating from) or more simply the presence or absence of something of interest to the agent. Thus in some setups s can take a continuum of values, $s \in \mathbb{R}$, while in others s can be discrete. For instance, $s \in \{H^-, H^+\}$, when an agent is deciding between two hypotheses.

To determine what s is, agents will gather both private and social information, which we can generically denote by I . Private information could come from observations (measurements) of the state of the world. We will frequently assume that private information is independent between agents conditioned on the state of the world. Thus, any measurement errors are uncorrelated between the agents. This assumption of independence is unlikely to hold in the real world, but makes the models much more tractable. Even though we assume the private observations are independent, when an agent receives social information from multiple neighbors this information will be dependent. How agents deal with such network induced correlations is the topic of a major part of this work.

We consider Bayesian agents who compute the posterior of the distribution over the parameter, s , given the information I . The goal of these agents is to decide what state s best matches the given information. Such agents can use a maximum likelihood estimate based on this posterior, $\max_s p(s|I)$, to obtain this estimate. Agents will perform actions based on which state maximizes their posterior distribution, and, in some instances, agents will not perform an action until they have sufficient evidence; that is, until the posterior probability of some state

is sufficiently large.

1.2 Underlying Assumptions and Aumann's Theorem

The decision-making processes we investigate all rely on a number of underlying assumptions. Even though agents do not always have access to the private signals of others they typically know how the signals are distributed, conditioned on the possible true states of the world. That is, while agents are not aware of the private measurements of the other agents unless these are communicated, they do know the probability distribution of those measurements, given a private signal s . Thus agents know each others' measurement error distributions.

Agents also know who communicates with whom in the entire network. Thus agents have knowledge of the entire network topology. Furthermore, agents assume that other agents are rational, and are optimally incorporating all the private and social information that reaches them. Thus all agents are aware of the decision making procedure of all other agents, or how each agent translates private to shared information (and this procedure will be common to all agents on the network). Moreover, each agent knows that the other agents know this, and so on, *ad infinitum*.

This collection of assumptions, that agents know what other agents know about them, etc., is an example of what is referred to as **common knowledge** [2]. Even though it strictly only applies to finite partition information spaces, where

the true state, s , is only known to lie in one of finitely many disjoint sets comprising the total state space, the concept of common knowledge is useful for understanding how agents incorporate each other's decisions and for formally stating how the decision process works. The main result in [2] shows that when two agents start with the same prior, and the posteriors of two agents about the correct state s are common knowledge, those posteriors have to be equal. This condition is quite restrictive, but has been extended in [27] to show how agents who have the same priors can sequentially announce posteriors and eventually converge to a common posterior. This process of sharing information and eventually converging to a belief can be generalized using martingales [41, 22, 58].

1.3 Herding

Even networks of rational agents can be dominated by the social information of a small subgroup. Banerjee presented a simple model of this phenomenon which has since been termed **herding behavior** [5]. In Banerjee's model, a sequence of agents make and announce decisions sequentially about true value of some variable $s \in [0, 1]$, based on a private observation and the previously announced decisions. Here each agent's private observation will be a private signal that can be correct, misleading, or uninformative. In the absence of other decisions, agents will make a decision according to their private information. By design, no two agents will have the same misleading private signal. Thus when two agents announce the same decision, either they both received the correct signal, or the first

agent had a misleading signal and the second agent copied it because it received an uninformative signal.

As a result agents will throw away their private information when two previous agents agree. In the worst case scenario, the first agent to announce can get a misleading signal, and the second agent can get an uninformative signal and thus copy the decision of the first agent. Then every other agent will act optimally by discarding its private information and copying the first decision, causing the entire network to be wrong.

In this model agents throw away their private information in favor of incorrect social information with nonzero probability. Moreover, the decisions of a finite number of agents will almost always dominate the decisions of the entire network even when it is infinitely large. Thus, after finitely many agents communicate their decisions all agents will throw away their private signal and go with the rest of the group with probability equal to 1. This behavior, where all agents make optimal decisions based on the available information, but infinitely many agents discard private information for the social information communicated by a finite subset is known as herding. In [26], Banerjee's work was extended to more general networks, and the dependence of herding behavior on network topology was explained. In [50], herding is described more generally, along with the occurrence of different undesirable behaviors such as cases of social information becoming uninformative forcing agents to rely only on their private information.

We will discuss the issues of herding. However, the herding model in [5] differs from the one we consider because of the structure of the private information.

Herding would not occur in a network that has the structure of that in the Banerjee model, where each agent infers and communicates the value of a variable s based on an unbiased measurement with normally distributed error. The difference here is due to the fact that in the prior setup, the decisions are not sufficient statistics for an agents' belief: An agent's decision does not always fully describe what private information the agent had. The issue of herding is more relevant for Chapters 4-6 where there is a combination of social and private information, but we will see a related example in Chapter 3, where the private observation of a single agent has a disproportionate impact on the estimate of all agents in a network.

1.4 Probabilistic Network Herding

While the example of herding behavior described above has generated much interest, it is fairly specific. In particular, the original study did not consider how network structure can impact herding. In [26] the process is detailed for more general topologies, including ones changing stochastically. This can be thought of as allowing uncertainty about the communication structure: an agent receives information from neighbors but does not know with certainty where this social information originated. All agents are aware of the probabilistic structure of the network topology, each agent gets a private signal, and each agent communicates its decision according to the network structure.

This, and similar studies, examined the impact of network topology on **asymptotic learning**: Assume no signal (measurement) received by an agent is perfectly

informative. This means that there is no signal ξ and state s such that $P(s|\xi) = 1$. Then for a fixed finite network size, regardless of how information is exchanged between agents, there is always some chance that the agents who decide last in the process will make the incorrect decision, even if this is a global agent with knowledge of all private signals. However, if the network size is allowed to grow asymptotically, asymptotic learning means that agents later in the process converge to the correct decision with probability approaching 1 as the network size increases.

Hence asymptotic learning is related to herding. If all agents only rely on a small subset of agents to base their decisions, then their probability of making the correct decision is bounded below 1. The conditions for which asymptotic learning occurs are discussed in [26] and depend on whether private beliefs are bounded and whether the network contains “expanding structures.”

The private beliefs of agents are bounded when anytime an agent receives a signal ξ about the state s , then $P(\xi|s) \in [\epsilon, 1 - \epsilon]$ for some $0 < \epsilon < 1$. One way to think of expanding structures is to consider their complement: networks where all agents only observe a small pockets of agents. Asymptotic learning will typically occur with weak requirements on the structure when private beliefs are not bounded. When they are bounded, then the network has to have expanding structures, otherwise the beliefs of a small group of agents take over and prevent asymptotic learning with high probability.

1.5 Sequential Parameter Estimation in Feedforward Networks

In Chapter 2 we look at information sharing in networks with no cycles, i.e., directed networks where there are no paths from an agent back to itself, and all agents have a path to a common final agent. Our goal is to investigate how rational agents estimate the parameter s when the information they receive from other agents is redundant.

Early models of information sharing relied on computationally tractable interactions, such as the majority rule assumed in Condorcet's Jury Theorem [17], or local averaging assumed in the DeGroot model [19]. More recent models rely on the assumption of rational (Bayesian) agents who use private signals, measurements or observations of each other's actions to maximize utility. Such models of information sharing are often used in the economics literature, sometimes in combination with ideas from game theory. For instance, in a series of papers Mossel, Tamuz, and collaborators considered the propagation of information on an undirected network of rational agents, and showed that all agents on an irreducible graph integrate information optimally in a finite number of steps [40]. A similar setup was used by Acemoglu *et al.* to examine herd behavior in a network [1]. Mueller-Frank considered model social networks where private information of each agent is represented by a finite partition of the state space [44], and showed that in networks of non-Bayesian agents information is typically not aggregated optimally, but optimality is achieved in the presence of a single Bayesian

agent [45]. These, and related works [29], refer to such abstract models as “social networks”, and we follow this convention for simplicity. However, we note that this is at odds with the more traditional definition of this term [56].

Simplified models about how information is exchanged are also used in the political science literature to explain tendencies observed in social groups, and to fit to data. For example, Ortoleva and Snowberg used dependent Gaussian random variables to model the experimentally observed neglect of redundancies in information received by human observers [24]. They used this model to show how neglect of correlations can explain overconfidence in a sample of 3000 adults from the 2010 Cooperative Congressional Election Study (CCES) [46]. On the other hand, Levy and Razin show that similar correlation neglect can also lead to positive outcomes, as observers rely on actual information in forming opinions, rather than political orientation [36].

Such social network models of information propagation are generally either sequential or iterative. In sequential models, agents are ordered and act in turn based on a private signal and the observed action of their predecessors [5, 9]. In iterative models, agents make a single or a sequence of measurements, and iteratively exchange information with their neighbors [26, 40]. Sequential models have been used to illustrate information cascades [10], while iterative models have been used to illustrate agreement and learning [42].

In Chapter 2, we use a sequential model of information propagation and we identify conditions on the connection matrix (a subset of the adjacency matrix, which encodes the network topology) that hinder information propagation. and

relate this to graphical conditions on the network. We then identify which networks display the worst performance and, along with the simulations described in Chapter 3, we investigate performance as the size of the network grows.

1.6 Evidence Accumulation in Networks

In the last chapters, we review a model of evidence integration that has been commonly used to describe how a rational agent chooses between two options, and generalize it to describe a network of decision makers. Here, each agent is assumed to accumulate information through a sequence of measurements in order to choose one of two options when it has acquired a sufficient amount of evidence. These measurements are conditionally independent samples from one of two possible evidence distributions which the agent integrates to obtain the log likelihood ratio between the posterior probabilities of the two options. The agent then uses the classical Sequential Probability Ratio Test to make a decision ([20, 55]). This process is defined for measurements occurring in discrete time, but can be approximated by a continuous, drift-diffusion process [13]. The continuous approximation has been used to describe the activity of single neurons [28] and populations of neurons [37] that perform a similar computation.

In Chapter 4, we discuss this Drift Diffusion Model (DDM) of evidence accumulation in more detail. We then extend the DDM to the case of agents communicating their choices, but not individual observations, to their neighbors. We next describe the behavior of a pair of agents interacting unidirectionally, that is, when

1.6. EVIDENCE ACCUMULATION IN NETWORKS

only one agent in the pair is sharing its actions. In Chapter 5 we describe the interaction when the pair of agents are bidirectionally coupled so that both are sharing their actions. Then in Chapter 6, we consider larger networks, and investigate two important contrasting network structures: a fully connected clique and agents arranged in a line. Finally, we look at results of simulations and use them to provide several conjectures.

Chapter 2

Feedforward Networks

While there are billions of people on the planet, we exchange information with only a small fraction of them. How does information propagate through such social networks, shape our opinions, and influence our decisions? How do our interactions impact our choice of career or candidate in an election? More generally, how do we as agents in a network aggregate noisy signals to infer the state of the world?

These questions have a long history. The general problem is not easy to describe using a tractable mathematical model, as it is difficult to provide a reasonable probabilistic description of the state of the world. We also lack a full understanding of how perception [15, 6], and the information we exchange [3] shapes our decisions. Progress has therefore relied on tractable idealized models that mimic some of the main features of information exchange in social networks, some of which we described in Section 1.5.

In this chapter, we consider social networks in which information propagates directionally across layers of rational agents. We assume that agents want to estimate the true value of some parameter. Thus their utility function could be the negative square-difference seen in (1.0.1). To do so they use measurements (private information), as well as estimates communicated by their neighbors (social information). We assume that each agent makes a locally optimal estimate of the parameter based on this information, and communicates this estimate to agents downstream.

However, when agents receive information from a common source their estimates are correlated. We show that the resulting redundancy can lead to the loss of information about the parameter across layers of the network, even when all agents have full knowledge of the network's structure and make optimal use of all available information. A simple algebraic condition identifies networks in which information loss occurs, and we show that all such networks must contain a particular network motif.

2.1 Setup

Here we consider a sequential model in which information propagates directionally through layers of rational agents. The agents are part of a structured network, rather than a simple chain. As in the sequential model, we assume that information transfer is directional, and the recipient does not communicate information back to its source or sources. This assumption could describe the propagation of

2.1. SETUP

information via print or any other fixed medium.

We assume that at each step, a layer of agents receive information from those in a previous layer. This is different from previous sequential models where agents received information in turn from all their predecessors as in [5, 23, 57] and [7]. Importantly, the same information can reach an agent via multiple paths. Therefore, information received from agents in the previous layer can be redundant. Unlike in models of information neglect [46], we assume that agents take into account these redundancies in making decisions. We show that, depending on the network structure, even rational agents with full knowledge of the network structure cannot always resolve these redundancies. As a result, an estimate of the state of the world can degrade over layers. We also show that network architectures that lead to information loss can amplify an agent's bias in subsequent layers.

As an example, consider the network in Fig. 2.1.1(a). We assume that the first-layer agents make measurements x_1, x_2 , and x_3 of the state of the world, s , and that these measurements are unbiased, normally distributed, and have equal variance. This assumption means that minimum-variance unbiased estimators for these parameters are always linear combinations of individual measurements [32]. Each agent makes an estimate, $\hat{s}_1^{(1)}, \hat{s}_2^{(1)}$, and $\hat{s}_3^{(1)}$, of s . The superscript and subscript refer to the layer and agent number, respectively. An agent with global access to all first-layer estimates would be able to make the optimal (minimum-variance) estimate $\hat{s}_{\text{ideal}} = \frac{1}{3} \left(\hat{s}_1^{(1)} + \hat{s}_2^{(1)} + \hat{s}_3^{(1)} \right)$ of s .

All agents in the first layer then communicate their estimates to one or both of

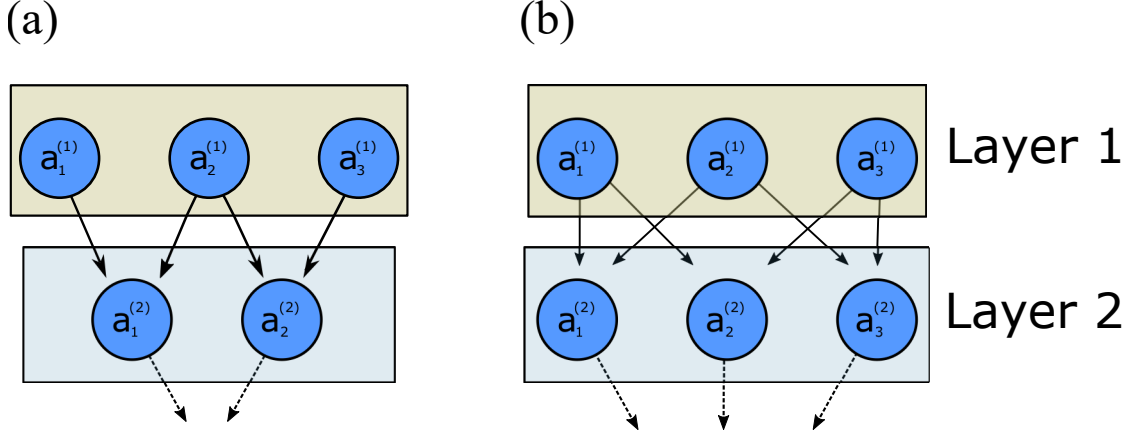


Fig. 2.1.1: Illustration of the general setup. Agents in the first layer (top layer in the figure) make measurements, x_1, x_2 , and x_3 , of a parameter s . In each layer agents make an estimate of this parameter, and communicate it to agents in the subsequent layer. Arrows indicate the direction in which information is propagated. We show that information about s degrades across layers in the network in panel (a), but not in the network in (b).

the second-layer agents. These in turn use the received information to make their own estimates, $\hat{s}_1^{(2)} = \frac{1}{2}(\hat{s}_1^{(1)} + \hat{s}_2^{(1)})$ and $\hat{s}_2^{(2)} = \frac{1}{2}(\hat{s}_2^{(1)} + \hat{s}_3^{(1)})$. An agent receiving the two estimates from the second layer then takes their linear combination to estimate s . However, in this network no linear combination of the locally optimal estimates, $\hat{s}_1^{(2)}$ and $\hat{s}_2^{(2)}$, equals the best estimate, $\hat{s}_{\text{ideal}}$, obtainable from all measurements in the first layer. Indeed,

$$\begin{aligned} \hat{s} &= \beta_1 \hat{s}_1^{(2)} + \beta_2 \hat{s}_2^{(2)} = \beta_1 \left(\hat{s}_1^{(1)} + \hat{s}_2^{(1)} \right) + \beta_2 \left(\hat{s}_2^{(1)} + \hat{s}_3^{(1)} \right) \\ &\neq \hat{s}_{\text{ideal}} = \frac{1}{3} \left(\hat{s}_1^{(1)} + \hat{s}_2^{(1)} + \hat{s}_3^{(1)} \right), \end{aligned}$$

with the inequality holding for any choice of β_1, β_2 . Moreover, assume the estimates of first-layer agents are biased, and $\hat{s}_i^{(1)} = x_i + b_i$. If the other agents are unaware of this bias, then, as we will show, the final estimate is $\hat{s} = \left(\frac{1}{4}, \frac{1}{2}, \frac{1}{4} \right) \cdot (\hat{s}_1^{(1)} + b_1, \hat{s}_2^{(1)} + b_2, \hat{s}_3^{(1)} + b_3) = \left(\frac{1}{4}, \frac{1}{2}, \frac{1}{4} \right) \cdot \hat{s}^{(1)} + \left(\frac{1}{4}, \frac{1}{2}, \frac{1}{4} \right) \cdot (b_1, b_2, b_3)$. Thus the bias

of the second agent in the first layer, $a_2^{(1)}$, has disproportionate weight in the final estimate.

In this example the information about the state of the world (parameter), s , available from second-layer agents is less than that available from first-layer agents. In the preceding example the measurement x_2 is used by both agents in the second layer. The estimates of the two second-layer agents are therefore correlated, and the final agent cannot disentangle them to recover the ideal estimate. We will show that the type of subgraph shown in Fig. 2.1.1(a), which we call a *W-motif*, provides the main obstruction to obtaining the best estimate in subsequent layers.

2.2 The Model

We consider feedforward networks having n layers and identify each node of a network with an agent. The structure of the network is thus given by a directed graph with agents occupying the vertices. Agents in each layer only communicate with those in the next layer. For convenience, we will assume that layer n consists of a single agent that receives information from all agents in layer $n - 1$. This final agent in the last layer therefore makes the best estimate based on all the estimates in the next-to-last layer. We will use this last agent's estimate to quantify information loss in the network. Two example networks are given in Fig. 2.1.1, with the single agent in the final, third layer not shown.

We assume that all agents are Bayesian, and know the structure of the network.

2.2. THE MODEL

Every agent estimates an unknown parameter, $s \in \mathbb{R}$, but only the agents in the first layer make a measurement of this parameter. Each agent makes the best possible estimate given the information it receives and communicates this estimate to a subset of agents in the next layer. We also assume that measurements, x_i , made by agents in the first layer are independent and normally distributed with mean s , and variance σ_i^2 , that is $x_i \sim \mathcal{N}(s, \sigma_i^2)$. Furthermore, every agent in the network knows the variance of each measurement in the first layer, σ_i^2 . Also, for simplicity, we will assume that all agents share an improper, flat prior over s . This assumption does not affect the main results.

An agent with access to all of the measurements, $\{x_i\}_i$, has access to all the information available about s in the network. This agent can make an **ideal** estimate, $\hat{s}_{\text{ideal}} = \operatorname{argmax}_s p(s|x_1, \dots, x_n)$. We assume that the actual agents in the network are making locally optimal, maximum-likelihood estimates of s , and ask when the estimate of the final agent equals the ideal estimate, $\hat{s}_{\text{ideal}}$.

Individual Estimate Calculations Each agent in the first layer only has access to its own measurement, and makes an estimate equal to this measurement. We therefore write $\hat{s}_i^{(1)} = x_i$. We denote the j^{th} agent in layer k by $a_j^{(k)}$. Each of these agents makes an estimate, $\hat{s}_j^{(k)}$ of s , using the estimates communicated by its neighbors in the previous layer. Under our assumptions, the posterior computed by any agent is normal and the vector of estimates in a layer follows a multivariate Gaussian distribution. As agents in the second layer and beyond can share upstream neighbors, the covariance between their estimates is typically nonzero.

2.2. THE MODEL

We show that under the assumption that the variance of the initial measurements and the structure of the network are known to all agents, each agent knows the full joint posterior distribution over s for all agents it receives information from.

Weight Matrices We define the connectivity matrix $C^{(k)}$ for $1 \leq k \leq n - 1$ as,

$$C_{ij}^{(k)} = \begin{cases} 1, & \text{if } a_j^{(k)} \text{ communicates with } a_i^{(k+1)} \\ 0, & \text{otherwise.} \end{cases} \quad (2.2.1)$$

An agent receives a subset of estimates from the previous layer determined by this connectivity matrix. The agent then uses this information to make its own, maximum-likelihood estimate of s . By our assumptions, this estimate will be a linear combination of the communicated estimates [32]. Denoting by $\hat{\mathbf{s}}^{(k)}$ the vector of estimates in the k^{th} layer, we can therefore write $\hat{\mathbf{s}}_i^{(k+1)} = \mathbf{w}_i^{(k+1)} \cdot \hat{\mathbf{s}}^{(k)}$, and

$$\hat{\mathbf{s}}^{(k+1)} = W^{(k+1)} \hat{\mathbf{s}}^{(k)}.$$

Here $W^{(k+1)}$ is a matrix of weights applied to the estimates in the k^{th} layer.

Weighting by Precision We can write $\hat{\mathbf{s}}^{(1)} = W^{(1)} \mathbf{x}$ where $W^{(1)}$ is the identity matrix and $\mathbf{x}$ is the vector of measurements made in the first layer. We assume that all measurements have finite, nonzero variance. Using standard estimation theory results [32], we can compute the optimal estimates for agents in the second layer. Defining $w_i := \frac{1}{\sigma_i^2}$, we can calculate $W^{(2)}$ entrywise: $w_{ij}^{(2)}$ is 0 if agent $a_i^{(2)}$ does *not* communicate with $a_j^{(1)}$. Otherwise $w_{ij}^{(2)} = \frac{w_j^{(1)}}{\sum_{k \rightarrow i} w_k^{(1)}}$, where the sum is

2.2. THE MODEL

taken over all agents in the first layer that communicate with agent $a_i^{(2)}$. Therefore,

$$\hat{\mathbf{s}}^{(2)} = W^{(2)} \hat{\mathbf{s}}^{(1)} = W^{(2)} W^{(1)} \mathbf{x}. \quad (2.2.2)$$

Covariance Matrices The estimates in the second layer and beyond can be correlated. Let L_k be the number of agents in the k^{th} layer and for $2 \leq k \leq n-1$ define $\Omega^{(k)} = (\xi_{ij}^{(k)})$ as the $L_k \times L_k$ covariance matrix of estimates in the k^{th} layer,

$$\xi_{ij}^{(k)} = \text{Cov}(\hat{s}_i^{(k)}, \hat{s}_j^{(k)}).$$

When all of the weights are known, we have

$$\hat{\mathbf{s}}^{(k)} = W^{(k)} \hat{\mathbf{s}}^{(k-1)} = W^{(k)} W^{(k-1)} \hat{\mathbf{s}}^{(k-2)} = \dots = \left(\prod_{l=0}^{k-2} W^{(k-l)} \right) \hat{\mathbf{s}}^{(1)}. \quad (2.2.3)$$

The i^{th} row of $\left(\prod_{l=0}^{k-2} W^{(k-l)} \right)$ is the vector of weights that the agent $a_i^{(k)}$ applies to the first-layer estimates, since its entries are the coefficients in $s_i^{(k)}$.

The complete covariance matrix, $\Omega^{(k)}$, can therefore be written as

$$\begin{aligned} \Omega^{(k)} &= \text{Cov}(\hat{\mathbf{s}}^{(k)}) = \text{Cov}(W^{(k)} \hat{\mathbf{s}}^{(k-1)}) = W^{(k)} \text{Cov}(\hat{\mathbf{s}}^{(k-1)}) \left(W^{(k)} \right)^T \\ &= \left(\prod_{l=0}^{k-2} W^{(k-l)} \right) \text{Cov}(\hat{\mathbf{s}}^{(1)}) \left(\prod_{l=0}^{k-2} W^{(k-l)} \right)^T \\ &= \left(\prod_{l=0}^{k-2} W^{(k-l)} \right) \text{Diag} \left(\frac{1}{w_1}, \dots, \frac{1}{w_{L_1}} \right) \left(\prod_{l=0}^{k-2} W^{(k-l)} \right)^T. \end{aligned} \quad (2.2.4)$$

Now the i^{th} agent in layer $k \geq 3$, $a_i^{(k)}$, can use $\Omega^{(k-1)}$ to calculate $\mathbf{w}_i^{(k)}$. If the agent is not connected to all agents in the $(k-1)^{\text{th}}$ layer, it uses the submatrix of $\Omega^{(k-1)}$ with rows and columns corresponding to the agents in the previous

2.2. THE MODEL

layer that communicate their estimates to it. We denote this submatrix $R_i^{(k-1)}$. As in [43], we assume that we remove edges from the graph so that all submatrices $R_i^{(k-1)}$ are invertible, but all estimates are the same as in the original network.

An agent thus receives estimates that follow a multivariate normal distribution, $\mathcal{N}(\hat{\mathbf{s}}_{j \rightarrow i}^{(k-1)}, R_i^{(k-1)})$, see [32]. The weights assigned by agent $a_i^{(k)}$ to the estimates of agents in the previous layer are therefore (see also [43]),

$$\tilde{\mathbf{w}}_i^{(k)} = \frac{\mathbf{1}^T (R_i^{(k-1)})^{-1}}{\mathbf{1}^T (R_i^{(k-1)})^{-1} \mathbf{1}}. \quad (2.2.5)$$

We define $\mathbf{w}_i^{(k)}$ by using the corresponding entries from $\tilde{\mathbf{w}}_i^{(k)}$ and setting the remainder to zero. In the following we describe the maximum-likelihood estimate that can be made from all the estimates in a layer. For simplicity, we denote this final estimate by $\hat{s}$. The following results are standard [32].

Proposition 1. *The posterior distribution over s of the final agent is normal with*

$$\hat{s} = \frac{\mathbf{1}^T (\Omega^{(n-1)})^{-1}}{\mathbf{1}^T (\Omega^{(n-1)})^{-1} \mathbf{1}} \hat{\mathbf{s}}^{(n-1)} \quad \text{and} \quad \text{Var} [\hat{s}] = \frac{1}{\mathbf{1}^T (\Omega^{(n-1)})^{-1} \mathbf{1}} \quad (2.2.6)$$

where $\Omega^{(n-1)}$ is defined by Eq. (2.2.4) and Eq. (2.2.5). Here $\hat{s}$ is the maximum-likelihood, as well as minimum-variance, unbiased estimate of s .

It follows from Eq. (2.2.3) that the estimate of any agent in the network is a convex linear combination of the estimates in the first layer.

Examples Returning to the example in Fig. 2.1.1(a) we have

$$C^{(1)} = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix}, W^{(2)} = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} & 0 \\ 0 & \frac{1}{2} & \frac{1}{2} \end{pmatrix}, \Omega^{(2)} = \begin{pmatrix} \frac{1}{2} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{2} \end{pmatrix}, (\Omega^{(2)})^{-1} = \frac{16}{3} \begin{pmatrix} \frac{1}{2} & -\frac{1}{4} \\ -\frac{1}{4} & \frac{1}{2} \end{pmatrix}$$

The final agent applies the weights $W^{(3)} = \begin{pmatrix} \frac{1}{2}, \frac{1}{2} \end{pmatrix}$ to the estimates from the second layer. We thus have the final estimate $\hat{s} = \begin{pmatrix} \frac{1}{4}, \frac{1}{2}, \frac{1}{4} \end{pmatrix} \cdot \hat{\mathbf{s}}^{(1)}$ with $\text{Var} [\hat{s}] = \frac{3}{8}$. The variance of the ideal estimate is $\frac{1}{3}$.

On the other hand, the final agent in the example in Fig. 2.1.1(b) makes an ideal estimate: Here $W^{(2)} = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} & 0 \\ \frac{1}{2} & 0 & \frac{1}{2} \\ 0 & \frac{1}{2} & \frac{1}{2} \end{pmatrix}$, $\Omega^{(2)} = \begin{pmatrix} \frac{1}{2} & \frac{1}{4} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{2} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{4} & \frac{1}{2} \end{pmatrix}$, and after inverting $\Omega^{(2)}$ we see that applying a weight of $\frac{1}{3}$ to every agent in the second layer gives the ideal estimate, $\hat{s} = \begin{pmatrix} \frac{1}{3}, \frac{1}{3}, \frac{1}{3} \end{pmatrix} \cdot \hat{\mathbf{s}}^{(1)}$.

Remark. If the agents have a proper normal prior with mean χ and variance σ_p^2 , then agents in the first layer make the estimate,

$$\hat{s}_i^{(1)} = \frac{\sigma_i^{-2}}{\sigma_i^{-2} + \sigma_p^{-2}} x_i + \frac{\sigma_p^{-2}}{\sigma_i^{-2} + \sigma_p^{-2}} \chi,$$

with a similar form in the following layers. This does not change the subsequent results as long as all agents have the same prior. Also, if each agent in the network makes a measurement, then it makes an estimate based on both this private information and the social information communicated by its upstream neighbors. However, the general ideas that follow remain unchanged.

2.3 Results

We ask what graphical conditions need to be satisfied so that the agent in the final layer makes an ideal estimate. That is, when does knowing all estimates of the agents in the $(n - 1)^{\text{st}}$ layer give an estimate that is as good as possible given the measurements of all first-layer agents. We refer to a network in which the final estimate is ideal as an **ideal** network.

Proposition 2. *A network with n layers and $\sigma_i^2 \neq 0$ for $i = 1, \dots, L_1$, is ideal if and only if the vector of inverse variances, $(w_1, \dots, w_{L_1})$, is in the row space of the weight matrix product $(\prod_{l=0}^{n-3} W^{(n-1-l)})$.*

Proof. In this setting the ideal estimate is

$$\hat{s}_{\text{ideal}} = \frac{1}{\sum_i w_i} \sum_{i=1}^{L_1} w_i \hat{s}_i^{(1)}. \quad (2.3.1)$$

The network is ideal if and only if there are coefficients $\beta_j \in \mathbb{R}$ such that

$$\hat{s}_{\text{ideal}} = \sum_{j=1}^{L_{n-1}} \beta_j \hat{s}_j^{(n-1)}.$$

Matching coefficients with Eq. (2.3.1), we need

$$\frac{1}{\sum_j w_j} \sum_{i=1}^{L_1} w_i \hat{s}_i^{(1)} = (\beta_1, \dots, \beta_{L_{n-1}}) \cdot \hat{\mathbf{s}}^{(n-1)},$$

or equivalently,

$$\begin{aligned} \frac{1}{\sum_j w_j} (w_1, \dots, w_{L_1}) \cdot \hat{\mathbf{s}}^{(1)} &= (\beta_1, \dots, \beta_{L_{n-1}}) \cdot W^{(n-1)} \hat{\mathbf{s}}^{(n-2)} \\ &= (\beta_1, \dots, \beta_{L_{n-1}}) \cdot \left(\prod_{l=0}^{n-3} W^{(n-1-l)} \right) \hat{\mathbf{s}}^{(1)}. \end{aligned}$$

2.3. RESULTS

Equality holds exactly when $(w_1, \dots, w_{L_1})$ is in the row space of $\left(\prod_{l=0}^{n-3} W^{(n-1-l)}\right)$.

□

In particular, a three-layer network with $\sigma_i^2 = \sigma$ for all $i \in \{1, \dots, L_1\}$ is ideal if and only if the vector $\vec{1} = (1, 1, \dots, 1)$ is in the row space of the connectivity matrix $C^{(1)}$ defined by Eq. (2.2.1). We will use and extend this observation below.

2.3.1 Graphical Conditions for Ideal Networks

We say that a network contains a **W-motif** if two agents downstream receive common input from a first-layer agent, as well as private input from two distinct first-layer agents. Examples are shown in Fig. 2.1.1(a) and Fig. 2.3.1. A rigorous definition follows.

We will show that *all networks that are not ideal contain a W-motif*. However, the converse is not true: The network in Fig. 2.1.1(b) contains many W-motifs, but is ideal. Therefore ideal networks can contain a W-motif, as the redundancy introduced by a W-motif can sometimes be resolved. Hence, additional graphical conditions determine if the network is ideal.

As shown in Fig. 2.3.1, in a W-motif there is a directed path from a single agent in the first layer to two agents in the third layer. There are also paths from distinct first-layer agents to the two third-layer agents. This general structure is captured by the following definitions.

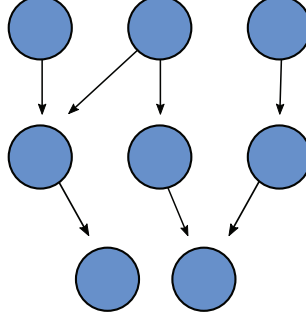


Fig. 2.3.1: A W-motif spanning three layers.

Definition 1. The path matrix P^{kl} , $l < k$, from layer l to layer k is defined by,

$$P_{ij}^{kl} = \begin{cases} 1, & \text{if there is a directed path from agent } a_j^{(l)} \text{ to agent } a_i^{(k)} \\ 0, & \text{otherwise.} \end{cases}$$

Definition 2. A network contains a W-motif if a path matrix from the first layer, P^{k1} , has a 2×3 submatrix equal to $\begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix}$ (modulo column permutation). Graphically, two agents in layer k are connected to one common, and two distinct agents in layer 1.

Theorem 1. A non-ideal network in which every agent communicates its estimate to the subsequent layer must contain a W-motif. Equivalently, if there are no W-motifs, then the network is ideal.

Proof of Theorem 1 We will spend the remainder of this subsection proving the theorem. Intuitively, any agent receives estimates that are a linear combination of first-layer measurements. If there are no W-motifs, any two estimates are either

2.3. RESULTS

obtained from disjoint sets of measurements, or the measurements in the estimate of one agent contain the measurements in the estimate of another. When measurements are disjoint, there are no correlations between the estimates and thus no degradation of information. When one set of measurements contains the other, then the estimates in the subset are redundant and can be discarded. Therefore, this redundant information does not cause a degradation of the final estimate.

We start with the simpler case of a W-motif between the first two layers and then extend it to the general case. We begin with definitions that will be used in the proof.

Let g be the **input-map** which maps an agent to the subset of agents in the first layer that it receives information from (through some path). That is, $g(a_i^{(j)})$ is the set of agents in the first layer that provide input to $a_i^{(j)}$. It is intuitive – and we show it formally in Lemma 1 – that a network contains a W-motif if each of the inputs to two agents, A and B are not contained in the other, and their intersection is not empty. That is, $g(A) \not\subseteq g(B)$ and $g(B) \not\subseteq g(A)$, but $g(A) \cap g(B) \neq \emptyset$. If these conditions are met, we also say that the inputs of A and B have a **nontrivial intersection**. If $g(A) \subseteq g(B)$, we say that the input of B **overlaps** the input of A : every agent which contributes to the estimate of A also contributes to the estimate of B .

Similarly, we let f be the **output-map** which maps an agent, $a_i^{(j)}$, to the set of all agents in the next, $j + 1^{\text{st}}$, layer that receive input from $a_i^{(j)}$. We first prove a few lemmas essential to the proof of Theorem 1.

2.3. RESULTS

Lemma 1. *Assume a network does not contain a W-motif and there are two agents, $a_{i_1}^{(k)}$ and $a_{i_2}^{(k)}$, with $g(a_{i_1}^{(k)}) \cap g(a_{i_2}^{(k)})$ nonempty. Then $g(a_{i_1}^{(k)})$ overlaps or is overlapped by $g(a_{i_2}^{(k)})$.*

Proof. We prove the claim by contradiction. If one input does not overlap the other, then there are two distinct first-layer agents $a_{n_1}^{(1)}$ and $a_{n_2}^{(1)}$ such that $a_{n_1}^{(1)} \in g(a_{i_1}^{(k)}) \setminus g(a_{i_2}^{(k)})$ and $a_{n_2}^{(1)} \in g(a_{i_2}^{(k)}) \setminus g(a_{i_1}^{(k)})$. This means $P_{i_1 n_1}^{k1} = P_{i_2 n_2}^{k1} = 1$ and $P_{i_1 n_2}^{k1} = P_{i_2 n_1}^{k1} = 0$. Since the inputs of the agents have nonempty intersection, we also have $P_{i_1 m}^{k1} = P_{i_2 m}^{k1} = 1$ for some m . Thus there is a 2×3 submatrix of P^{k1} which, up to rearrangement of the columns, is equal to $\begin{pmatrix} 1 & 1 & 0 \\ 1 & 0 & 1 \end{pmatrix}$ and the network contains a W-motif, contrary to assumption. $\square$

Every agent's estimate is a convex linear combination of estimates in the first layer, given by Eq.(2.2.3). We will use the corresponding weight vectors in the following proofs. We show that in networks without W-motifs, agents will only be receiving collections of estimates with weight vectors which pairwise either have disjoint support (nonzero indices) or the support is contained in the support of the other agent. Thus, with no W-motifs, no two agents have inputs with nontrivial intersection. The next two lemmas will allow us to easily calculate the estimates of such agents.

Lemma 2. *Let r, s, t be positive integers, $w_i = \sigma_i^{-2}$, and consider three weight vectors*

2.3. RESULTS

applied by three agents in layer k , $a_1^{(k)}$, $a_2^{(k)}$, and $a_3^{(k)}$, to the estimates of the first layer:

$$\begin{aligned} v_1 &= \left(\frac{w_1}{\sum_{i=1}^r w_i}, \dots, \frac{w_r}{\sum_{i=1}^r w_i}, 0, \dots, 0 \right) \\ v_2 &= \left(\frac{w_1}{\sum_{i=1}^{r+s} w_i}, \dots, \frac{w_{r+s}}{\sum_{i=1}^{r+s} w_i}, 0, \dots, 0 \right) \\ v_3 &= \left(0, \dots, 0, \frac{w_{r+s+1}}{\sum_{i=r+s+1}^{r+s+t} w_i}, \dots, \frac{w_{r+s+t}}{\sum_{i=r+s+1}^{r+s+t} w_i}, 0, \dots, 0 \right). \end{aligned}$$

An agent $a_i^{(k+1)}$ in $f(a_1^{(k)}) \cap f(a_2^{(k)})$, but not in $f(a_3^{(k)})$, will use weight vector v_2 . An agent $a_i^{(k+1)}$ in $f(a_2^{(k)}) \cap f(a_3^{(k)})$, but not $f(a_1^{(k)})$, will use weight vector

$$v_4 = \left(\frac{w_1}{\sum_{i=1}^{r+s+t} w_i}, \dots, \frac{w_{r+s+t}}{\sum_{i=1}^{r+s+t} w_i}, 0, \dots, 0 \right).$$

Proof. First, consider an agent receiving the first two estimates with weights v_1 and v_2 . Suppose that a fictitious agent receives a collection of estimates with weight vectors $\{z_1, \dots, z_{r+s}\}$, where $z_i = (0, \dots, 0, 1, 0, \dots, 0)$, i.e., each estimate equals the measurement of agent $a_i^{(1)}$. This fictitious agent can obtain any linear combination of the first $r+s$ measurements. The linear combination with lowest variance has weights given by v_2 . Therefore, an agent receiving measurements corresponding to the weight vectors v_1 and v_2 cannot do better than the estimate of agent $a_2^{(k)}$ with weights given by v_2 .

A similar argument works when estimates are received from agents $a_2^{(k)}$ and $a_3^{(k)}$. Since these two agents make locally optimal estimates based on non-overlapping sets of measurements in the first layer, the best estimate is obtained by combining the two sets of measurements. This is precisely the estimate corresponding to the weights given by vector v_4 . $\square$

2.3. RESULTS

Lemma 3. *Suppose an agent, $a_i^{(k)}$, receives a collection of estimates such that for any pair, there is a relabeling of agents in the first layer that makes the pair look like v_1 and v_2 or like v_2 and v_3 in Lemma 2. Then, up to some relabeling of the agents in the first layer, that agent will make an estimate with corresponding weight vector*

$$v = \left(\frac{w_1}{\sum_{i=1}^r w_i}, \dots, \frac{w_r}{\sum_{i=1}^r w_i}, 0, \dots, 0 \right).$$

Proof. Let the vectors z_i be defined as in the proof of Lemma 2. Relabel the first-layer agents so that only the first r entries of the weight vector applied by agent $a_i^{(k)}$ are non-zero. Then a fictitious agent receiving estimates with weight vectors $z_i, 1 \leq i \leq r$ can construct any estimate that agent $a_i^{(k)}$ can obtain. The optimal estimate of this fictitious agent has weight vector v . Hence if some linear combination of the weight vectors of estimates communicated to agent $a_i^{(k)}$ equals v , this linear combination defines the best estimate.

Then for each $j = 1, \dots, r$, we can find a weight vector, v_j , which is nonzero in the j^{th} entry with support that contains the support of every other weight vector which is nonzero in the j^{th} entry. Such a vector exists by the assumption that any two vectors have disjoint support or the support of one contains the other. Therefore, we can find the weight vector with maximal support for each entry. If we take the distinct elements of $\{v_j : 1 \leq j \leq r\}$, then these maximal weight vectors will have disjoint support that partitions the first r indices. Therefore,

$$v = \frac{1}{\sum_{i=1}^r w_i} \sum_{v_j \text{ distinct}} \left(\sum_{i=1, v_j^i \text{ nonzero}}^r w_i \right) v_j,$$

2.3. RESULTS

which shows the lemma. $\square$

We now state and prove the three-layer case of Theorem 1 and then use it to finish the proof of Theorem 1.

Proposition 3. *If a three-layer network is not ideal and every first-layer agent communicates with at least one second-layer agent, then the network must contain a W-motif.*

Proof. Assume the network does not contain a W-motif. Given a first-layer agent $a_i^{(1)}$, Lemma 1 says that for any two agents in $f(a_i^{(1)})$, one agent's input must overlap the other. Two second-layer agents thus receive estimates with input sets where one overlaps the other, or the sets do not intersect. Thus the set of weight vectors in the second layer satisfies the assumptions of Lemma 3. As all agents from the first layer communicate with the final agent, the network is ideal. $\square$

To obtain the proof of Theorem 1, we use induction with Proposition 3 as a base case.

Proof of Theorem 1. Assume the network has n layers, there are no W-motifs, and every agent (except those in the first layer) receives input from at least one other agent. Lemma 1 implies that in the second layer each pair of agents has either disjoint input or one overlaps the other. Thus in the third layer, by relabeling the agents, each agent makes an estimate with weight vector of the form:

$$\frac{1}{\sum_{i=1}^r w_i} (w_1, \dots, w_r, 0, \dots, 0).$$

Now assume that any estimate in layer k can be put in this form by relabeling

2.4. SUFFICIENT CONDITIONS FOR IDEAL THREE-LAYER NETWORKS

the agents. Since there are no W -motifs, Lemma 1 implies that set of measurements used by agents $a_{i_1}^{(k)}$ and $a_{i_2}^{(k)}$ is disjoint or overlapping. This again allows us to apply Lemma 3 and any agent in layer $k + 1$ makes an estimate whose weight vector again has the form $\frac{1}{\sum_{i=1}^r w_i}(w_1, \dots, w_r, 0, \dots, 0)$. Applying the same argument to the final agent, where every entry will be nonzero in some penultimate-layer agent's weight vector, we have that the network is ideal.

□

2.4 Sufficient Conditions for Ideal Three-Layer Networks

We next consider only three-layer networks. This allows us to give a graphical interpretation of the algebraic condition describing ideal networks in Proposition 2. To do so, we will use the following corollary of the proposition.

Corollary 1. *Let $C^{(1)}$ be defined as in Eq. (2.2.1). Then a three-layer network is ideal if and only if the vector $m\vec{1}$ is in the row space of $C^{(1)}$ over $\mathbb{Z}$ for some nonzero $m \in \mathbb{N}$.*

Proof. We will show that a three-layer network is ideal if and only if $m\vec{1}$ is in the row space of $C^{(1)}$ over $\mathbb{Z}$ for some $m \in \mathbb{N}$. We do this by first showing that the network is ideal if and only if $\vec{1}$ is in the row space of $C^{(1)}$ over $\mathbb{R}$, and then we show that this is equivalent to $m\vec{1}$ being in the row space of $C^{(1)}$ over $\mathbb{Z}$.

By Proposition 2, a three-layer network is ideal if and only if $(w_1, \dots, w_{L_1})$ is

2.4. SUFFICIENT CONDITIONS FOR IDEAL THREE-LAYER NETWORKS

in the row space of $W^{(2)}$. We claim that this is equivalent to $\vec{1}$ being in the row space of $C^{(1)}$: Multiplying each row of $W^{(2)}$ by the common denominator of the nonzero entries gives

$$\mathcal{R}(W^{(2)}) = \mathcal{R}(C^{(1)}\text{Diag}(w_1, \dots, w_{L_1})),$$

where $\mathcal{R}$ denotes the row space. By definition, $\vec{1}$ is a linear combination of the rows of $C^{(1)}$ if and only if

$$1 = \sum_i \beta_i C_{ij}^{(1)}, \quad \forall j.$$

This holds if and only if

$$w_j = \sum_i \beta_i w_j C_{ij}^{(1)}, \quad \forall j.$$

The last equality is equivalent to

$$(w_1, \dots, w_{L_1}) = \sum_i \beta_i (C^{(1)}\text{Diag}(w_1, \dots, w_{L_1}))_i,$$

which means $(w_1, \dots, w_{L_1})$ is in the row space of $W^{(2)}$. Hence, for three-layer networks, the network is ideal if and only if the vector $\vec{1}$ is in the row space of $C^{(1)}$ over $\mathbb{R}$.

Thus it remains to show that $\vec{1} \in \mathcal{R}(C^{(1)})$ over $\mathbb{R}$ is equivalent to $\vec{1} \in \mathcal{R}(C^{(1)})$ over $\mathbb{Z}$. If $m\vec{1} \in \mathcal{R}(C^{(1)})$ over $\mathbb{Z}$, then it is a linear combination of the rows of $C^{(1)}$ with integer coefficients. Multiplying the coefficients of this linear combination by $\frac{1}{m}$ shows that $\vec{1}$ is in the row space of $C^{(1)}$ and hence the network is ideal.

If $\vec{1}$ is in the row space of $C^{(1)}$ over $\mathbb{R}$, then by closure of $\mathbb{Q}^n$ this means there is some linear combination of the rows of $C^{(1)}$ over $\mathbb{Q}$ which is equal to $\vec{1}$:

$$\sum_{i=1}^{L_2} \alpha_i C_i^{(1)} = \vec{1}, \quad \alpha_i \in \mathbb{Q}.$$

2.4. SUFFICIENT CONDITIONS FOR IDEAL THREE-LAYER NETWORKS

Multiplying both sides by the absolute value of the product of the denominators of the nonzero α_i shows that

$$\sum_{i=1}^{L_2} \beta_i C_i^{(1)} = m\vec{1}, \quad \beta_i \in \mathbb{Z}$$

for some $m \in \mathbb{N}$ and thus $m\vec{1}$ is in the row space of $C^{(1)}$ over $\mathbb{Z}$. $\square$

Note that the corollary is not restricted to the case where first-layer agents have equal variance measurements; whether the network is ideal or not depends entirely on the connection matrix $C^{(1)}$. The i^{th} row of the matrix $C^{(1)}$ corresponds to the inputs of agent $a_i^{(2)}$, and the sum of the j^{th} column is the out-degree of agent $a_j^{(1)}$. Therefore, Corollary 1 is equivalent to the following: If each second-layer agent applies equal integer weights to all of its received estimates, then a three-layer network is ideal if and only if, for some choice of weights, the weighted out-degrees of all agents in the first layer are equal. Hence, we have the following special case:

Corollary 2. *A three-layer network is ideal if all first-layer agents have equal out-degree in each connected component of the network restricted to the first two layers.*

In the connected network in Fig. 2.1.1(a), the second agent in the first layer has greater out-degree than the others, while the agents in the first layer of the connected network in Fig. 2.1.1(b) have equal out-degree.

Some row reduction operations can be interpreted graphically. Let g be the **input-map** which maps an agent, $a_i^{(2)}$, to the subset of agents in the first layer that it receives estimates from. Formally, let $\mathcal{P}(A)$ denote the power set of a set A ,

2.4. SUFFICIENT CONDITIONS FOR IDEAL THREE-LAYER NETWORKS

then $g: \{a_1^{(2)}, \dots, a_{L_2}^{(2)}\} \rightarrow \mathcal{P}\{a_1^{(1)}, \dots, a_{L_1}^{(1)}\}$ is defined by $a_j^{(1)} \in g(a_i^{(2)})$ if agent $a_j^{(1)}$ communicates with agent $a_i^{(2)}$, i.e. if $C_{ij}^{(1)} = 1$.

If $g(a_i^{(2)}) \subseteq g(a_j^{(2)})$ for some $i \neq j$, then some of the information received by $a_j^{(2)}$ is redundant, as it is already contained in the estimate of agent $a_i^{(2)}$. We can then reduce the network by eliminating the directed edges from $g(a_i^{(2)})$ to $a_j^{(2)}$, so that in the reduced network $g(a_i^{(2)}) \cap g(a_j^{(2)}) = \emptyset$. This reduction is equivalent to subtracting row i from row j of $C^{(1)}$ resulting in a connection matrix with the same row space. By Proposition 2, the reduced network is ideal if and only if the original network is ideal. This motivates the following definition.

Definition 3. A three-layer network is said to be **reduced** if $g(a_i^{(2)})$ is not a subset of $g(a_j^{(2)})$ for all $1 \leq i \neq j \leq L_2$.

Reducing a network eliminates edges, and results in a simpler network structure. In a three-layer network, this will not affect the final estimate: Since reduction leaves the row space of $C^{(1)}$ unchanged, the final estimate in the reduced and unreduced network is the result of applying the same weights to the first-layer estimates. This reduction procedure often simplifies identification of ideal networks to a counting of out-degrees (see Corollary 2).

Example In Fig. 2.4.1, we illustrate a two-step reduction of a network. In both steps, an agent (in yellow) has an input set which is overlapped by the input sets of some other agents (bolded). We use this to cancel the common inputs to the bolded agents and simplify the network. In the first step, note that the yellow

2.4. SUFFICIENT CONDITIONS FOR IDEAL THREE-LAYER NETWORKS

agent receives input (in red) from a single first-layer agent. We use this to remove all of the other connections (in green) emanating from this first-layer agent. In the second step, we again see that the yellow agent receives input (red) that is overlapped by input to the agent next to it. We can thus remove the redundant inputs (in green) to the bolded agent. The reduced network has 5 connected components all containing vertices with equal out-degree. Hence, this network is ideal by Corollary 2.

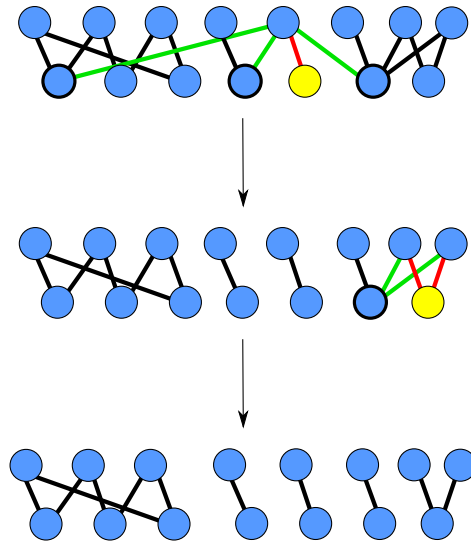


Fig. 2.4.1: Example of a two step network reduction. It is difficult to tell whether the network on top is ideal. However, after two steps of reduction, all first-layer agents in each of the five connected components have equal out-degree. The network is therefore ideal.

2.5 Conclusion

We examined how information about the world propagates through layers of rational agents. We assumed that at each step, a group of agents makes an inference about the state of the world from information provided by their predecessors. The setup is related, but different from information cascades where a chain of rational agents make decisions in turn [5, 23, 57, 7], or recurrent networks where agents exchange information iteratively [40]. The assumption that the observed variables in our analysis follow a Gaussian distribution simplified the analysis considerably. However, we believe that the main results hold under more general assumptions. Our preliminary work shows that when agents in the first layer make a Boolean measurement the presence of W -motif is necessary to prevent ideal information propagation. For more general measurements, for instance a sample from the exponential family of distribution, a nonlinear estimator would be needed, and the analysis becomes more complicated.

Related results have been obtained by Acemoglu, *et al.* [1] who considered social networks in which individuals receive information from a random neighborhood of agents. They show that agents can make the right choice, or infer the correct state of the world as network size increases when a finite group of agents does not account for most of the information that is propagated through the network. However, the setting of this study is somewhat different from ours: Agents are assumed to only observe each other's actions, but do not share their belief about the binary state of the world.

2.5. CONCLUSION

We translated the question about whether the estimate of the state of the world degrades across layers in the network to a simple algebraic condition. This allowed us to use results of random matrix theory in the case of random networks, find equivalent networks through an intuitive reduction process, and identify a class of networks in which estimates do not degrade across layers, and another class in which degradation is maximal.

Networks in which estimates degrade across layers must contain a W-motif. This motif introduces redundancies in the information that is communicated downstream and may not be removed. Such redundancies, also known as “bad correlations,” are known to limit the information that can be decoded from neural responses [39, 8]. This suggests that agents with large out-degrees and small in-degrees can hinder the propagation of information, as they introduce redundant information in the network. On the other hand, agents with large in-degrees integrate information from many sources, which can help improve the final estimate. However, the detailed structure of a network is important: For example, an agent with large in-degree in the second layer can have a large out-degree without hindering the propagation of information as it has already integrated most available first-layer measurements.

To make the problem tractable, we have made a number of simplifying assumptions. We made the strong assumption that agents have full knowledge of the network structure. Some agents may have to make several calculations in order to make an estimate, so we also do not assume bounded rationality [4]. This is unlikely to hold in realistic situations. Even when making simple decisions, pairs

2.5. CONCLUSION

of agents are not always rational [3]: When two agents each make a measurement with different variance, exchanging information can degrade the better estimate.

The assumption that only agents in the first layer make a measurement is not crucial. We can obtain similar results if all agents in the network make independent measurements, and the information is propagated directionally, as we assume here. However, in such cases, the confidence (inverse variance of the estimates) typically becomes unbounded across layers.

Asymptotics for Feedforward Networks

3.1 Variance and Bias of the Final Estimate

We next consider how the variance and bias of the estimate in layer n depend on the network structure. By definition, the variance of the ideal estimate is $\text{Var}(\hat{s}) = \left(\sum_{i=1}^{L_1} w_i\right)^{-1}$. If the variances of the individual estimates are bounded above as the size of the network increases, the final estimate in an ideal network is **consistent**: As the number of measurements increases the final estimate converges in probability to the true value of s [32]. We next show that the final estimate in non-ideal networks is not necessarily consistent. We also show that biases of certain first-layer agents can have a disproportionate impact on the bias of the final estimate.

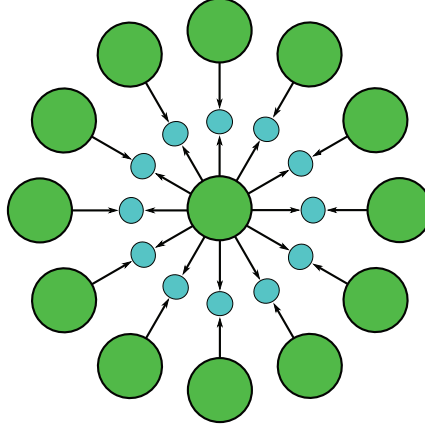


Fig. 3.1.1: Example of a network with an inconsistent final estimate. The green and blue nodes represent agents in the first and second layer, respectively. Each second-layer agent receives input from the common, central agent and a distinct first-layer agent, and thus $L_2 = L_1 - 1$.

Example (Variance Maximizing Network Structure) Fig. 3.1.1 shows an example of a network structure for which the variance of the final estimate converges to a positive number as the number of agents in the first layer increases. We assume that all first-layer agents make measurements with unit variance. We will show that as the number of agents in both layers increases, the variance of the final estimate approaches $1/4$. Let the estimate of the central agent be $s_1^{(1)}$. Then each agent in the second layer makes an estimate $\frac{1}{2}(s_1^{(1)} + s_i^{(1)})$ for some $i \neq 1$. By symmetry the single agent in the last layer averages all estimates from the second layer to obtain $\hat{s} = \frac{1}{2}(s_1^{(1)} + \frac{1}{L_1-1} \sum_{i=2}^{L_1} s_i^{(1)})$. Therefore, the estimate of the central agent (which communicates with all agents in the second layer) receives a much higher weight than all other estimates from the first layer. The variance of the final estimate thus equals

$$\text{Var}(\hat{s}) = \frac{1}{4} + \frac{1}{4(L_1 - 1)}.$$

3.1. VARIANCE AND BIAS OF THE FINAL ESTIMATE

Hence, the final estimate is not consistent, as its variance remains positive as the number of first-layer agents, L_1 , diverges. Given a restriction on the number of second-layer agents, we show that this network leads to the highest possible variance of the final estimate:

Proposition 4. *The final estimate in the network in Fig. 3.1.1 has the largest variance among all three-layer networks with a fixed number $L_1 \geq 4$ of first-layer, and $L_2 \geq L_1 - 1$ second-layer agents, assuming that every first-layer agent makes at least one connection.*

The idea of the proof is to limit the possible out-degrees of the agents in the first layer and show that the structure in Fig. 3.1.1 has the highest variance for this restriction.

Proof. We will show that the network architecture that maximizes the variance of the final estimate for a given number of first and second-layer agents is the one shown in Fig. 3.1.1. To simplify notation we write $L_1 = n$ and $L_2 = m$.

Lemma 4. *If $\mathbf{d} = (d_1, \dots, d_n)$ is the vector of out-degrees in the first layer, so $d_i = |f(a_i^{(1)})|$, then to maximize the variance of the final estimate, $\mathbf{d}$ must equal $(m, 1, \dots, 1)$, up to relabeling.*

Proof of Claim. Given a network structure consider the naïve estimate:

$$\frac{1}{Z} \sum_i |g(a_i^{(2)})| \hat{s}_i^{(2)} = \frac{1}{\sum_{ij} C_{ij}^{(1)}} \sum_i C_i^{(1)} \cdot \hat{\mathbf{s}}^{(1)}, \quad (3.1.1)$$

where Z is a normalizing factor that makes the entries of the corresponding vector of weights sum to 1. This estimate can always be made and is the same as using

3.1. VARIANCE AND BIAS OF THE FINAL ESTIMATE

a linear combination of estimates of agents $a_j^{(1)}$ with weights $\frac{d_i}{\sum_{j=1}^n d_j}$. Thus the variance of the optimal estimate of the agent in the final layer is bounded above by the variance of the naïve estimate in Eq. (3.1.1). By assumption $1 \leq d_j \leq m$ for all j . For the network in Fig. 3.1.1, this naïve estimate equals the final estimate. Thus it is sufficient to show that the naïve estimate has maximal variance when $\mathbf{d} = (m, 1, \dots, 1)$, up to relabeling.

The variance, V , of the naïve estimate is:

$$V(d_1, \dots, d_n) = \sum_j \left(\frac{d_j}{\sum_{k=1}^n d_k} \right)^2.$$

If we treat the degrees as continuous variables then V is continuous on $\mathbf{d} \in [1, m]^n$ and we can calculate the gradient of V to find the critical points.

$$\frac{\partial V}{\partial d_i} = 2 \left(\frac{d_i}{\sum_k d_k} \right) \frac{\sum_k d_k - d_i}{(\sum_k d_k)^2} + \sum_{j \neq i} 2 \left(\frac{d_j}{\sum_k d_k} \right) \frac{-d_j}{(\sum_k d_k)^2}$$

Setting $\frac{\partial V}{\partial d_i} = 0$ and multiplying both sides by $\frac{1}{2} (\sum_{k=1}^n d_k)^3$ gives

$$0 = d_i \left(\sum_{k \neq i} d_k \right) - \sum_{j \neq i} d_j^2 = \sum_{j \neq i} d_j (d_i - d_j).$$

This shows that $d = k\vec{1}$ for $k = 1, \dots, m$ are the only critical points, since if there exist $d_i \leq d_j$, for all $j \neq i$ and $d_i < d_k$ for some $k \neq i$ then the right hand side would be negative. These critical points are the first-layer out-degrees of ideal networks by Corollary 2, hence they are minima. This implies that V takes on its maximum values on the boundary.

The boundary of $[1, m]^n$ consists of points where at least one coordinate is 1 or m . Since V is invariant under permutation of the variables, we set d_1 equal to one of these values and investigate the behavior of V on this restricted set.

3.1. VARIANCE AND BIAS OF THE FINAL ESTIMATE

First set $d_1 = m$. Setting $\frac{\partial V}{\partial d_i}$ to 0 on this boundary gives:

$$0 = m(d_i - m) + \sum_{j \neq i, 1} d_j(d_i - d_j)$$

One critical point is thus $\vec{m}$. If $d_i \leq d_j$ for $j \neq i$ and $d_i < m$ then again the right hand side would be negative. Hence $d_i = m$ for all i , and there are no critical points on the interior of $\{m\} \times [1, d]^{n-1}$.

Next if $d_1 = 1$, setting $\frac{\partial V}{\partial d_i}$ to 0 on this boundary and multiplying by -1 gives:

$$0 = 1 - d_i + \sum_{j \neq i, 1} d_j(d_j - d_i)$$

Here a critical point is $\vec{1}$. If $d_i \leq d_j$ for $j \neq i$ and $1 < d_i < m$ then again the right hand side would be negative. Hence $d_i = 1$ for all i , and there are no critical points on the interior of $\{1\} \times [1, d]^{n-1}$. If we iterate this procedure, we see that the maximum value of V must occur on the corners of the hypercube $[1, d]^n$.

Choose one of these corners, $\mathbf{c}$, and, without loss of generality, assume that the first l coordinates are m and the last $n - l$ coordinates are 1, $1 \leq l < n$. Then

$$\begin{aligned} V(\mathbf{c}) &= \sum_{j=1}^l \left(\frac{m}{\sum_{k=1}^n d_k} \right)^2 + \sum_{j=l+1}^n \left(\frac{1}{\sum_{k=1}^n d_k} \right)^2 \\ &= \left(\frac{1}{lm + (n-l)} \right)^2 (lm^2 + (n-l)) \\ &= \frac{lm^2 + n - l}{l^2m^2 + 2lm(n-l) + (n-l)^2} \\ &= \frac{l(m^2 - 1) + n}{l^2(m-1)^2 + l2n(m-1) + n^2} \end{aligned}$$

Under the assumption that $m \geq n - 1$, a lengthy algebra calculation that we omit shows that this is maximized for $l = 1$. Hence the maximum value of V is achieved at $(m, 1, \dots, 1)$, or any of its coordinate permutations. $\square$

3.1. VARIANCE AND BIAS OF THE FINAL ESTIMATE

Finally, to have $\mathbf{d} = (m, 1, \dots, 1)$, one first-layer agent, $a_1^{(1)}$, communicates with all second-layer agents and every other agent has exactly one output. Since there are at least $n - 1$ agents in the second layer, this means that each first-layer agent must communicate with a distinct second-layer agent and each second-layer agent must receive input from $a_1^{(1)}$. Otherwise, some agent in the second layer would receive only the input from $a_i^{(1)}$ and thus the final estimate could use that estimate to decorrelate all of the second-layer estimates.

So, the naïve estimate for an alternative network has smaller variance than the ideal estimate for the ring network in Fig. 3.1.1. Hence the final estimate in any alternative network will have smaller variance. Since the only network with $\mathbf{d} = (m, 1, \dots, 1)$ is the network in Fig. 3.1.1, we have shown that this structure maximizes the variance of the final estimate among all networks with $L_2 \geq L_1 - 1$. $\square$

In general, we conjecture that for the final estimate to have large variance, some agents upstream must have a disproportionately large out-degree, with the remaining agents making few connections. On the other hand, as the in-degree of a second-layer agent increases, the variance of its estimate shrinks. Thus when a few agents communicate information to many, the resulting redundancy is difficult to resolve downstream. But when downstream agents receive many estimates, we expect the estimates to be good. We next show that the biases of the agents with the highest out-degrees can have an outsized influence on the estimates downstream.

3.1. VARIANCE AND BIAS OF THE FINAL ESTIMATE

Propagation of Biases We next ask how biases in the measurements of agents in the first layer propagate through the network. Ideally, such biases would be averaged out in subsequent layers. To simplify the analysis we assume constant, additive biases, $\hat{s}_i^{(1)} = x_i + b_i$, with the constant bias, b_i . Downstream agents are unaware of these biases, and therefore assume them to be zero. Since all estimates in the network are convex linear combinations of first-layer measurements, the final estimate will have the form

$$\hat{s} = \sum \alpha_i (x_i + b_i) = \sum \alpha_i x_i + \sum \alpha_i b_i, \quad (3.1.2)$$

and thus will have finite bias bounded by the maximum of the individual biases.

We have provided examples of network structures where the estimate of a first-layer agent was given higher weight than others, even when all first-layer measurements had equal variance. Eq. (3.1.2) shows that this agent's bias will also be disproportionately represented in the bias of the final estimate. Indeed, in the example in Fig. 2.1.1(a), the estimate of second agent in first layer has weight $\frac{1}{2}$, and its bias will have twice the weight of the other agents in the final estimate. Similarly, the bias of the central agent in Fig. 3.1.1 will account for half the bias of the final estimate as $n \rightarrow \infty$. Thus even if the biases, b_i , are distributed randomly with zero mean, the asymptotic bias of the final estimate does not always disappear as the number of measurements increases.

More generally, networks that contain W-motifs can result in biases of first-layer agents with disproportionate impact on the final estimate. As with the variance, we conjecture that the bias of agents that communicate their estimates to

many agents downstream will be disproportionately represented in the final estimate. Equivalently, if the network contains agents that receive many estimates, we expect the bias of the final estimate to be reduced.

3.2 Inference in random feedforward networks

We have shown that networks with specific structures can lead to inconsistent and asymptotically biased final estimates. We now consider networks with randomly and independently chosen connections between layers. Such networks are likely to contain many W -motifs, but it is unclear whether these motifs are resolved and whether the final estimate is ideal. We will use results of random matrix theory to show that there is a sharp transition in the probability that a network is ideal when the number of agents from one layer exceeds that of the previous layer [14].

We assume that connections between agents in different layers are random, independent and made with fixed probability, p . We will use the following result of [35], also discussed by [14]:

Theorem 2 (Komlos). *Let ξ_{ij} , $i, j = 1, \dots, n$ be i.i.d. with non-degenerate distribution function $F(x)$. Then the probability that the matrix $X = (\xi_{ij})$ is singular converges to 0 with the size of the matrix,*

$$\lim_{n \rightarrow \infty} P(\det X = 0) = 0.$$

Corollary 3. *For a three-layer network with independent, random, equally probable ($p =$*

3.2. INFERENCE IN RANDOM FEEDFORWARD NETWORKS

1/2) connections from first to second-layer, as the number of agents L_1 and L_2 increases,

$$\frac{L_1}{L_2} \leq 1 \implies P(\hat{s} = \hat{s}_{ideal}) \rightarrow 1,$$

and

$$\frac{L_1}{L_2} > 1 \implies P(\hat{s} = \hat{s}_{ideal}) \rightarrow 0.$$

Proof. Whether or not $\hat{s}_{ideal} = \hat{s}$ is determined by $C^{(1)}$. For simplicity, we drop the superscript and refer to this connectivity matrix as C . By our assumption, this is a random matrix with $P(C_{ij} = 0) = P(C_{ij} = 1) = 1/2$.

First assume that there are at least as many second-layer agents as there are first-layer agents: $L_2 \geq L_1$ or $\frac{L_1}{L_2} \leq 1$. Then C is a random $L_2 \times L_1$ matrix with i.i.d. non-degenerate entries that has more rows than columns. By Theorem 2, this means that the $L_1 \times L_1$ submatrix formed by the first L_1 rows and columns is nonsingular with probability approaching 1 as $L_1, L_2 \rightarrow \infty$. Thus the probability that the row space of C contains the vector $\vec{1}$ converges to 1 with the size of the network.

Next assume that there are fewer second-layer agents than first-layer agents, that is $L_2 < L_1$ or $\frac{L_1}{L_2} > 1$. We will show that the probability that the row space of C contains $\vec{1}$ goes to zero as $L_1, L_2 \rightarrow \infty$. Since increasing the number of rows will not decrease the probability that C contains a vector in its row space we assume that $L_2 = L_1 - 1$ and let $L_1 = n$:

$$\lim_{L_1, L_2 \rightarrow \infty} P(\hat{s} = \hat{s}_{ideal}) \leq \lim_{n \rightarrow \infty} P(\vec{1} \in R(C(n-1, n)))$$

3.2. INFERENCE IN RANDOM FEEDFORWARD NETWORKS

where $C(n-1, n)$ refers to the random matrix as before, and identifies that it has $n-1$ rows and n columns. We first use:

$$P(\vec{1} \in R(C(n-1, n))) \leq P\left(\begin{pmatrix} \vec{1} \\ C \end{pmatrix} \text{ is singular}\right)$$

since if $\vec{1}$ is the row space of C , then attaching that row of ones to it would create a singular matrix.

Lemma 1. $P\left(\det\left(\begin{pmatrix} \vec{1} \\ C \end{pmatrix}\right) = 0\right) \rightarrow 0 \text{ as } n \rightarrow \infty.$

We can rewrite $C = \begin{pmatrix} B & \mathbf{v} \end{pmatrix}$, where $\mathbf{v}$ is the n^{th} column of C and B is the remaining submatrix. We claim

$$\det\left(\begin{pmatrix} \vec{1} \\ C \end{pmatrix}\right) = -1^k \det\left(\begin{pmatrix} \vec{1} & 1 \\ \tilde{B} & \vec{0} \end{pmatrix}\right) = -1^{k+n+1} * \det(\tilde{B}) \quad (3.2.1)$$

where $\tilde{B}$ is a random $(n-1) \times (n-1)$ matrix distributed like C . Assuming this claim, then by [35] :

$$P\left(\det\left(\begin{pmatrix} \vec{1} \\ C \end{pmatrix} = 0\right)\right) = P(\det(\tilde{B}) = 0) \rightarrow 0 \text{ as } n \rightarrow \infty.$$

Thus $P(\vec{1} \in R(M(n-1, n))) \rightarrow 0 \text{ as } n \rightarrow \infty.$

To prove the first equality in Eq. (3.2.1), we use row operations on $\begin{pmatrix} \vec{1} & 1 \\ B & \mathbf{v} \end{pmatrix}$: If $v_i = 1$ then subtract the first row from the i^{th} row, $(B_i \ v_i)$, to get a vector whose entries are all 0 and -1 . Then $(B_i \ v_i) \rightarrow -(\tilde{B}_i \ 0)$ where $(\tilde{B}_i \ 0)$ is a vector of

3.2. INFERENCE IN RANDOM FEEDFORWARD NETWORKS

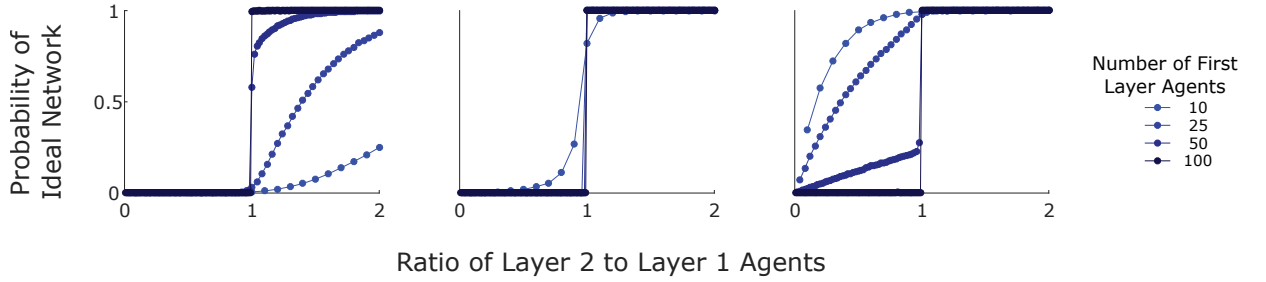


Fig. 3.2.1: The probability that a random, three-layer network is ideal for connection probabilities $p = 0.1$ (left), 0.5 (center), and 0.9 (right). In each panel, the different curves correspond to different, but fixed numbers of agents in the first layer. The number of agents in the second layer is varied. There is a sharp transition in the probability that a network is ideal when the number of agents in the second layer exceeds the number in the first. Simulation details can be found in Section 3.2.2.

entries which are again either 0 or 1 with equal probability. We do this for every row which has a 1 in its last entry and multiply the determinant a factor -1 and denote the number of these reductions as k . Since $P(C_{ij} = 0) = \frac{1}{2}$ we also have $P(\tilde{B}_{ij} = 0) = \frac{1}{2}$. $\square$

The same proof works when $L_1/L_2 \leq 1$ and the probability of a connection is arbitrary, $p \in (0, 1]$. We conjecture that the result also holds for $L_1/L_2 > 1$ and arbitrary p , but the present proof relies on the assumption that $p = 1/2$. Fig. 3.2.1 shows the results of simulations which support this conjecture: The different panels correspond to different connection probabilities, and the curves to different numbers of agents in the first layer. As the number of agents in the second layer exceeds that in the first, the probability that the network is ideal approaches 1 as the number first-layer agents increases. With 100 agents in the first layer, the curve is approximately a step function for all connection probabilities we tested.

3.2.1 More than 3 Layers

We conjecture that a similar result holds for networks with more than three layers:

Conjecture. *For a network with n layers with independent, random, equally probable connections between consecutive layers, as the total number of agents increases,*

$$L_k \leq L_{k+1} \text{ for } 1 \leq k < n - 1 \implies P(\hat{s} = \hat{s}_{ideal}) \rightarrow 1$$

and

$$L_1 > L_k \text{ for some } 1 < k < n \implies P(\hat{s} = \hat{s}_{ideal}) \rightarrow 0.$$

Evidence for the conjecture can be found in Fig. 3.2.2, which shows the results with four-layer networks with different connection probabilities across layers. The number of agents in the first and second layers are equal, and we varied the number of agents in the third layer. The results support our conjecture.

With multiple layers ($n \geq 4$), if $L_1 > L_2$ then the network will not be ideal as in the limit the estimate of s will not be ideal already in the second layer by Corollary 3. If the number of agents does not decrease across layers, we conjecture that the probability that information is lost across layers is small when the number of agents is large. Indeed, it seems reasonable that the products of the random weight matrices will be full rank with increasing probability allowing us to apply Proposition 2. However, the entries in these matrices are no longer independent, so classical results of random matrix theory no longer apply.

3.3. CONCLUSION

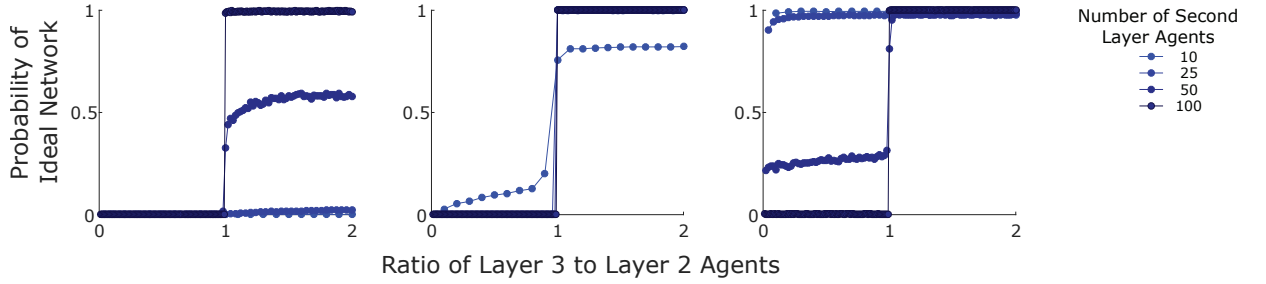


Fig. 3.2.2: The probability that a random, four-layer network is ideal for connection probabilities $p = 0.1$ (left), 0.5 (center), and 0.9 (right). Each curve corresponds to equal, fixed numbers of agents in the first two layers, with a changing number of agents in the third layer. Simulation details can be found in Section 3.2.2.

3.2.2 Simulation Details

All simulations were done in MATLAB. For the 3-layer networks we randomly generated binary connection matrices and tested whether or not the vector $\vec{1}$ was in the row space. Each point in the plots corresponds to the number of agents in the first two layers for a given connection probability and was generated using at least 10,000 samples. The code used for these simulations can be found at the repository <https://github.com/Spstolar/FFNetInfoLoss>.

3.3 Conclusion

Here we considered asymptotic results for the information sharing process detailed in Chapter 2. We gave conditions for when ideal networks were generic by showing the dependence on the ratio of the layer sizes and detailed the worst-case

3.3. CONCLUSION

scenario: an arbitrarily large network for which the final estimate is not consistent. In the viewpoint of asymptotic learning on networks, we showed something analogous of asymptotic learning: rather than the probability of making the exact correct estimate going to 1, we considered how the variance of that estimate shrunk. If the network is not consistent, then the variance does not go to 0 and in some sense the network does not exhibit asymptotic learning. However, if it is consistent, which we showed is a generic condition for three-layer networks where the second layer is larger than the first, then asymptotic learning occurs, because the variance of the final agent *does* go to 0.

Chapter 4

Evidence Accumulation on Networks

In the previous chapters we have assumed that agents make a single observation or measurement of a parameter (state of the world), and communicate information about this measurement to their neighbors. In many situations agents can integrate information from multiple private measurements to make an estimate or reach a decision. In an uncertain environment, agents on a network can use this private information along with the social information obtained from their neighbors. Intuitively, if the measurements of the agents in the network are conditionally independent on the state of the world, then the sharing of private information should lead to a better estimate or decision. This is especially important in scenarios like predator detection, where animals are trying to determine if there is a threat nearby [51, 18]. The group has a better chance of escaping a predator if the individuals within it observe each other's actions or communicate information to one another.

How to accumulate evidence from a sequence of measurements to decide between two choices is one of the fundamental problems in decision theory [55, 20]. It has long been known that optimal decisions can be based on the log-likelihood ratio of the posterior probabilities of the two choices given a sequence of measurements [55]. When there are many observations, each providing little evidence, this process is well-approximated by a drift-diffusion process [13]. Decisions that provide the best balance between speed and accuracy can be made by choosing an optimal threshold which, once crossed by the log-likelihood, triggers a choice.

A group of agents that makes sequential measurements and communicates the obtained private information to their neighbors can be modeled heuristically by a set of diffusively coupled drift-diffusion equations [51, 48]. However, there are two issues with this approach: First, it is not always clear whether the linear coupling between neighbors used in such models approximates the accumulation of evidence by agents that optimally use all available information (rational agents). Second, and more importantly for the discussion that follows, in many situations agents do not share their exact beliefs, or each individual measurement and observation with their neighbors. Rather, the agents may only observe each other's actions or decisions. While such actions do reveal the belief of each agent, they occur much less frequently than the individual measurements. Furthermore, an agent will typically not see the actions of all other agents in the network, which means it will have to take unobserved possibilities into account. This is in contrast to the model discussed in [33], where there is a global agent which sees all actions and then optimizes a final decision based on a group of independent agents.

Here we take a Bayesian approach and investigate the behavior of rational agents who can only observe each other's decisions or actions. We assume that these decisions occur at a single moment in time. Thus the social information that is communicated by each action is localized in time. The decision itself informs all those observing the acting agent about its belief. Interestingly, we will show that in some situations, even the absence of a decision or action can communicate information. We will begin by reviewing the drift-diffusion model in the single agent case. We next extend the model to a pair of agents, and continue with more general networks in the next chapter. We show how decision thresholds (which determine when an agent has sufficient information to act), network connectivity, and assumptions about boundedly rational agents affect group evidence accumulation.

Throughout we will use the terms individual, agent, and observer interchangeably.

4.1 Single-Agent Setup

The two-alternative forced choice task has been thoroughly studied for a single observer [13]. Here an observer is tasked with deciding on the true state of the world which can take one of two values $H \in \{\pm 1\}$. Equivalently, the observer needs to decide which of two hypotheses is true. To do so, the observer makes a sequence of noisy observations, ξ_t . The observations lie in a set Ξ , and in what follows we will assume that $\Xi = \mathbb{R}$, or that Ξ is a discrete subset of $\mathbb{R}$. We will assume

4.1. SINGLE-AGENT SETUP

that the observations are independent and identically distributed conditioned on the state of the world, H . For two observations this means

$$P(\xi_1, \xi_2 | H) = P(\xi_1 | H)P(\xi_2 | H) = f_H(\xi_1)f_H(\xi_2),$$

where $f_+(\xi) = P(\xi | H = 1)$ and $f_-(\xi) = P(\xi | H = -1)$ are the conditional distributions of the observations or measurements, which we will call the **evidence distributions**. We denote the set of observations until time T by $\xi_{0:T}$. We denote by I the totality of information about H available to an agent. This is initially $I = \xi_{0:T}$, but will be more general later, when agents begin to incorporate evidence (social information) from observing the decisions of their neighbors.

Given information, I , the agent compares $P(H = 1 | I)$ to $P(H = -1 | I)$. How a choice is made depends on the setup: An agent can be asked which state is more likely at some point in time (interrogation paradigm), or an agent can be allowed to freely make a selection once they have a sufficient amount of evidence in favor of one of the options (free response paradigm). Under the interrogation paradigm, the agent is allowed a fixed number of observations, or a fixed amount of time, to reach a decision. If we denote this set of observations by $\xi_{0:T}$, at the end of the given time a rational agent will choose

$$\arg \max_H P(H | \xi_{0:T}).$$

Here, we will focus on the free response paradigm. A possible utility function could be simply the probability of being correct. However, this would incentivize an agent to wait infinitely long before deciding. Thus we assume our rational agents choose $H = 1$ when $\log \frac{P(H=1|I)}{P(H=-1|I)}$ is sufficiently large, and $H = -1$ when

4.1. SINGLE-AGENT SETUP

it is sufficiently small. More precisely, we assume that an agent wants to exceed an accuracy level α . Therefore, if the goal is to be correct more than 95% of the time, then $\alpha = 0.95$. To achieve this, the agent can check when $\log \frac{P(H=1|I)}{P(H=-1|I)}$ exceeds the threshold $\theta_+ = \log \frac{\alpha}{1-\alpha}$ and select $H = 1$, and similarly select $H = -1$ when the quantity falls below $\theta_- = \log \frac{1-\alpha}{\alpha}$. In this case we say the thresholds are **symmetric**, since $\theta_- = -\theta_+$. We will extend this and more generally allow the thresholds to be asymmetric, so that we do not necessarily have $|\theta_-| = |\theta_+|$. Thus we generally assume the agent is using thresholds $\theta_- < 0 < \theta_+$.

The agent knows both evidence distributions f_- and f_+ , corresponding to $H = -1$ and $H = 1$, respectively, but not which one is being sampled from. An agent computes how likely it was that an observation came from either distribution and applies the optimal Sequential Probability Ratio Test [55]. To compute the posterior probability $P(H|\xi_t)$ agents use Bayes' Rule to compute $P(\xi_t|H)$, that is, the probability that the observation was generated from one of the two evidence distributions. For simplicity we assume a flat prior, $P(H = 1) = P(H = -1)$, but our arguments are easily extended to the case of unequal priors.

An agent making a single observation at time t applies Bayes' Rule and computes:

$$\begin{aligned} \log \left(\frac{P(H = 1|\xi_t)}{P(H = -1|\xi_t)} \right) &= \log \left(\frac{P(\xi_t|H = 1)P(H = 1)P(\xi_t)}{P(\xi_t|H = -1)P(H = -1)P(\xi_t)} \right) \\ &= \log \left(\frac{P(\xi_t|H = 1)}{P(\xi_t|H = -1)} \right). \end{aligned} \quad (4.1.1)$$

Conditioned on the true state, H , the observations are independent. Thus, if we define y_t as the log-likelihood ratio at time t , which is a measure of the agent's

belief at time t , we have

$$\begin{aligned}
 y_t &:= \log \left(\frac{P(H = 1 | \xi_{0:t})}{P(H = -1 | \xi_{0:t})} \right) \\
 &= \log \left(\frac{P(\xi_{0:t} | H = 1) P(H = 1) P(\xi_{0:t})}{P(\xi_{0:t} | H = -1) P(H = -1) P(\xi_{0:t})} \right) \\
 &= \log \left(\frac{P(\xi_{0:t} | H = 1)}{P(\xi_{0:t} | H = -1)} \right) \\
 &= \log \left(\frac{\prod_{s=0}^t P(\xi_s | H = 1)}{\prod_{s=0}^t P(\xi_s | H = -1)} \right) \\
 &= \log \left(\prod_{s=0}^t \frac{P(\xi_s | H = 1)}{P(\xi_s | H = -1)} \right) \\
 &= \sum_{s=0}^t \log \left(\frac{P(\xi_s | H = 1)}{P(\xi_s | H = -1)} \right).
 \end{aligned}$$

Hence to get the log-likelihood ratio at time T , the agent adds the result of (4.1.1) to a running total:

$$y_t = y_{t-1} + \log \left(\frac{P(\xi_t | H = 1)}{P(\xi_t | H = -1)} \right),$$

where

$$y_0 = \log \left(\frac{P(\xi_0 | H = 1)}{P(\xi_0 | H = -1)} \right)$$

is the evidence from the initial observation. The agent continues until the sum reaches one of the two pre-determined decision thresholds, $\theta_- < 0 < \theta_+$, and then chooses $H = -1$ if the evidence $y_t \leq \theta_-$ or $H = 1$ if $y_t \geq \theta_+$. An example of what the process look like is given in Figure 4.1.1.

Note that this process can be extended to more than two possible choices. Optimal decision strategies with more than two alternatives are harder to define and analyze, as the threshold crossing procedure we use here is not optimal [38]. However, many interesting scenarios (such as predator detection) are modeled with

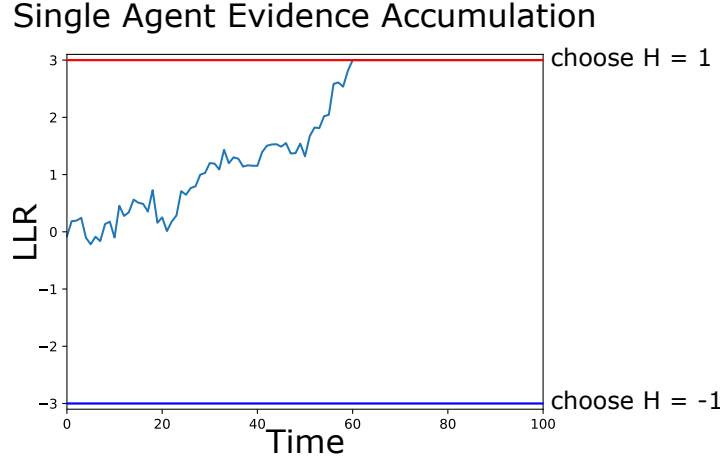


Fig. 4.1.1: An example of the evolution of the log-likelihood ratio when the evidence distributions are normal with means ± 0.1 corresponding to $H = \pm 1$ and equal variance 1. The true state is $H = 1$, the thresholds are $\theta_- = -3$ and $\theta_+ = 3$, and the agent eventually accumulates enough evidence to make a correct decision.

just two alternatives.

4.2 Two-Agent Setup

Our goal is to extend the single-agent model of evidence integration to a directed network with N agents, where the direction of the edges will correspond to the flow of information. Each agent has the same goal as a single observer: determine the true state of the world, $H \in \{-1, 1\}$. We will assume that agents are making a series of noisy observations of this unknown state. These observations are identically distributed, and independent from each other in time and between agents, conditioned on the true state of the world, H . As in the case of a single agent, we assume that each agent makes a decision and communicates this decision exactly when its log-likelihood ratio crosses a threshold.

4.2. TWO-AGENT SETUP

Unlike an isolated observer, agents on a network receive additional social information when their neighbors decide or act. We assume that after every observation, agents communicate whether they have integrated sufficient evidence to make a decision, in which case they communicate their choice to all their neighbors downstream, according to the network topology. When they have not accumulated enough evidence to reach a decision they do not communicate anything. Thus the absence of a decision or an action communicates that an agent has not gathered sufficient evidence to decide. This is equivalent to assuming that agents observe the actions of all of their upstream neighbors, and that each action is determined completely by an agent's belief about the state of the world. We will show how an ideal observer integrates its own observations along with social information communicated by its upstream neighbors to reach a decision.

We begin with the simplest possible setup with two agents and social information flowing in one direction. We use this simple example to show how decision thresholds determine the evidence accumulation dynamics, and derive a continuum limit of the process. In the next chapter, we then move to the case where two agents are exchanging social information, and each is observing the actions of the other. We will investigate how this bidirectional flow introduces additional complications in the analysis, which are absent in the following setup.

We start with the simplest nontrivial network, depicted in Figure 4.2.1.



Fig. 4.2.1: A pair of agents with unidirectional coupling.

4.2. TWO-AGENT SETUP

We assume both agents accumulate evidence about H : For agent 1 this information consists entirely of its own observations. Agent 2 makes its own observations, and obtains social information from agent 1. Let $I_t^{(i)}$ be the total information available to agent i at time t . For instance, $I_t^{(2)}$ consists of all private observations of agent 2, and the history of decisions of agent 1 up to time t . We denote the decision boundaries for both agents by $\theta_- < 0 < \theta_+$. We assume that when the evidence causes the log-likelihood ratio,

$$y_t^{(i)} = \log \frac{P(H = 1 | I_t^{(i)})}{P(H = -1 | I_t^{(i)})},$$

to exceed θ_+ , agent i chooses $H = 1$ and stops collecting information. Equivalently, when the log-likelihood ratio, $y_t^{(i)}$, falls below θ_- the agent chooses $H = -1$ and stops collecting information. If neither condition is satisfied, the agent continues accumulating evidence. We refer to $y_t^{(i)}$ as the **belief** of agent i at time t . We will always assume that the agents have the same decision boundaries, θ_- and θ_+ , but this assumption can be relaxed.

Agent 1 makes only private observations, so behaves as a lone observer. At each time step, t , agent 1 makes an observation, $\xi_t^{(1)} \in \Xi$, which is a sample from $P(\xi | H)$ with H the true state of the world. Then the agent updates its belief, $y_t^{(1)}$, equivalently to an isolated observer:

$$y_t^{(1)} = y_{t-1}^{(1)} + \log \left(\frac{P(\xi_t^{(1)} | H = 1)}{P(\xi_t^{(1)} | H = -1)} \right),$$

and

$$y_0^{(1)} = \log \left(\frac{P(\xi_0 | H = 1)}{P(\xi_0 | H = -1)} \right).$$

4.2. TWO-AGENT SETUP

Next, the agent communicates to agent 2 the **decision state** $d_t^{(1)}$ where

$$d_t^{(1)} = \begin{cases} -1, & y_t^{(1)} \leq \theta^- \\ 0, & \log \frac{P(\xi_t^{(1)}|H=1)}{P(\xi_t^{(1)}|H=-1)} \in (\theta^-, \theta^+) \\ 1, & y_t^{(1)} \geq \theta^+ \end{cases}$$

So if agent 1 has sufficient evidence to make a decision, it communicates its belief about H . Otherwise the agent indicates that it is still undecided by communicating 0. We emphasize that the absence of a decision can provide information about the belief of an agent.

At each time step, t , agent 2 makes an observation $\xi_t^{(2)} \in \Xi$ from the same distribution as agent 1: $P(\xi|H)$, and updates its belief before receiving the latest piece of information from agent 1:

$$\tilde{y}_t^{(2)} = \log \frac{P(H = 1|\xi_{0:t}^{(2)}, d_0^{(1)}, \dots, d_{t-1}^{(1)})}{P(H = -1|\xi_{0:t}^{(2)}, d_0^{(1)}, \dots, d_{t-1}^{(1)})}.$$

If this evidence is sufficient to make a decision, then agent 2 stops here. Otherwise, it receives the communicated decision, $d_t^{(1)}$, from agent 1 and updates its log-likelihood ratio again:

$$y_t^{(2)} = \log \frac{P(H = 1|\xi_{0:t}^{(2)}, d_0^{(1)}, \dots, d_{t-1}^{(1)}, d_t^{(1)})}{P(H = -1|\xi_{0:t}^{(2)}, d_0^{(1)}, \dots, d_{t-1}^{(1)}, d_t^{(1)})}.$$

Agents continue this process until they have both made a decision.

As stated before, it is important that we assume both agents know that they share the same distribution of observations, $P(\xi|H = 1)$ and $P(\xi|H = -1)$. Alternatively, we could assume that agents know each other's measurement distributions, but the notation becomes more cumbersome. We also assume that both

4.2. TWO-AGENT SETUP

agents know they are acting rationally, and that once an agent has made a decision, it cannot change it. Agents make a decision as soon as they have sufficient evidence, that is, once the log-likelihood ratio given all the evidence exceeds one of the two thresholds.

For agent 1, the process is identical to the case of a single observer, except that it is also communicating its decision state, $d_t^{(1)}$, at every time step. However, the second agent has additional information. We next show that the log-likelihood ratio of the two choices can be separated into parts corresponding to evidence from private observations and evidence from agent 1's decisions.

Consider the following computation that agent 2 makes after a private measurement at step $t = 1$, and having observed the decision of agent one at $t = 0$,

$$\begin{aligned}
 \log \frac{P(H = 1 | I_1^{(2)})}{P(H = -1 | I_1^{(2)})} &= \log \left(\frac{P(H = 1 | \xi_{0:1}^{(2)}, d_0^{(1)})}{P(H = -1 | \xi_{0:1}^{(2)}, d_0^{(1)})} \right) \\
 &= \log \left(\frac{P(\xi_{0:1}^{(2)}, d_0^{(1)} | H = 1)}{P(\xi_{0:1}^{(2)}, d_0^{(1)} | H = -1)} \right) \\
 &= \log \left(\frac{P(\xi_{0:1}^{(2)} | H = 1) P(d_0^{(1)} | H = 1)}{P(\xi_{0:1}^{(2)} | H = -1) P(d_0^{(1)} | H = -1)} \right) \\
 &= \log \left(\frac{P(\xi_0^{(2)} | H = 1) P(\xi_1^{(2)} | H = 1) P(d_0^{(1)} | H = 1)}{P(\xi_0^{(2)} | H = -1) P(\xi_1^{(2)} | H = -1) P(d_0^{(1)} | H = -1)} \right) \\
 &= \sum_{t=0}^1 \log \left(\frac{P(\xi_t^{(2)} | H = 1)}{P(\xi_t^{(2)} | H = -1)} \right) + \log \left(\frac{P(d_0^{(1)} | H = 1)}{P(d_0^{(1)} | H = -1)} \right)
 \end{aligned}$$

The first equation follows from Bayes' rule, assuming equal prior probabilities over the two states. The second and third equality follow from the assumption

4.2. TWO-AGENT SETUP

that the observations across time and between agents are conditionally independent. Thus the second agent's log-likelihood ratio splits into a sum of the log-likelihood ratio change due solely to observations (private information), and a log-likelihood ratio obtained from observations of agent 1. Agent 2 then uses the result of this computation to update its belief at $t = 1$.

Remark. Even though $P(\xi_0^{(2)}, \xi_1^{(2)}, d_0^{(1)} | H = 1)$ can be written as the product of $P(\xi_0^{(2)} | H = 1)$, $P(\xi_1^{(2)} | H = 1)$, and $P(d_0^{(1)} | H = 1)$ it is not necessarily the case that $P(\xi_0^{(2)}, \xi_1^{(2)}, d_0^{(1)}) = P(\xi_0^{(2)})P(\xi_1^{(2)})P(d_0^{(1)})$. Observations are only independent from each other, and the decisions of the other agent, when conditioned on the state of the world.

Because agent 2 knows the conditional distribution of measurements and assumes the first agent is acting optimally, it knows that the value of agent 1's first decision state, $d_0^{(1)}$, indicates something about the first agent's belief, $y_0^{(1)}$. At this first step, agent one has only reached a decision if its single observation provided a sufficient amount of evidence for $H = 1$ or $H = -1$.

The following proposition shows that the belief of the second agent can be split into two parts at any other point in time. For simplicity, we will use $\text{OB}_{0:t}^{(j)}$ to denote the evidence agent j has received from its own observations up to time t :

$$\text{OB}_{0:t}^{(j)} = \sum_{l=0}^t \log \frac{P(\xi_l^{(j)} | H = 1)}{P(\xi_l^{(j)} | H = -1)} \quad (\text{Observation Evidence})$$

and $\text{DEC}(d_{0:t}^{(1)})$ to denote the evidence agent 2 gets from knowing the decision states of agent 1 up to time t :

$$\text{DEC}(d_{0:t}^{(1)}) = \log \frac{P(d_{0:t}^{(1)} | H = 1)}{P(d_{0:t}^{(1)} | H = -1)}. \quad (\text{Decision Evidence})$$

4.2. TWO-AGENT SETUP

Proposition 1. *Assume that in the network depicted in Fig. 4.2.1 agent 1 chooses $H = 1$ at time T . Then, if agent 2 has not yet made a decision, its belief can be written as a sum of observational and decisional evidence acquired up to time t :*

$$y_s^{(2)} = \begin{cases} \text{OB}_{0:s}^{(2)} + \text{DEC}(d_s^{(1)} = 0) & , s < T \\ \text{OB}_{0:s}^{(2)} + \text{DEC}(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1) & , s \geq T \end{cases} \quad (4.2.1)$$

Proof. The decisions of agent 1 are independent from the observations of agent 2, when conditioned on the state. Thus

$$P(\xi_{0:s}^{(2)}, d_{0:s}^{(1)} | H = 1) = P(\xi_{0:s}^{(2)} | H = 1) P(d_{0:s}^{(1)} | H = 1).$$

The same thing holds for conditioning on $H = -1$, hence taking the log ratio gives

$$y_s^{(2)} = \text{OB}_{0:s}^{(2)} + \text{DEC}(d_{0:s}^{(1)})$$

thus it just remains to show that $\text{DEC}(d_{0:s}^{(1)})$ simplifies as shown in (4.2.1).

Consider $P(d_0^{(1)} = i_1, \dots, d_s^{(1)} = i_s | H)$ where $i_1, \dots, i_s \in \{-1, 0, 1\}$. Before agent 1 makes a decision, it communicates a decision state of 0, so if $s < T$:

$$P(d_0^{(1)} = 0, \dots, d_s^{(1)} = 0 | H) = P(d_s^{(1)} = 0 | H)$$

because once an agent makes a decision it cannot change it.

Similarly, when agent 1 chooses state $H = 1$ at time T , we can write

$$P(d_0^{(1)} = 0, \dots, d_{T-1}^{(1)} = 0, d_T^{(1)} = 1, \dots, d_s^{(1)} = 1 | H).$$

Note $d_0^{(1)} = 0, \dots, d_{T-2}^{(1)} = 0$ is implied by $d_{T-1}^{(1)} = 0$ and the values of the decision states after time T , $d_{T+1}^{(1)} = 1, \dots, d_s^{(1)} = 1$ are implied by $d_T^{(1)} = 1$. Thus

$$\text{DEC}(d_0^{(1)} = 0, \dots, d_{T-1}^{(1)} = 0, d_T^{(1)} = 1, \dots, d_t^{(1)} = 1) = \text{DEC}(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1)$$

and we note that the evidence from the decision state depends on the value of the decision state and the time when it was first non-zero, i.e., when agent 1 made a choice. $\square$

4.3 Thresholds and Decision Evidence

First, we will see what evidence a decision, $d_t^{(1)} = \pm 1$, provides. We will then show that when $|\theta_-| = |\theta_+|$, non-decisions, $d_t^{(1)} = 0$, are uninformative, provided an assumption on the evidence distributions.

4.3.1 Decision Evidence

At each point in time each agent computes its belief, a log-likelihood ratio (LLR), $y_t^{(i)}, i = 1, 2$. As we have shown, since the agents are making independent observations, $y_t^{(2)}$ can be broken into two parts:

$$y_t^{(2)} = \text{OB}_{0:t}^{(2)} + \text{DEC}(d_0^{(1)}, \dots, d_t^{(1)}).$$

After agent 1 has made a decision, the evidence from the decision state remains fixed. Assume agent 1 chooses $H = 1$ at time T . Then the belief of agent 2 at time $t \geq T$ is:

$$y_t^{(2)} = \text{OB}_{0:t}^{(2)} + \log \left(\frac{P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1)}{P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = -1)} \right).$$

Hence if we can compute $P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = \pm 1)$, we can determine the evidence provided by a decision from agent 1. Intuitively, if agent 1 communicates

4.3. THRESHOLDS AND DECISION EVIDENCE

the decision $d_T^{(1)} = 1$, then agent 2 knows that $y_T^{(1)} \geq \theta_+$. Since this information was acquired through independent observations, agent 2 can use all of it to update its belief:

Proposition 5. *A choice of $H = 1$ at time T by agent 1 results in an increase of at least θ_+ in the belief of agent 2, since*

$$\log \left(\frac{P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1)}{P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = -1)} \right) \geq \theta_+.$$

If the evidence distributions $f_{\pm}$ are sufficiently close, so that individual observations small amounts of evidence on average, then:

$$\text{DEC}(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1) \approx \theta_+.$$

Proof. We have

$$P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1) = P(\text{OB}_{0:T-1}^{(1)} \in (\theta_-, \theta_+), \text{OB}_{0:T}^{(1)} \geq \theta_+ | H = 1).$$

We observe that $d_{T-1}^{(1)} = 0, d_T^{(1)} = 1$ imply $y_{T-1}^{(1)} \in (\theta_-, \theta_+)$ and $y_T^{(1)} \geq \theta_+$.

Hence

$$\begin{aligned} \theta_- &< \sum_{t=0}^{T-1} \log \frac{P(\xi_t^{(1)} | H = 1)}{P(\xi_t^{(1)} | H = -1)} < \theta_+ \\ &\sum_{t=0}^T \log \frac{P(\xi_t^{(1)} | H = 1)}{P(\xi_t^{(1)} | H = -1)} \geq \theta_+ \end{aligned}$$

Following [12], we have

$$\begin{aligned} \prod_{t=0}^T \frac{P(\xi_t^{(1)} | H = 1)}{P(\xi_t^{(1)} | H = -1)} &\geq e^{\theta_+} \\ \prod_{t=0}^T P(\xi_t^{(1)} | H = 1) &\geq e^{\theta_+} \prod_{t=0}^T P(\xi_t^{(1)} | H = -1) \end{aligned} \tag{4.3.1}$$

4.3. THRESHOLDS AND DECISION EVIDENCE

Hence $P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1)$ is the probability of making a chain of observations $\xi_{0:T}^{(1)}$ from f_+ whose log-likelihood probability surpasses θ_+ at time T , but remains in (θ_-, θ_+) prior to T . We call the collection of such “legal” observation chains $\mathcal{L}$ and note that $\mathcal{L} \subseteq \Xi^T$. We then have:

$$P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1) = \int_{\mathcal{L}} f_+(\xi_{0:T}^{(1)}) d\xi_{0:T}^{(1)}.$$

Every legal chain satisfies (4.3.1), hence we can integrate:

$$\int_{\mathcal{L}} f_+(\xi_{0:T}^{(1)}) d\xi_{0:T}^{(1)} \geq \int_{\mathcal{L}} e^{\theta_+} f_-(\xi_{0:T}^{(1)}) d\xi_{0:T}^{(1)} = e^{\theta_+} P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = -1).$$

Thus

$$P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1) \geq e^{\theta_+} P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = -1)$$

and so

$$\log \frac{P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = 1)}{P(d_{T-1}^{(1)} = 0, d_T^{(1)} = 1 | H = -1)} \geq \theta_+.$$

When the evidence distributions are close together $\log \frac{P(\xi_t^{(1)} | H=1)}{P(\xi_t^{(1)} | H=-1)}$ will be small and since $\sum_{t=0}^{T-1} \log \frac{P(\xi_t^{(1)} | H=1)}{P(\xi_t^{(1)} | H=-1)} < \theta_+$, we will have $\sum_{t=0}^T \log \frac{P(\xi_t^{(1)} | H=1)}{P(\xi_t^{(1)} | H=-1)} \approx \theta_+$. $\square$

The equivalent results hold when agent 1 chooses $H = -1$ and provides $\theta_- < 0$ evidence to agent 2. Thus, we have fully described the belief of the second agent when it receives a decision from its upstream neighbor. This computation shows that a decision from a neighboring agent gives a kick of at least the threshold size, as compared with a heuristic fraction of the threshold size used in [16]. Next, we look at the behavior of the evidence before agent 1 makes a decision: $\text{DEC}(d_t^{(1)} = 0)$.

4.3.2 Non-decision evidence and Symmetry

When $|\theta_-| \neq |\theta_+|$ we say that the thresholds are **asymmetric** and when $\theta_- = -\theta_+$ we say that the thresholds are **symmetric**. The farther the thresholds are away from zero, the more accurate we expect an agent's decision to be, but at the cost of response time [12, 51]. This is intuitive because we can think of $P(d_t^{(1)} = \pm 1 | H = \pm 1) = \alpha$ as the probability that agent 1 makes the correct decision and $P(d_t^{(1)} = \mp 1 | H = \pm 1) = 1 - \alpha$ as the probability that it makes an error. If we set $\theta_+ = \log \frac{\alpha}{1-\alpha}$, then the larger α is, the more accurate the agent will be, but it takes more observations to achieve that level of accuracy.

We call the evidence agent 2 has before agent 1 makes a decision the **non-decision evidence**, given by

$$\text{DEC}(d_t^{(1)} = 0) = \log \frac{P(d_t^{(1)} = 0 | H = 1)}{P(d_t^{(1)} = 0 | H = -1)}. \quad (4.3.2)$$

We will show that non-decisions are uninformative, i.e., the non-decision evidence is zero, given symmetric thresholds and a symmetry condition on the evidence distributions. Furthermore, this will be true for more general network structures, which we examine in the next chapters. Thus, we will assume these symmetries in the next chapters to simplify our examination of decision making dynamics in larger networks.

The quantity $P(d_t^{(1)} = 0 | H = 1)$ is the probability that the observations of

4.3. THRESHOLDS AND DECISION EVIDENCE

agent 1 have not caused it to make a decision. The absence of a decision is related to the **survival probability**:

$$S_{\pm}(t) = P(d_t^{(1)} = 0 | H = \pm 1) = P(\text{OB}_{0:s}^{(1)} \in (\theta_-, \theta_+), 0 \leq s \leq t | H = \pm 1).$$

When the thresholds are asymmetric we typically have $S_-(t) \neq S_+(t)$. Thus, when no decision is made this should give evidence for the evidence distribution associated with the larger threshold. We think of $d_t^{(1)}$ as providing partial information of the stochastic process $y_t^{(1)}$: It reveals whether or not the process $y_t^{(1)}$ has hit a threshold yet. This means the non-decision information is the log ratio of the survival probabilities:

$$\text{DEC}(d_t(1) = 0) = \log \frac{S_+(t)}{S_-(t)}.$$

We will do some explicit calculations for evidence accumulation with asymmetric thresholds in the next section.

Now let us assume that the thresholds are symmetric. Intuitively, let the evidence distributions $f_{\pm}$ be symmetric about the y -axis ($f_+(\xi) = f_-(-\xi)$) then the belief of agent 1 should remain bounded by the thresholds (survive) with the same probability given $H = 1$ or $H = -1$. To see this, as in the proof of 5, we let $\mathcal{L}_S$ be the collection of chains of observations $\xi_{0:t}^{(1)}$ that survive, i.e., do not lead to a threshold crossing. Then if $\xi_{0:t}^{(1)}$ survives, so does $-\xi_{0:t}^{(1)}$ since for $0 \leq s \leq t$:

$$\sum_{l=0}^s \log \frac{f_+(\xi_l^{(1)})}{f_-(\xi_l^{(1)})} = \sum_{l=0}^s \log \frac{f_-(-\xi_l^{(1)})}{f_+(-\xi_l^{(1)})} = - \sum_{l=0}^s \log \frac{f_+(-\xi_l^{(1)})}{f_-(-\xi_l^{(1)})}$$

4.3. THRESHOLDS AND DECISION EVIDENCE

So, if $\sum_{l=0}^s \log \frac{f_+(\xi_l^{(1)})}{f_-(\xi_l^{(1)})} \in (\theta_-, \theta_+)$, so is $\sum_{l=0}^s \log \frac{f_+(-\xi_l^{(1)})}{f_-(-\xi_l^{(1)})}$. Hence,

$$\begin{aligned} S_+(t) &= \int_{\mathcal{L}_S} f_+(\xi_{0:t}^{(1)}) d\xi_{0:t}^{(1)} \\ &= \int_{\mathcal{L}_S} f_-(-\xi_{0:t}^{(1)}) d\xi_{0:t}^{(1)} \\ &= \int_{\mathcal{L}_S} f_-(\xi_{0:t}^{(1)}) d\xi_{0:t}^{(1)} \\ &= S_-(t). \end{aligned}$$

We generalize this definition to any pair of evidence distributions that satisfy this property with the following.

Definition 4. The conditional measurement distributions $P(\xi|H = 1) = f_+(\xi)$ and $P(\xi|H = -1) = f_-(\xi)$ are **symmetric** if $\forall z \in (\theta_-, \theta_+)$

$$P(\log \frac{f_+(\xi)}{f_-(\xi)} = z | H = 1) = P(\log \frac{f_+(\xi)}{f_-(\xi)} = -z | H = -1)$$

whenever $-z \in (\theta_-, \theta_+)$.

For example, if we let $\mathcal{N}(\xi; \mu, \sigma^2)$ denote the probability density function of the normal distribution with mean μ and variance σ^2 , then when $f_+(\xi) = \mathcal{N}(\xi; \mu_+, \sigma^2)$ and $f_-(\xi) = \mathcal{N}(\xi; \mu_-, \sigma^2)$ with $\mu_+, \mu_- \in \mathbb{R}$ and $\sigma^2 \in \mathbb{R}^+$, we get distributions are symmetric.

Thus, we finish by noting that when the thresholds and decisions are symmetric we have

$$\text{DEC}(d_t^{(1)} = 0) = \log \frac{P(d_t^{(1)} = 0 | H = 1)}{P(d_t^{(1)} = 0 | H = -1)} = \log \frac{S_+(t)}{S_-(t)} = 0,$$

i.e., non-decisions are uninformative.

Unidirectional Coupling Evidence Accumulation

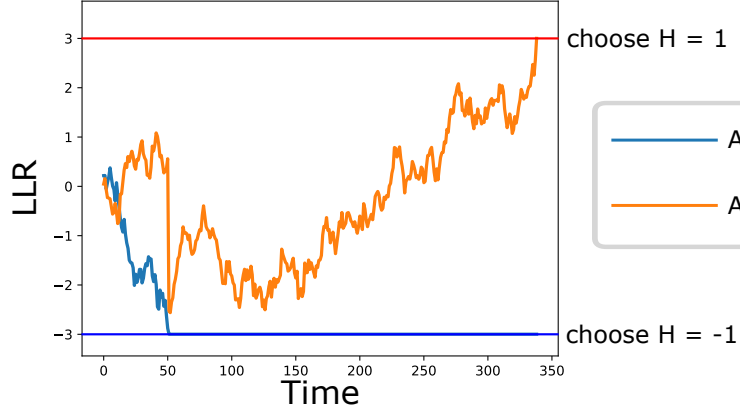


Fig. 4.3.1: An example of evidence accumulation with two agents, and normal evidence distributions with means ± 0.1 corresponding to $H = \pm 1$ and equal variance 1. The true state is $H = 1$. Agent 1 makes the wrong decision, but this does not mislead the second agent.

Summary Assume the decision thresholds and observation distributions are symmetric. Then before agent 1 makes a decision at time T , agent 2 only has evidence from its own observations:

$$y_t^{(2)} = \text{OB}_{0:t}^{(2)}, \quad t < T.$$

After the first agent makes a decision, the log-likelihood ratio for agent 2 will jump by θ_+ if agent 1 chooses $H = 1$ and θ_- if it chooses $H = -1$. Hence

$$y_t^{(2)} = \text{OB}_{0:t}^{(2)} + d_T^{(1)} \theta_+, \quad t \geq T.$$

An example of this process is given in Figure 4.3.1. The true state is $H = 1$. Even though agent 1 chooses $H = -1$ around $t = 50$ and this causes agent 2 to update its belief by θ_- , agent 2 still ends up accumulating enough evidence to make the correct decision.

4.4 Discrete Example

To gain some intuition we look at this process when the measurement distributions f_+ and f_- are discrete and then we do explicit calculations for an example where the decision thresholds are small and asymmetric.

4.4.1 Setup

We assume that each agent can make one of three observations, A and B , with

$$\begin{aligned} P(A|H = 1) &= p & P(A|H = -1) &= q \\ P(B|H = 1) &= q & P(B|H = -1) &= p \\ P(C|H = 1) &= s & P(C|H = -1) &= s \end{aligned} \tag{4.4.1}$$

where $p + q + s = 1$ and $p \geq q \geq 0, s \geq 0$. Thus observation A supplies stronger evidence for $H = 1$, observation B provides stronger evidence for $H = -1$, and observation C is uninformative. Equivalently, each informative measurement has only two values and has some probability of being incorrect. These equations define the confusion matrix [34]. Note that these discrete evidence distributions are symmetric, where f_+ is defined by the left equations and f_- the right equations in (4.4.1).

For simplicity we let $\theta = \log \frac{p}{q} > 0$, and define the positive threshold as $\theta_+ = m\theta$ for some $m \in \mathbb{N}$, and the negative threshold as $\theta_- = -n\theta$ for some $n \in \mathbb{N}$. These assumptions simplify the analysis, as the thresholds, as well as the beliefs of the two agents are restricted to lie on the lattice $\theta\mathbb{Z} = \{k\theta : k \in \mathbb{Z}\}$. The

4.4. DISCRETE EXAMPLE

informative private observations A and B change the belief of an agent by θ and $-\theta$, respectively. Thus an A measurement and a B measurement can cancel each other out, so it takes m cumulative measurements of A to choose $H = 1$, and n cumulative measurements of B to choose $H = -1$. Thus for an isolated agent, repeated measurements result in a biased random walk on the $\theta\mathbb{Z}$ lattice.

4.4.2 Non-Decision Evidence

As we discussed in Section 4.2, the belief of agent 2 at any point in time can be split into a part corresponding to private observations, and another part corresponding to social information obtained from agent 1. Thus, to determine the behavior of the unidirectional coupling for discrete distributions, we first compute the amount of information the second agent receives before agent 1 makes a decision. For $t \geq 0$ the change in belief of agent 2 is the non-decision evidence defined in (4.3.2).

We are looking at a biased random walk on a lattice where the boundaries are absorbing. Thus the terms in the numerator and denominator of the log-ratio in Eq. (4.3.2) are the probabilities that the belief of agent 1 has not reached either threshold by time i , given that the true states is $H = 1$ or $H = -1$, respectively. Thus the update of agent 2 is based on an inference about the unobserved state of agent 1's belief. This belief evolves as a biased random walk on a lattice with absorbing boundaries (See Figure 4.4.1). Therefore agent 2 can explicitly calculate the relative probability that a walk has not escaped, for example:

$$P(d_t^{(1)} = 0 | H = 1) = \sum_{l=-(n-1)}^{m-1} P(y_t^{(1)} = l | H = 1).$$

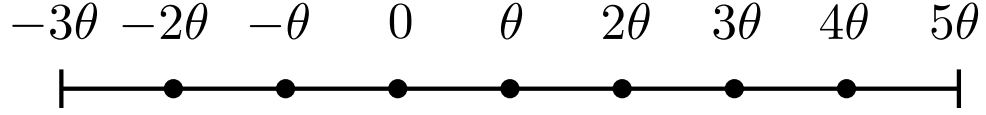


Fig. 4.4.1: An example lattice for a random walk where $\theta_- = -3\theta$ and $\theta_+ = 5\theta$ are the absorbing boundaries.

The random walk starts at 0 and each observation corresponds to a one step movement on the lattice. Thus $d_t^{(1)}$ must equal zero until a walk can hit the boundary, and thus $\text{DEC}(d_t^{(1)} = 0) = 0$ for $t < \min \{m, n\}$. Hence, not observing a decision during this initial time provides no information about agent 1's belief, since no decision could have occurred in either $H = 1$ or $H = -1$.

Survival Trajectories

From here on we assume $m \geq n$. Let $\mathbf{v}(t)$ be the vector of length $(m + n + 1) - 2$ which catalogues the probability that a walk starting at 0 ends up at each state without hitting an absorbing boundary by time t . We index $\mathbf{v}(t)$ using the lattice positions:

$\mathbf{v}_j(t) =$ probability a non-escaping walk is located at lattice position j at time t ,

with $t \geq 0$ and $-n < j < m$. Since our walks start at the origin we set the initial condition:

$$\mathbf{v}(0) = (0, \dots, 0, 1, 0, \dots, 0)^T$$

where the 1 is at the 0 index, i.e., $\mathbf{v}_0(0) = 1$.

4.4. DISCRETE EXAMPLE

Define

$$A_{p,q,s} = \begin{bmatrix} s & q & 0 & 0 \\ p & s & q & 0 \\ 0 & p & s & q \\ 0 & 0 & p & s \\ \vdots & & \ddots & \ddots & \ddots \end{bmatrix}.$$

Note that $A_{1,1,1}$ is the adjacency matrix for walks. Using $A_{p,q,s}$ we can define the update equation

$$\mathbf{v}(t) = A_{p,q,s} \mathbf{v}(t-1)$$

or

$$\mathbf{v}(t) = A_{p,q,s}^t \mathbf{v}(0).$$

Now we have a formula for the non-decision evidence:

$$\text{DEC}(d_t^{(1)} = 0) = \log \left(\frac{(A_{p,q,s}^t \mathbf{v}(0)) \cdot \mathbf{1}}{(A_{q,p,s}^t \mathbf{v}(0)) \cdot \mathbf{1}} \right).$$

This is simple to compute numerically for relatively small decision boundary sizes, but unfortunately we cannot get an explicit formula for $(A_{p,q,s}^t \mathbf{v}(0)) \cdot \mathbf{1}$, the survival probability of the biased random walk with absorbing boundaries, for general $m, n \in \mathbb{N}$. Next we look at a couple of cases where we can do a direct computation and then produce simulations for more general examples.

Symmetric Absorbing Boundaries

When the decision thresholds are symmetric, that is, when $m = n$, then because the distributions for $P(\xi|H = 1)$ and $P(\xi|H = -1)$ are symmetric we know that

4.4. DISCRETE EXAMPLE

non-decision evidence is zero in this case, as shown in Section 4.3.2. This makes sense because $(A_{p,q,s}^i \mathbf{v}(0)) \cdot \mathbf{1} = (A_{q,p,s}^i \mathbf{v}(0)) \cdot \mathbf{1}$ since $A_{p,q,s}$ is the transpose of $A_{q,p,s}$. As a result, $\text{DEC}(d_i^{(1)} = 0)$. When the bounds are not symmetric, $m \neq n$, the survival probabilities will eventually differ because probability should leak out of the closer boundary faster.

Analytic Example: Four State Biased Random Walk

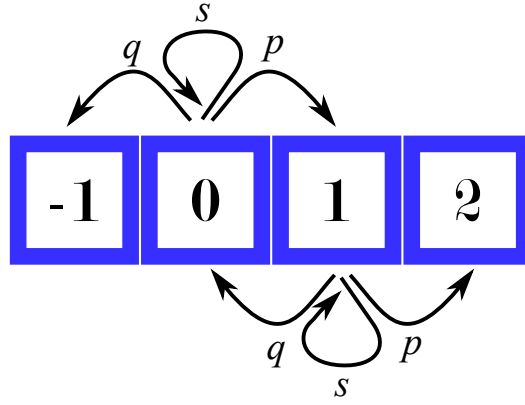


Fig. 4.4.2: A simple example lattice with two transient states, two absorbing states, and transition probabilities indicated by arrows.

We compute explicit log-likelihood ratios for the simplest nontrivial example with non-symmetric boundaries: a biased random walk on 4 states. To simplify things even more, we let $\theta = 1$, so that we get the lattice in Figure 4.4.2,

Every walk starts at location 0, and locations -1 and 2 are the absorbing boundaries. Stationary steps correspond to uninformative observations.

We will use standard first-step analysis [53]. Let $y_n^{(1)}$ be the location of agent 1's belief at time n . For simplicity let $P_+(y_n^{(1)} = l) = P(y_n^{(1)} = l | H = 1)$ denote

4.4. DISCRETE EXAMPLE

the probability that agent 1's belief is at location $l \in \{-1, 0, 1, 2\}$ given that the true state is $H = 1$. Then the probability of the belief reaching different locations is given by:

$$\begin{aligned} P_+(y_n^{(1)} = -1) &= P_+(y_{n-1}^{(1)} = -1) + qP_+(y_{n-1}^{(1)} = 0) \\ P_+(y_n^{(1)} = 0) &= sP_+(y_{n-1}^{(1)} = 0) + qP_+(y_{n-1}^{(1)} = 1) \\ P_+(y_n^{(1)} = 1) &= pP_+(y_{n-1}^{(1)} = 0) + sP_+(y_{n-1}^{(1)} = 1) \\ P_+(y_n^{(1)} = 2) &= pP_+(y_{n-1}^{(1)} = 1) + P_+(y_{n-1}^{(1)} = 2) \\ P_+(y_0^{(1)} = 0) &= 1 \end{aligned}$$

where $p + q + s = 1$ as before with $p > q > 0$ and $s \geq 0$, so that the belief is biased towards $H = 1$.

So we have:

$$\mathbf{v}(n) = \begin{pmatrix} P_+(y_n^{(1)} = 0) \\ P_+(y_n^{(1)} = 1) \end{pmatrix} = A_{p,q,s}^n \mathbf{v}(0) = \begin{pmatrix} s & q \\ p & s \end{pmatrix}^n \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

We can use eigenvector analysis to compute

$$\text{DEC}(d_n^{(1)} = 0) = \log \frac{(A_{p,q,s}^n \mathbf{v}(0)) \cdot \mathbf{1}}{(A_{q,p,s}^n \mathbf{v}(0)) \cdot \mathbf{1}}$$

explicitly for this case. $A_{p,q,s}^n$ has the eigenvalues $\lambda_1 := s - \sqrt{pq}$, $\lambda_2 := s + \sqrt{pq}$.

Then we can write the corresponding eigenvectors:

$$\mathbf{e}_1 = \begin{pmatrix} \frac{-q}{\sqrt{pq}} \\ 1 \end{pmatrix}, \mathbf{e}_2 = \begin{pmatrix} \frac{q}{\sqrt{pq}} \\ 1 \end{pmatrix},$$

so that

$$\mathbf{v}(0) = \frac{-\sqrt{pq}}{2q} \mathbf{e}_1 + \frac{\sqrt{pq}}{2q} \mathbf{e}_2 = \begin{pmatrix} \frac{1}{2} \\ \frac{-\sqrt{pq}}{2q} \end{pmatrix} + \begin{pmatrix} \frac{1}{2} \\ \frac{\sqrt{pq}}{2q} \end{pmatrix}.$$

4.4. DISCRETE EXAMPLE

Which gives us a simple expression for $\mathbf{v}(n)$:

$$\mathbf{v}(n) = \lambda_1^n \begin{pmatrix} \frac{1}{2} \\ -\frac{\sqrt{pq}}{2q} \end{pmatrix} + \lambda_2^n \begin{pmatrix} \frac{1}{2} \\ \frac{\sqrt{pq}}{2q} \end{pmatrix}.$$

We get do the same for $A_{q,p,s}^n$ by swapping p and q . Hence:

$$\text{DEC}(d_n^{(1)} = 0) = \log \frac{\frac{1}{2}(\lambda_1^n + \lambda_2^n) + \frac{\sqrt{pq}}{2q}(\lambda_2^n - \lambda_1^n)}{\frac{1}{2}(\lambda_1^n + \lambda_2^n) + \frac{\sqrt{pq}}{2p}(\lambda_2^n - \lambda_1^n)} = \log \frac{1 + \frac{2\lambda_1^n}{\lambda_2^n - \lambda_1^n} + \frac{\sqrt{pq}}{q}}{1 + \frac{2\lambda_1^n}{\lambda_2^n - \lambda_1^n} + \frac{\sqrt{pq}}{p}}.$$

In order to see what the asymptotic non-decision information is, first note that

$$\frac{2\lambda_1^n}{\lambda_2^n - \lambda_1^n} = \frac{2}{\left(\frac{\lambda_2}{\lambda_1}\right)^n - 1}.$$

For our choice of s, p, q we have $\lambda_2/\lambda_1 \approx -7.98$ so we get fairly fast convergence

$$\lim_{n \rightarrow \infty} \frac{2}{\left(\frac{\lambda_2}{\lambda_1}\right)^n - 1} \rightarrow 0$$

which gives

$$\lim_{n \rightarrow \infty} \text{DEC}(d_n^{(1)} = 0) = \log \frac{1 + \sqrt{\frac{p}{q}}}{1 + \sqrt{\frac{q}{p}}}$$

Then $\sqrt{\frac{p}{q}} = \sqrt{e}$ so

$$\lim_{n \rightarrow \infty} \text{DEC}(d_n^{(1)} = 0) = \log \frac{1 + \sqrt{e}}{1 + \frac{1}{\sqrt{e}}} = \frac{1}{2}.$$

because

$$\frac{1}{\sqrt{e}} \frac{1 + \sqrt{e}}{1 + \frac{1}{\sqrt{e}}} = \frac{1 + \sqrt{e}}{\sqrt{e} + 1} = 1$$

implies

$$\frac{1 + \sqrt{e}}{1 + \frac{1}{\sqrt{e}}} = \sqrt{e}.$$

4.4. DISCRETE EXAMPLE

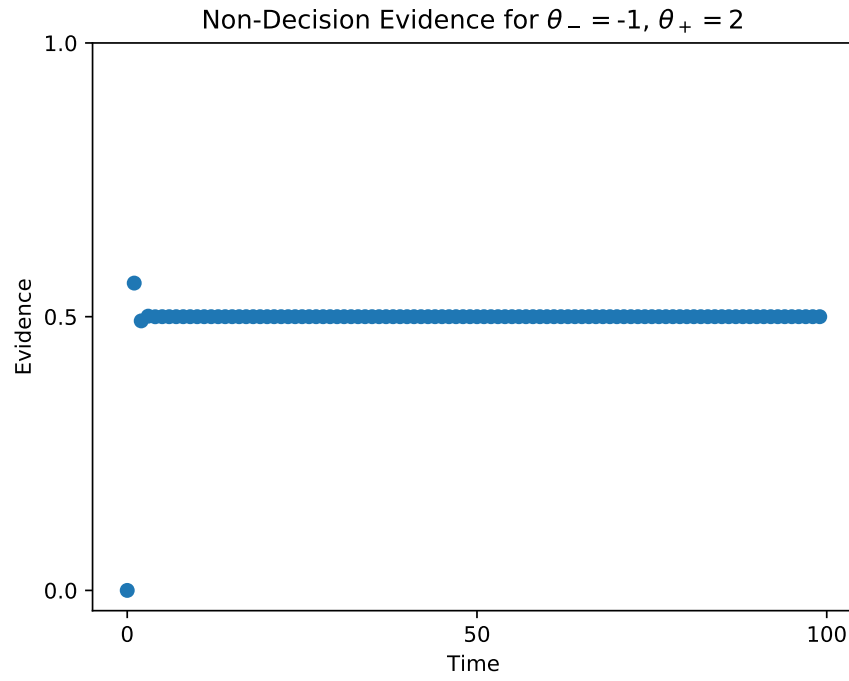


Fig. 4.4.3: The value of the non-decision evidence that computes for $\theta_- = -1, \theta_+ = 2, p = \frac{\epsilon}{5}, q = \frac{1}{5}, s = 1 - p - q$. We see that the evidence quickly converges to $\frac{1}{2}$.

4.4. DISCRETE EXAMPLE

We thus get that the non-decision information fairly quickly converges to $\frac{1}{2}$, a fact we confirm with simulations in Figure 4.4.3

We have a recursive formula for the belief distribution of agent 1. We also have a closed form for the belief distribution in a specific case. The closed form allows us to describe how the non-decision evidence from agent 1 behaves and we confirm that behavior more generally with simulations.

In particular, this example shows us that we might expect the evidence from a non-decision to saturate at a significantly large value that is bounded by the upper threshold, but not enough to cause the agent to decide without additional evidence. We show this holds using simulations.

4.4.3 Simulations

We next investigate the evolution of beliefs when coupling is unidirectional and the measurement distribution discrete. Using simulations we will show how non-decision evidence defined in (4.3.2) depends on threshold size. We will further investigate the size of the non-decision evidence by varying the evidence obtained from each observation.

Non-decision Evidence Choosing probabilities that allow us to have integer-valued thresholds, we plot the non-decision evidence for four pairs of thresholds in Figure 4.4.4. We see that the evidence saturates for all pairs to a value less than the threshold. This is related to a general fact about how much evidence social

4.4. DISCRETE EXAMPLE

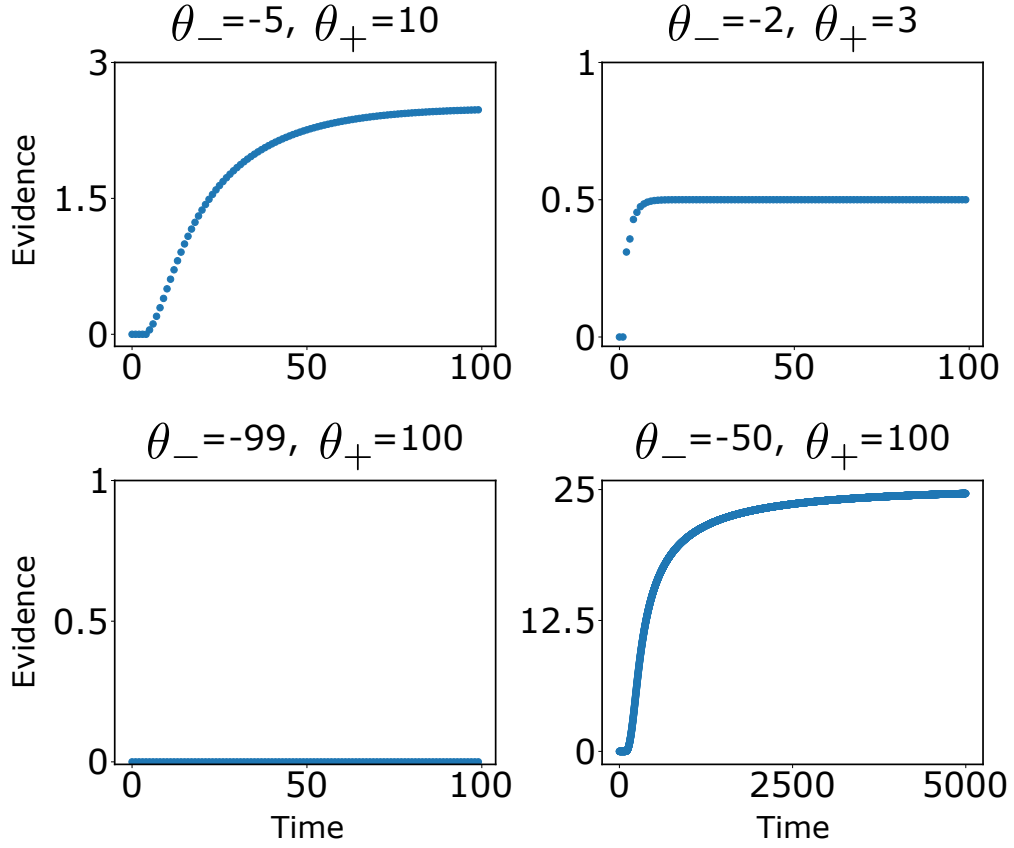


Fig. 4.4.4: Non-decision evidence saturates for different ratios of thresholds. We let $p = \frac{e}{5}, q = \frac{1}{5}, s = 1 - q - p$ so that $\log \frac{p}{q} = 1$ and $p + q + s = 1$.

information can provide shown in a claim in Chapter 6. Furthermore when the thresholds are close in size, we see that less evidence is provided than the evidence from an informative private observation. But, when the thresholds are different in size, the non-decision evidence becomes substantial; the equivalent of almost 25 A observation when $\theta_- = -50$ and $\theta_+ = 100$, which is shown in the bottom right panel of Figure 4.4.4.

Dependence on Bias For large thresholds, it will take several steps for a decision to be made. If $\theta_- = -n\theta$ and $\theta_+ = m\theta$, with $m < n$ then it will take at least n time steps for a decision to be possible, in which case it must have resulted from n consecutive B observations. As a consequence, the absence of a decision during these first $n - 1$ steps is not informative about the belief of agent 1. We can compute exactly how much evidence agent 2 obtains if agent 1 does not make a decision after observation n :

$$\text{DEC}(d_n^{(1)} = 0) = \log \frac{1 - p^n}{1 - q^n}.$$

The amount of evidence when the non-decision information is first informative is shown in Figure 4.4.5 for different values of p and n . This tells us that the initial evidence is smaller the farther away the boundaries are. Furthermore, it shows that the initial evidence for a non-decision is always smaller than the evidence provided by a decision. Thus the non-decision information will have a small impact early on. However, since it is positive, if agent 2 does not choose $H = -1$ before receiving $d_n^{(1)}$, then the non-decision evidence is enough to force an extra B observation before choosing $H = -1$ because $-n\theta + \text{DEC}(d_n^{(1)} = 0) > -n\theta$.

4.5 Continuum Limit

We finish our investigation of the simple network in Figure 4.2.1 by characterizing the continuum limit of the evidence accumulation process. We obtain this limit by letting the time between observations, as well as the change in the belief of each agent due to private observations go to 0 [12]. We start with the continuum limit

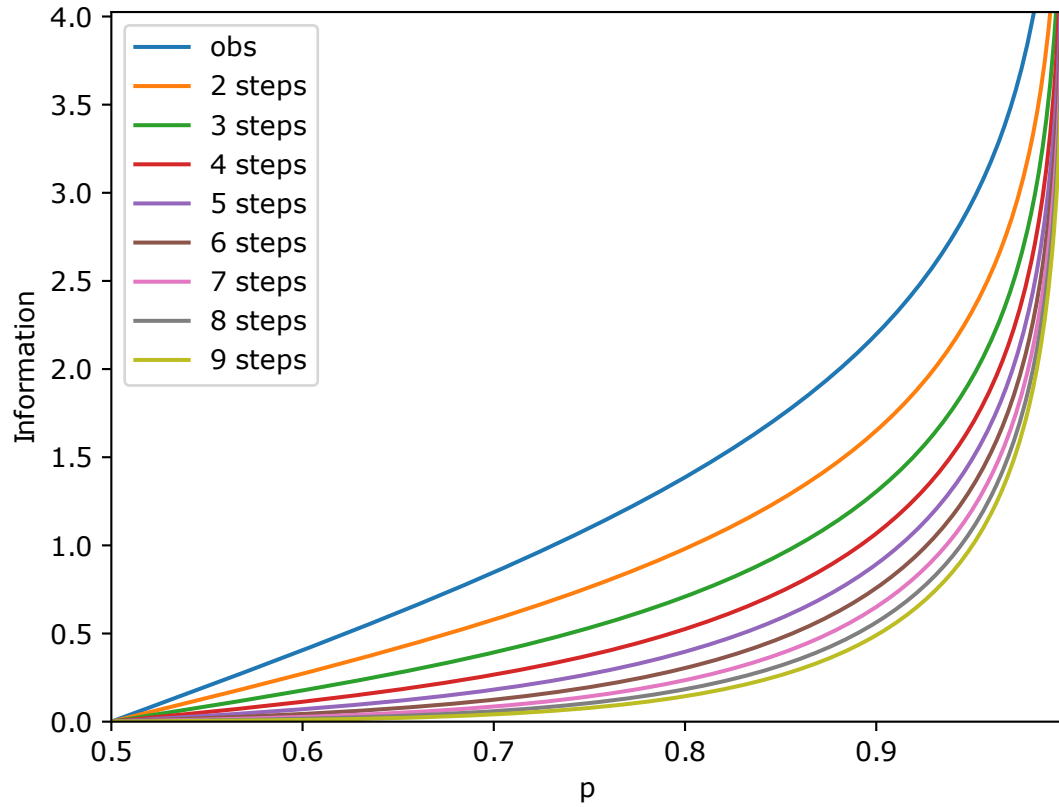


Fig. 4.4.5: The amount of evidence obtained by agent 2 after observing a non-decision at time step n where n is the minimal number of steps needed to reach the closer, negative threshold. The parameter p defining the observational distribution is varied in the interval $p \in [0.5, 1)$ to obtain the different curves. For comparison, the amount of evidence obtained from a single A observation is given by the blue curve.

4.5. CONTINUUM LIMIT

of agent 1. As this agent only makes observations this is equivalent to previous derivations that use the Functional Central Limit Theorem to describe the evolution of the belief as a drift–diffusion process [54, 12]. We next turn to agent 2 by first assuming it does not make private observation and only receives the decision states of agent 1. We conclude by describing the evolution of the belief of an agent that combines both social and private information to update its belief.

For simplicity we give the results for the concrete evidence distributions N_+ and N_- , defined by

$$N_{\pm}(\xi) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(\frac{-1}{2\sigma^2}(\xi \mp \mu)^2\right) \quad (4.5.1)$$

with $\mu, \sigma^2 > 0$. So, the likelihood functions are determined by the normal distributions, $N_-(\xi)$ and $N_+(-\xi)$ so they are symmetric about the y -axis and thus are symmetric as in Definition 4. The results hold for more general distributions, but require more theoretical detail without much distinction in the resulting behavior.

Observations Only

As shown in [54, 12], the first agent's belief evolves approximately as a realization of a process described by the SDE:

$$dy_t^{(1)} = gdt + \rho dW \quad (4.5.2)$$

where W_t is a standard Wiener process and g and ρ^2 are defined as:

$$g := E_{\xi} \left[\ln \frac{f_+(\xi)}{f_-(\xi)} \middle| H_T \right], \quad \rho^2 := \text{Var}_{\xi} \left[\ln \frac{f_+(\xi)}{f_-(\xi)} \middle| H_T \right].$$

4.5. CONTINUUM LIMIT

where H_T is the true state.

When the observations follow normal distributions, then

$$\ln \frac{N_+(\xi)}{N_-(\xi)} = \frac{2\mu}{\sigma^2} \xi,$$

and so

$$g = H_T \frac{2\mu^2}{\sigma^2}, \quad \rho^2 = \frac{4\mu^2}{\sigma^2}.$$

Decision States Only

To simplify the computation of $y_t^{(2)}$, we first compute z_t , the amount of evidence agent 2 receives from agent 1 alone, in the absence of its own observations. Assume the observations are made at discrete times $\{t_n\}_{n \in \mathbb{N}_0} = \{t_0, t_1, \dots\}$ and that no decision has been made by time t_n :

$$z_n = \text{DEC}(d_n^{(1)} = 0) = \log \frac{P(d_n^{(1)} = 0 | H = 1)}{P(d_n^{(1)} = 0 | H = -1)}.$$

Then define the time increment between any two consecutive decisions:

$$\Delta t = t_n - t_{n-1}.$$

To define the continuum limit process we assume Δt is small and introduce the difference:

$$\Delta z_n = z_n - z_{n-1} = \log \frac{P(d_n^{(1)} = 0 | H = 1)}{P(d_n^{(1)} = 0 | H = -1)} - \log \frac{P(d_{n-1}^{(1)} = 0 | H = 1)}{P(d_{n-1}^{(1)} = 0 | H = -1)}.$$

Given Δt , we can compute a differential equation for Δz_n . Note that since these increments are defined in terms of *probabilities* that an agent has not made

4.5. CONTINUUM LIMIT

a decision up to a certain time, they are not random. The probabilities in these ratios are survival probabilities:

$$P(d_n^{(1)} = 0 | H = \pm 1) = S_{\pm, t_n} = \int_{\theta_-}^{\theta_+} P(y_n^{(1)} = x | H = \pm 1) dx.$$

We assume, but do not rigorously show, that as $\Delta t \rightarrow 0$ the discrete step survival probabilities approximate the continuous survival probabilities $S_{\pm}(t)$, which are the probabilities that the solution to (4.5.2) has not hit a threshold by continuous time $t \in \mathbb{R}^{\geq 0}$ given $H = \pm$. Therefore,

$$\lim_{\Delta t \rightarrow 0} S_{\pm, t_n} \rightarrow S_{\pm}(t).$$

Under this assumption for $t = t_{n-1}$:

$$\begin{aligned} \Delta z_n &= \log \frac{S_+(t_n)}{S_-(t_n)} - \log \frac{S_+(t_{n-1})}{S_-(t_{n-1})} \\ &= \log \frac{S_+(t + \Delta t)}{S_-(t + \Delta t)} - \log \frac{S_+(t)}{S_-(t)} \\ &= [\log S_+(t + \Delta t) - \log S_-(t + \Delta t)] - [\log S_+(t) - \log S_-(t)]. \end{aligned}$$

After dividing both sides by Δt , and taking the limit $\Delta t \rightarrow 0$ we get

$$z'(t) = \frac{d}{dt} \ln \frac{S_+(t)}{S_-(t)} = \frac{S'_+(t)}{S_+(t)} - \frac{S'_-(t)}{S_-(t)}$$

or

$$dz = \left(\frac{S'_+(t)}{S_+(t)} - \frac{S'_-(t)}{S_-(t)} \right) dt.$$

Thus in the case agent 2 only observes the decisions of agent 1, but makes no private observations, its belief evolves deterministically according to

$$z(t) = \ln \frac{S_+(t)}{S_-(t)},$$

until the time of a decision by agent 1.

Assuming that agent 1 has made a decision at time T : $d_T^{(1)} \neq 0$, we have

Claim 1.

$$dz_t = \left((1 - H_T(t)) \left(\frac{S'_+(t)}{S_+(t)} - \frac{S'_-(t)}{S_-(t)} \right) + \delta_T \theta_{d_T^{(1)}} \right) dt$$

Observations and Decision States

To finish, we combine the previous two process to fully describe the behavior of agent 2 in the continuous case. Assume the first agent has not reached a decision, and the second agent has made the observations $\xi_{0:n}^{(2)}$. Then the log-likelihood ratio for agent 2 at step n is:

$$y_n^{(2)} = \sum_{j=0}^n \ln \frac{P(\xi_j^{(2)} | H = 1)}{P(\xi_j^{(2)} | H = -1)} + \ln \frac{P(d_n^{(1)} = 0 | H = 1)}{P(d_n^{(1)} = 0 | H = -1)}.$$

Then, by independence of the observations, we can combine the previous equations to get the full SDE for agent 2.

Claim 2.

$$dy_t^{(2)} = \left(g + (1 - H_T(t)) \left(\frac{S'_+(t)}{S_+(t)} - \frac{S'_-(t)}{S_-(t)} \right) + \delta_T \theta_{d_T^{(1)}} \right) dt + \rho dW$$

Fokker-Planck Equation and Simulations In order to understand the behavior of the additional drift term, $\log \frac{S_+(t)}{S_-(t)}$, we note that if

$$dy_t^{(1)} = gdt + \rho dW$$

4.5. CONTINUUM LIMIT

and the process terminates when $y_t^{(1)} = \theta_-$ or θ_+ , then we can define

$$p_{\pm}(x, t) = p(y_t^{(1)} = x | H = \pm 1),$$

so that the probabilities $p_{\pm}(x, t)$ satisfy

$$\frac{\partial p(x, t)}{\partial t} = \mp a \frac{\partial p_{\pm}(x, t)}{\partial x} + \frac{b}{2} \frac{\partial^2 p_{\pm}(x, t)}{\partial x^2},$$

for some drift rate a and diffusion rate b , with absorbing boundaries

$$p_{\pm}(\theta_-, t) = p_{\pm}(\theta_+, t) = 0,$$

and initial condition

$$p_{\pm}(x, 0) = \delta(x).$$

It is possible to use a Fourier series to solve for $p_{\pm}(x, t)$ explicitly, but the coefficients are difficult to solve for and the analytic expression does not help our analysis much further. Instead we note that

$$S_{\pm}(t) = \int_{\theta_-}^{\theta_+} p_{\pm}(x, t) dx.$$

Hence, we numerically compute the solutions $p_{\pm}(x, t)$ and plot them in Fig. 4.5.1. We see that the probability mass between the absorbing boundaries stays high longer when the belief diffuses towards the farther boundary (higher threshold). Then to compute the survival probabilities we sum up our simulation of $p_{\pm}(x, t)$ at each time step and plot the results, including the resulting non-decision evidence, in Figure 4.5.2.

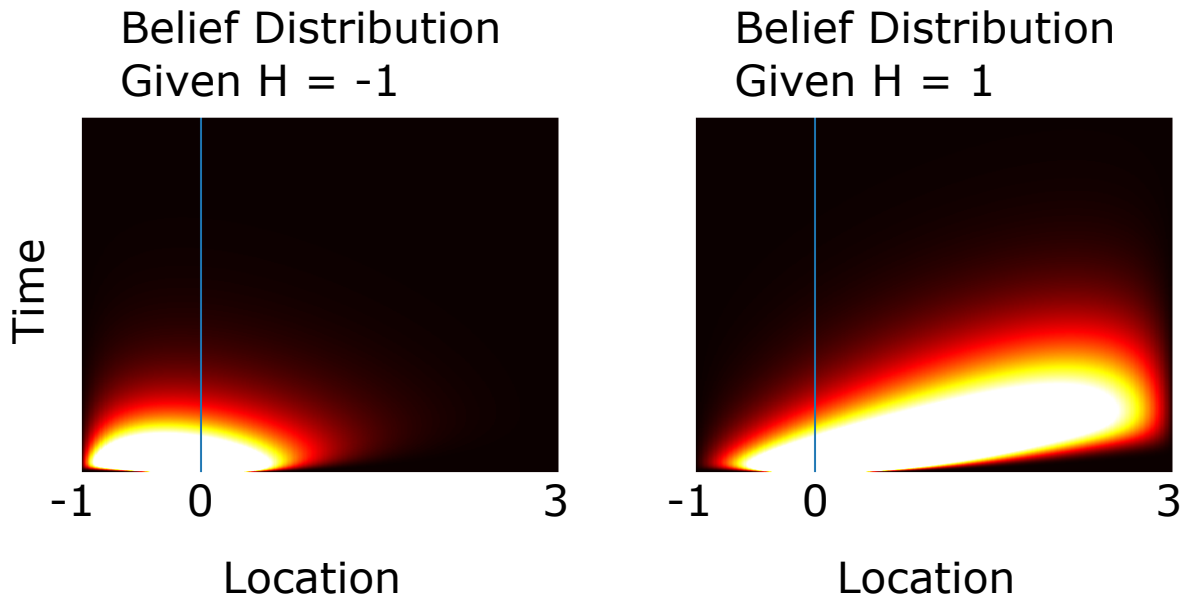


Fig. 4.5.1: A simulation of the conditional belief distributions for $H = -1$ (left) and $H = 1$ (right) using Crank-Nicolson discretization. Each plot shows how the belief is distributed and evolves in time where brightness corresponds to probability. Here we let the boundaries be $\theta_- = -1$ and $\theta_+ = 3$. We used a drift rate $a = 1$, diffusion rate $b = 1$, space step size $dx = 0.01$, time step size $dt = 0.001$, and total time $T = 10$.

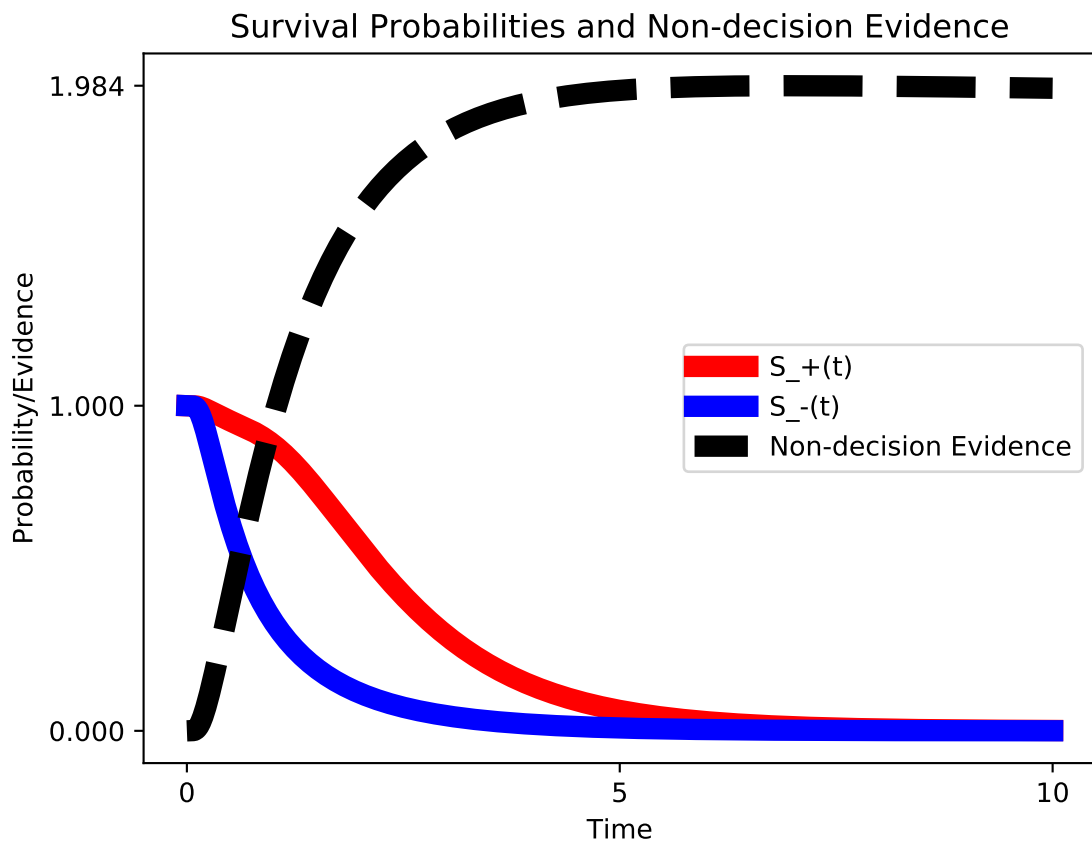


Fig. 4.5.2: Using the simulations from Figure 4.5.1, we compute the survival probabilities and resulting non-decision evidence and plot them on the same graph.

4.6 Conclusion

We derived the behavior for the simplest evidence accumulation network. This interaction will give us insight when things become more complex as we allow for more general networks which can have recurrent connections, but we will see that many of the concepts carry over. We showed that for symmetric pairs of evidence distributions and symmetric thresholds, the social information the downstream agent receives before agent 1 makes a decision is uninformative. If the thresholds are asymmetric or the evidence distributions are not symmetric, then we no longer have this property. The non-decision evidence the downstream agent receives can be substantial and will depend on the distributions and how long agent 1 has gone without making a decision. We gained intuition for the process by looking at discrete distributions. When we took a continuum limit and numerically computed the survival probabilities for our process, we saw that the same properties hold for more general distributions.

Chapter 5

Bidirectional Coupling

When people look to one another to try and make decisions, they are not merely observing each other's actions. When you observe another person's gaze, you also realize that the other person is looking at you. Moreover, the other person is acknowledging your gaze, and you realize this too. Thus we know that the other person knows that we are observing them, and we know that the other person knows that we know this. In theory, this recursive process does not end, and it is at the heart of the mathematical notion of "common knowledge." A similar recursive process also makes the analysis of a bidirectionally coupled pair of agents more difficult to study than the unidirectionally coupled pair of the previous chapter.

In this chapter we explore the dynamics of evidence accumulation and decision making for two agents that observe each other's decisions. We will show that

5.1. SETUP AND SYMMETRIC BOUNDARIES

this process is simple when the agents have symmetric boundaries and symmetric evidence distributions, as was the case of one-way communication. However, with asymmetries the absence of a decision is informative about an agent's belief. A rational observing agent will use this information to update its belief, leading to a recursive process which is difficult to describe explicitly.

In the first section we will describe the setup, introducing new notation because the process is more complicated. We will show that the process is equivalent to the unidirectional coupling case when boundaries and evidence distributions are symmetric. We then describe the case of asymmetric boundaries and give update equations for discrete time. We will then look at an example with discrete evidence distributions, and discuss simulations of the process that illustrate how the behavior differs from the unidirectional case.

5.1 Setup and Symmetric Boundaries

Figure 5.1.1, shows the graph corresponding to two agents who make independent, private observations, and share their decision states after each observation.

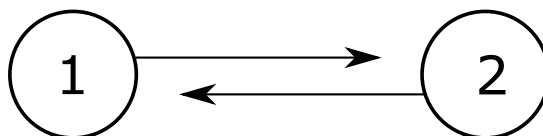


Fig. 5.1.1: Two bidirectionally coupled agents.

5.1. SETUP AND SYMMETRIC BOUNDARIES

We will assume time, t , is discrete and an observation is made during each time step, t_i . Moreover, we divide each time step into multiple parts. At the beginning of a time step, t_i , agents make a private observation, and update their belief. They then observe whether the neighboring agent has made a decision. As we will show, this leads to a process that can take a number of substeps. We assume that the agents continue to observe each other's decision state until there are no further changes in their beliefs. Only at this point do they make another private observation (this corresponds to time step t_{i+1}).

As in Section 4.2, an observation is a sample from the conditional evidence distribution f_+ or f_- . Thus agent $j = 1, 2$ makes the observation $\xi_t^{(i)}$, at time t_i , and updates its belief *from the end of the previous time step*, $y_{t-1}^{(i)}$. Therefore,

$$y_{t,0}^{(i)} = y_{t-1}^{(i)} + \log \left(\frac{P(\xi_t^{(i)} | H = 1)}{P(\xi_t^{(i)} | H = -1)} \right). \quad (5.1.1)$$

The second subscript in $y_{t,n}^{(i)}$ indexes the substeps between two observations. We will define $y_{t-1}^{(i)}$ iteratively using (5.1.3) below. As we will see, during the process, observers iteratively update their belief based on the observations of each other's decision state. At the beginning of this process, $n = 0$, each agent integrates evidence from a new observation (private information). For $n \geq 1$, each agent integrates the social information communicated by its neighbor.

The process proceeds as follows: For $n \geq 0$, after both agents compute $y_{t,n}^{(i)}$,

5.1. SETUP AND SYMMETRIC BOUNDARIES

they communicate their decision state to their counterpart:

$$d_{t,n}^{(i)} = \begin{cases} -1, & y_{t,n}^{(i)} \leq \theta^- \\ 0, & y_{t,n}^{(i)} \in (-\theta^-, \theta^+) \\ 1, & y_{t,n}^{(i)} \geq \theta^+. \end{cases}$$

Once an agent makes a decision, it cannot change it. So if $d_{t,n}^{(i)} \neq 0$ then $d_{s,m}^{(i)} = d_{t,n}^{(i)}$ for all $s \geq t, m \geq n$. At substep $n + 1$, between two observations agent i incorporates the observation of the decision $d_{t,n}^{(\neg i)}$ from their counterpart, at the previous substep. We first write:

$$y_{t,n+1}^{(i)} = \log \left(\frac{P(d_{t,n}^{(\neg i)}, h_{t-1}^{(i)}, d_{t,0:n-1}^{(\neg i)}, d_{t,0:n-1}^{(i)}, \xi_{0:t}^{(i)} | H = 1)}{P(d_{t,n}^{(\neg i)}, h_{t-1}^{(i)}, d_{t,0:n-1}^{(\neg i)}, d_{t,0:n-1}^{(i)}, \xi_{0:t}^{(i)} | H = -1)} \right).$$

Where we define the index $\neg i = 2, 1$ when $i = 1, 2$, respectively, $d_{t,0:n-1}^{(j)} = \{d_{t,0}^{(j)}, \dots, d_{t,n-1}^{(j)}\}$ are the decision states communicated by agent j up to substep $n - 1$, $d_{0:t}^{(i)}$ are the decision states at the end of each time step, and $h_{0:t-1}^{(i)}$ is the entire record of decision states of agent i knows up to time step $t - 1$:

$$h_{t-1}^{(i)} = \{d_{s,i}^{(1)}, d_{s,i}^{(2)}\}_{0 \leq s \leq t-1, i \in \mathbb{N}_0}.$$

Then we are able to split the belief into evidence from observations and evidence from decision states, but the latter must be conditioned on the observations:

$$y_{t,n+1}^{(i)} = \text{OB}_{0:t}^{(i)} + \log \left(\frac{P(d_{t,n}^{(\neg i)}, h_{t-1}^{(i)}, d_{t,0:n-1}^{(\neg i)}, d_{t,0:n-1}^{(i)} | \xi_{0:t}^{(i)}, H = 1)}{P(d_{t,n}^{(\neg i)}, h_{t-1}^{(i)}, d_{t,0:n-1}^{(\neg i)}, d_{t,0:n-1}^{(i)} | \xi_{0:t}^{(i)}, H = -1)} \right). \quad (5.1.2)$$

We note a key difference here from the unidirectional coupling case: decision states are no longer necessarily independent from observations. In the asymmetric case some information about agent i 's private observations is communicated

5.1. SETUP AND SYMMETRIC BOUNDARIES

through its decisions, and in general

$$P(d_{t,n}^{(-i)}, h_{t-1}^{(i)}, d_{t,0:n-1}^{(-i)}, d_{t,0:n-1}^{(i)} | \xi_{0:t}^{(i)}, H = \pm) \neq P(d_{t,n}^{(-i)}, h_{t-1}^{(i)}, d_{t,0:n-1}^{(-i)}, d_{t,0:n-1}^{(i)} | H = \pm).$$

Information that through the absence of the decision of an agent is then used by the other agent in the computation of its own decision states. This dependence makes it necessary to condition on observations in the numerator and denominator of the second term in Eq. (5.1.2).

If neither agent has made a decision at substep n of the observational step t , then

$$y_{t,n+1}^{(i)} = \text{OB}_{0:t}^{(i)} + \log \left(\frac{P(d_{t,n}^{(-i)} = 0, d_{t,n-1}^{(i)} = 0 | \xi_{0:t}^{(i)}, H = 1)}{P(d_{t,n}^{(-i)} = 0, d_{t,n-1}^{(i)} = 0 | \xi_{0:t}^{(i)}, H = -1)} \right)$$

because we do not allow decisions to change, and hence $d_{t,n}^{(-i)} = 0$ implies $d_{s,i}^{(-i)} = 0$ for $0 \leq s \leq t, 0 \leq i \leq n-1$.

To move from time step $t-1$ to t we let

$$y_t^{(i)} = \lim_{n \rightarrow \infty} y_{t,n}^{(i)} \tag{5.1.3}$$

be the belief of agent i at after communicating its decision states back and forth with its neighbor. We say that two agents are **equilibrating** their beliefs during this exchange of decision states, and we will refer to the repeated exchange of decision states as **equilibration** steps. We next show why this process simplifies when the decision thresholds and evidence distributions are symmetric, and in the next section we show that equilibration process converges even when the thresholds are asymmetric.

5.2. EQUILIBRATION PROCESS

Proposition 6. *When the distributions f_+ and f_- are symmetric and the agents have the same symmetric thresholds, then if agent j decides before agent i at time T , then*

$$y_t^{(i)} = \begin{cases} \text{OB}_{0:t}^{(i)} & , t < T \\ \text{OB}_{0:t}^{(i)} + d_T^{(j)}\theta & , t \geq T. \end{cases}$$

Thus, decision states are uninformative and both agent integration private observation evidence until one of the makes a decision. Once agent j chooses $H = \pm 1$, agent i 's belief changes by $\pm\theta$.

Proof. The argument is similar to that in Section 4.3.2. If the two agents have not made a decision then this does not provide any evidence for either choice $H = \pm$:

$$\text{DEC}(d_{t,n}^{(i)} = 0) = \log \frac{P(d_{t,n}^{(i)} = 0 | H = 1)}{P(d_{t,n}^{(i)} = 0 | H = -1)} = 0.$$

Thus non-decision evidence is zero, so when an agent communicates $d_{t,n}^{(i)} = 0$ it is uninformative and $d_{t,n+1}^{(-i)} = d_{t,n}$, and the equilibration process terminates at the first step. When an agent decides it can provide no further information to its counterpart. Therefore, agent i updates its belief by $\pm\theta$ as in the unidirectional case. $\square$

5.2 Equilibration Process

We show that the equilibration process converges so that at the end of each time step, prior to the next observation, both agents settle on the amount of evidence they gain from their private and social information.

5.2. EQUILIBRATION PROCESS

Proposition 7. *At the end of each time step $t \geq 0$, $\lim_{n \rightarrow \infty} y_{t,n}^{(i)}$ exists, and thus the beliefs of both agents equilibrate.*

Proof. Let $Y_{t,n} \subseteq (\theta_-, \theta_+)$ be the range of values that $y_{t,0}^{(i)}$ could have been equal to without causing the agent to make a decision by time t and equilibration step n . Note that each observation of a non-decision can only shrink this range of beliefs of the other agent. Since this is a finite interval, the process must converge. $\square$

Alternatively, we conjecture that following [49, 41], we can obtain convergence via martingale theory. The idea is to use conditional probabilities break up the probabilities comprising the social evidence. For instance we can break up $P(d_{t,n}^{(2)}, h_{t-1}^{(1)}, d_{t,0:n-1}^{(2)}, d_{t,0:n-1}^{(1)} | \xi_{0:t}^{(1)}, H = 1)$ into:

$$P(h_{t-1}^{(1)}, d_{t,0:n-1}^{(2)}, d_{t,0:n-1}^{(1)} | \xi_{0:t}^{(1)}, H = 1) P(d_{t,n}^{(2)} | h_{t-1}^{(2)}, h_{t-1}^{(1)}, d_{t,0:n-1}^{(2)}, d_{t,0:n-1}^{(1)}, \xi_{0:t}^{(1)}, H = 1).$$

Then the term on the left should cancel out with the corresponding term for $H = -1$ and then we just need to show

$$\lim_{n \rightarrow \infty} P(d_{t,n}^{(2)} | h_{t-1}^{(2)}, h_{t-1}^{(1)}, d_{t,0:n-1}^{(2)}, d_{t,0:n-1}^{(1)}, \xi_{0:t}^{(i)}, H = 1)$$

converges. If we define the history of an agent up to this equilibrium step:

$$\mathcal{H}_n^{(1)} = \{h_{t-1}^{(2)}, h_{t-1}^{(1)}, d_{t,0:n-1}^{(2)}, d_{t,0:n-1}^{(1)}\}$$

we observe that $\mathcal{H}_n^{(1)} \subseteq \mathcal{H}_{n+1}^{(1)}$ and thus, we can use martingale convergence to show that the agents eventually equilibrate.

Furthermore we conjecture that when both agents are undecided $\lim_{n \rightarrow \infty} \mathcal{H}_n^{(1)} = \lim_{n \rightarrow \infty} \mathcal{H}_n^{(2)}$ and the evidence agents gain from observing each other's decision

5.2. EQUILIBRATION PROCESS

states converges to the same value for both agents. This does not mean that the agents will have the same belief at the end of the equilibration process, because they made different private observations. However, the evidence they gain from social information should converge to the same value.

This result gives us convergence, but does not tell us what computations the agents use, or how quickly their beliefs equilibrate.

We will give the concrete calculation for the equilibration at time step $t = 0$. As we will see, showing the computation as time moves forward introduces another difficulty, so this isolates the calculation of the equilibration process.

We assume that both agents make their first observation at time step $t = 0$, and the evidence it provides is insufficient to cause either one to decide. If one of them does decide, the process is equivalent to the unidirectional case, as the agent that observes the decision updates its belief by an amount equal to the threshold that has been crossed, and continues accumulating information. We therefore describe only steps during which no decisions take place.

No decisions after the first observation imply $y_{0,0}^{(1)}, y_{0,0}^{(2)} \in (\theta_-, \theta_+)$. We now view the computation from the perspective of agent 1 because the computation is identical for the other agent.

First, because the evidence from their initial observations was insufficient to make a decision, it follows that $d_{0,0}^{(i)} = 0$ for $i = 1, 2$. This tells agent 1 that $y_{0,0}^{(2)} \in$

5.2. EQUILIBRATION PROCESS

$(\theta_-, \theta_+) =: \Theta$, and its first belief update during the equilibration process is:

$$y_{0,1}^{(1)} = \text{OB}_0^{(1)} + \log \frac{P(d_{0,0}^{(2)} = 0 | H = 1)}{P(d_{0,0}^{(2)} = 0 | H = -1)} = \text{OB}_0^{(1)} + \log \frac{P(y_{0,0}^{(2)} \in \Theta | H = 1)}{P(y_{0,0}^{(2)} \in \Theta | H = -1)},$$

with an equivalent expression for agent 2.

Assuming no decision has been made after the first observation, we have we have $d_{0,1}^{(i)} = 0$ and

$$\left| \log \frac{P(y_{0,0}^{(2)} \in \Theta | H = 1)}{P(y_{0,0}^{(2)} \in \Theta | H = -1)} \right| < |\theta_-| + \theta_+$$

otherwise the agent would have accumulated enough evidence to make a decision. Thus agent 1 knows the following

$$\theta_- < y_{0,0}^{(2)} < \theta_+ \quad (5.2.1)$$

$$\theta_- < y_{0,0}^{(2)} + \log \frac{P(y_{0,0}^{(1)} \in \Theta | H = 1)}{P(y_{0,0}^{(1)} \in \Theta | H = -1)} < \theta_+ \quad (5.2.2)$$

where we see in the second line how agent 2 incorporates its knowledge about of agent 1's belief. Let us denote this piece of evidence,

$$E_{0,0}^{(i)} := \log \frac{P(y_{0,0}^{(i)} \in \Theta | H = 1)}{P(y_{0,0}^{(i)} \in \Theta | H = -1)}.$$

Depending on the relative sizes of the thresholds and the evidence distributions $E_{0,0}^{(i)}$ may be positive, negative, or zero. If $E_{0,0}^{(i)} \geq 0$ then (5.2.2) becomes

$$\theta_- < y_{0,0}^{(\neg i)} < \theta_+ - E_{0,0}^{(i)}$$

and if $E_{0,0}^{(i)} < 0$ it is

$$\theta_- - E_{0,0}^{(i)} < y_{0,0}^{(\neg i)} < \theta_+.$$

5.2. EQUILIBRATION PROCESS

Thus, if there have been no decisions after the first equilibration step, agent i , knows the following about the belief of agent $\neg i$

$$\max\{\theta_- - E_{0,0}^{(i)}, \theta_-\} < y_{0,0}^{(\neg i)} < \min\{\theta_+ - E_{0,0}^{(i)}, \theta_+\}.$$

This allows agent i to update its belief again:

$$y_{0,2}^{(i)} = \text{OB}_0^{(i)} + \log \frac{P(y_{0,0}^{(\neg i)} \in (\max\{\theta_- - E_{0,0}^{(i)}, \theta_-\}, \min\{\theta_+ - E_{0,0}^{(i)}, \theta_+\}) | H = 1)}{P(y_{0,0}^{(\neg i)} \in (\max\{\theta_- - E_{0,0}^{(i)}, \theta_-\}, \min\{\theta_+ - E_{0,0}^{(i)}, \theta_+\}) | H = -1)}.$$

This illustrates the steps in the general equilibration process resulting in a sequence of intermediate belief updates, $y_{0,n}^{(i)}$. While the process may seem complex, it can be explained simply: At each equilibration step, the absence of a decision provides bounds on the neighbor's belief at the beginning of the process, that is, after the latest private observation. These bounds give a sequentially tighter bound on the belief of the opposite agent. When the bounds on the belief of the opposite agent do not change from one step to the next, then nothing new is gained by learning the non-decision and the process ends. Otherwise, new evidence is gained and the process is repeated.

Importantly, this process is *deterministic*: Given the initial observations of both agents, it will always produce the same sequence of bounds on their beliefs. We can describe the general steps in the form of an algorithm:

1. Denote

$$E_{0,l} = \log \frac{P(y_{0,l}^{(2)} \in \Theta | \xi_0^{(1)}, H = 1)}{P(y_{0,l}^{(2)} \in \Theta | \xi_0^{(1)}, H = -1)},$$

where the conditioning has to be stated because doing this calculation is dependent on $\xi_0^{(1)} \in \Theta$.

5.2. EQUILIBRATION PROCESS

2. From the previous step, agent 1 has bounds on agent 2's initial observation

$$\max_{0 \leq l \leq n-2} \{\theta_- - E_{0,l}, \theta_-\} < y_{0,0}^{(2)} < \min_{0 \leq l \leq n-2} \{\theta_+ - E_{0,l}, \theta_+\}.$$

3. Then if $E_{0,n-1}$ tightens either bound then $d_{0,n-1}^{(1)}$ was informative. Meaning if

$$\theta_- - E_{0,n-1} > \max_{0 \leq i \leq n-2} \{\theta_- - E_{0,i}, \theta_-\}$$

or

$$\theta_+ - E_{0,n-1} < \min_{0 \leq i \leq n-2} \{\theta_+ - E_{0,i}, \theta_+\}$$

then $y_{0,n}^{(1)} \neq y_{0,n-1}^{(1)}$.

4. We then compute

$$y_{0,n}^{(1)} = \text{OB}_0^{(1)} + \log \frac{P(y_{0,0}^{(2)} \in [a_{n-1}, b_{n-1}] | \xi_0^{(1)}, H = 1)}{P(y_{0,0}^{(2)} \in [a_{n-1}, b_{n-1}] | \xi_0^{(1)}, H = -1)}$$

with

$$Y_{0,n-1} = [\max_{0 \leq i \leq n-1} \{\theta_- - E_{0,i}, \theta_-\}, \min_{0 \leq i \leq n-1} \{\theta_+ - E_{0,i}, \theta_+\}],$$

the range of values that $y_{0,0}^{(i)}$ could have been equal to without causing the agent to make a decision by equilibration step $n - 1$.

5. Then, the agent uses $y_{0,n}^{(1)}$ to compute its decision state $d_{0,n-1}^{(1)}$. Then the process continues and only terminates when $E_{0,n} = E_{0,n+1}$ because in that case the belief update will no longer change.

We use simulations for the probability mass which allow us to approximate the probability that $y_0^{(i)}$ is in a certain interval. This allows us to simulate the

equilibration process and preliminary results show that it usually terminates after one step, which provides evidence that the process is finite and moreover stops after at most 2 steps. Hence we conjecture $y_{t,2}^{(i)} = y_t^{(i)}$.

5.3 Recursive Process in Time

In the previous section we described the equilibration process after the first private observation. We note that this process is only necessary when the thresholds or evidence distributions are not symmetric, as otherwise the absence of a decision provides no new information to an observing agent. We now want to give a general description of what happens after a private observation at an arbitrary time t .

As in the equilibration process, computing $y_{t,0}$ will require recursive inference, that is, it requires using how much a neighboring agent updates its belief in response to the decision state at previous time steps, which were in turn calculated using the decision states computed by that neighbor at previous time steps, and so on.

We can think of decision states as revealing partial information about an agent's belief. If an agent is undecided, it must have had evidence greater than θ_- and less than θ_+ , otherwise it would have picked $H = -1$ or $H = 1$, respectively. Therefore, the history of decisions gives a sequence of bounds on the belief of the opposing agent over time. Agents know how beliefs are computed, so this allows

them to infer bounds on the observation evidence alone by subtracting off this known amount of social evidence. This means the agents can convert the history of belief bounds into a history of bounds on the private evidence.

The initial part, for $n = 0$, of the evidence at time step t is simply what appeared in Eq. (5.1.1). We next show how to compute $y_{t,1}^{(1)}$, under the assumption that both agents have remained undecided after the observation at t . We see from Eq. (5.1.2) that we need to compute

$$y_{t,1}^{(1)} = \text{OB}_{0:t}^{(1)} + \log \frac{P(d_{t,0}^{(2)} = 0 | \xi_{0:t}^{(1)}, H = 1)}{P(d_{t,0}^{(2)} = 0 | \xi_{0:t}^{(1)}, H = -1)},$$

since

$$P(d_{t,0}^{(2)} = 0, h_{t-1}^{(i)}, d_{t,0:n-1}^{(\neg i)}, d_{t,0:n-1}^{(i)} | \xi_{0:t}^{(i)}, H = 1) = P(d_{t,0}^{(2)} = 0 | \xi_{0:t}^{(1)}, H = 1)$$

because decisions are immutable. Because agents do not communicate their exact evidence this can be simplified

$$y_{t,1}^{(1)} = \text{OB}_{0:t}^{(1)} + \log \frac{P(d_{t,0}^{(2)} = 0 | d_{0:t}^{(1)} = 0, H = 1)}{P(d_{t,0}^{(2)} = 0 | d_{0:t}^{(1)} = 0, H = -1)}.$$

To go further we rewrite the numerator of the decision state term:

$$P(d_{t,0}^{(2)} = 0 | d_{0:t}^{(1)} = 0, H = 1) = P(d_{t,0}^{(2)} = 0, d_{t-1,n}^{(2)} = 0 | d_{0:t}^{(1)} = 0, H = 1).$$

This can be written as the product

$$P(d_{t,0}^{(2)} = 0 | d_{t-1,n}^{(2)} = 0, d_{0:t}^{(1)} = 0, H = 1) P(d_{t-1,n}^{(2)} = 0 | d_{0:t}^{(1)} = 0, H = 1).$$

Next, we again have $d_t^{(1)} = 0 \implies d_{0:t}^{(1)} = 0$ so this becomes

$$P(d_{t,0}^{(2)} = 0 | d_{t-1,n}^{(2)} = 0, d_t^{(1)} = 0, H = 1) P(d_{t-1,n}^{(2)} = 0 | d_t^{(1)} = 0, H = 1).$$

If the probability of a non-decision is not altered significantly by the fact that it will result in a non-decision for their counterpart, i.e., if

$$P(d_{t-1,n}^{(2)} = 0 | d_t^{(1)} = 0, H = 1) \approx P(d_{t-1,n}^{(2)} = 0 | d_{t-1}^{(1)} = 0, H = 1)$$

then

$$y_{t,1}^{(1)} \approx \text{OB}_t^{(1)} + \log \frac{P(d_{t,0}^{(2)} = 0 | d_{t-1,n}^{(2)} = 0, d_t^{(1)} = 0, H = 1)}{P(d_{t,0}^{(2)} = 0 | d_{t-1,n}^{(2)} = 0, d_t^{(1)} = 0, H = -1)} + y_{t-1,1}^{(1)}.$$

However, it is not tractable to go much further than this, and all that can really be done is computing complicated equations using concrete equations, which is beyond the scope of what we do here. Rather, we will generalize the interactions described in Chapter 4 and this chapter under the assumption of symmetric evidence distributions and decision thresholds.

Chapter 6

General Network Accumulation

We have examined interactions between two rational agents accumulating information and observing each other's resulting decisions. We next extend this analysis to larger networks of agents. We first discuss different three agent networks to motivate the theory. Information sharing in three agent networks has aspects that are not observed with two agents, but are characteristic of larger networks of agents. In particular, a rational agent in a network with more than two agents may need to take into account the impact of unobserved decisions of next nearest neighbors. The required computation can be complex, and we only consider the symmetric case.

We next consider arbitrarily large fully connected networks (cliques). We provide an explicit description of the evidence update process and obtain conjectures for asymptotic results. We compare this to conjectured behavior for an arbitrarily large directed line.

As noted above, throughout this chapter we assume the decision boundaries are symmetric. This results in more tractable computations as the absence of a decision is not informative about the belief of any agent, hidden or visible. Therefore, social information is only communicated when an agent makes a decision. As a decision leads to a jump in the belief of neighboring agents, it typically ignites a cascade of decisions across the network. Although we assume symmetric boundaries and symmetric conditional evidence distributions, once an agent makes a decision symmetry can be broken. As a result, after an agent observes a decision of a neighbor, non-decisions of its other neighbors may become informative about their beliefs. In this case, which we do not analyze in detail here, rational agents must employ survival probabilities to update their own beliefs, as discussed in Section 6.1.3.

6.1 Accumulation of Evidence on General Networks

We will consider general directed networks with N rational agents accumulating measurements, and observing each other's decisions, as in the previous chapters. After introducing terminology and notation we first prove an important result on non-decisions: As in the case of two agents, when agents on a larger network all have symmetric decision thresholds and measurement distributions non-decisions of their neighbors are uninformative. Finally, we describe the main challenge posed by the need to take into account all possible decisions of unobserved agents via marginalization.

6.1.1 Terminology and Notation

Each agent i , for $1 \leq i \leq N$, makes a private observation at every time step. After incorporating the evidence from this observation the agent then updates its decision state and shares it with its neighbors. A directed edge from agent i to j , denoted by $i \rightarrow j$, means that agent i communicates its decisions to agent j . The set of neighbors that agent i observes (receives information from) is denoted by $N^{(i)}$:

$$N^{(i)} = \{j : j \rightarrow i\}.$$

Agent i thus receives social information from all agents in $N^{(i)}$, but agent i also needs to take into account decisions of *unobserved* agents, and we therefore define

$$U^{(i)} = \{j : j \notin N^{(i)} \cup \{i\}\}.$$

The entire set of agents – the set of all nodes in the network – can therefore be partitioned as $N^{(i)} \cup \{i\} \cup U^{(i)}$.

Recall that an agent makes decisions based on private information obtained from a sequence of measurements, and social information obtained from observing the decision states of its neighbors. We denote by $\text{OB}_{0:t}^{(i)}$ the information agent i has received from its own observations up to time t :

$$\text{OB}_{0:t}^{(i)} = \sum_{l=0}^t \ln \frac{P(\xi_l^{(i)} | H = 1)}{P(\xi_l^{(i)} | H = -1)}. \quad (\text{Observation Information})$$

We will denote the set of decisions by the neighbors of agent i at time t by $d_t^{N^{(i)}} = \{d_l^{(k)} : k \in N^{(i)}\}$. Similarly, the set of the decisions by unobserved agents

is denoted by $d_t^{U(i)}$. More generally, $d_{0:t}^{N(i)}$ denotes the sequence of decision states of the neighbors of agent i up to and including time t : $d_{0:t}^{N(i)} = \{d_s^{N(i)} : 0 \leq s \leq t\}$.

Finally, we have seen that the evidence from decisions depends on the time the decision was first made, i.e., the time step T such that $d_T^{(i)} \neq 0$ but $d_{T-1}^{(i)} = 0$. Hence will use $\bar{d}_T^{(i)} = \pm 1$ to denote $d_{T-1}^{(i)} = 0, d_T^{(i)} = \pm 1$ as a notational convenience.

6.1.2 Marginalization

Suppose that agents i, j , have neighborhoods such that $N^{(j)} \not\subset N^{(i)} \cup \{i\}$, that is agent j has neighbors that agent i does not observe directly. Agent i will therefore not know if a decision of agent j was solely due to j 's private observations, or has been influenced by social information from agents not directly observed by i . If j makes a decision $d_t^{(j)}$, we can only conclude that the resulting belief update of i should be $\text{DEC}_t^{(j)} = d_t^{(j)}\theta$ if $N^{(j)} \subset N^{(i)} \cup \{i\}$.

As the decision of agent j may have been caused by unobserved agents, agent i must marginalize over ("take into account") all possible decision states of the unobserved neighbors in order to accurately update its belief, $P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)} | H = 1)$. We can think of $d_s^{U(i)}$ as a decision tuple at time s . Therefore $d_s^{U(i)} = \vec{c} \in \{-1, 0, 1\}^{\#U(i)}$ where the l -th component $c_l = d_{s,l}^{U(i)} = d_s^{(l)}$ corresponds to the decision of agent $l \in U(i)$ at time s . However, in order to fully marginalize we have to take into account the full decision history of each unobserved agent. To do so we denote by $\mathcal{C}_t^{U(i)}$ all distinct tuples of histories of the unobserved agents:

$$\mathcal{C}_t^{U(i)} = \{d_{0:t}^{(k)} : k \in U(i), d_{0:t}^{(k)} \in \{-1, 0, 1\}^{t+1}\}.$$

Thus we can write the general belief update of agent i as:

$$\begin{aligned}
 y_t^{(i)} &= \log \frac{P(I_t^{(i)} | H = 1)}{P(I_t^{(i)} | H = -1)} \\
 &= \log \frac{P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)} | H = 1)}{P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)} | H = -1)} \\
 &= \log \frac{\sum_{d_{0:t}^{U(i)} \in \mathcal{C}_t} P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)}, d_{0:t}^{U(i)} | H = 1)}{\sum_{d_{0:t}^{U(i)} \in \mathcal{C}_t} P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)}, d_{0:t}^{U(i)} | H = -1)}.
 \end{aligned}$$

Therefore to update its belief, an agent must compute the joint probability of its observed evidence and marginalize over all possible unobserved decision histories. We will give an example of such marginalization in Section 6.2.2.

6.1.3 Pre-decision

We now show an important result on networks where all agents have symmetric thresholds and are drawing observations from symmetric evidence distributions. Before an agent or any of its neighbors has made a decision, the agent's belief $y_t^{(i)}$ has two properties that make computations somewhat simpler. First, the belief is a sum of terms corresponding to private and social information. Second, the social information is uninformative until a decision is made.

Proposition 8. *Assume all agents have symmetric thresholds and evidence distributions. At time t , agent i 's information consists of $I_t^{(i)} = \{\text{OB}_{0:t}^{(i)}, d_{0:t}^{N(i)}\}$. If neither agent i nor any of its neighbors in $N^{(i)}$ have made a decision by time t then*

1. $I_t^{(i)} = \{\text{OB}_{0:t}^{(i)}\} \cup \{d_t^{(k)} = 0 : k \in N^{(i)}\},$

$$2. y_t^{(i)} = \text{OB}_{0:t}^{(i)}$$

Proof. The first claim follows as in Proposition 1. Agents cannot change a non-zero decision state and thus a non-decision at time t , $d_t^{(k)} = 0$, implies a non-decision for previous times before as well, $d_{0:t}^{(k)} = 0$.

The second claim follows similarly to reasons discussed in Sections 4.3.2 and 5.1. However, we need to marginalization over unobserved decisions to establish the result.

First we can split the observation evidence from the social evidence conditioned on the agent's observations. Hence:

$$\begin{aligned} y_t^{(i)} &= \log \frac{P(I_t^{(i)} | H = 1)}{P(I_t^{(i)} | H = -1)} \\ &= \log \frac{P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)} = 0 | H = 1)}{P(\xi_{0:t}^{(i)}, d_{0:t}^{N(i)} = 0 | H = -1)} \\ &= \log \frac{P(\xi_{0:t}^{(i)} | H = 1)}{P(\xi_{0:t}^{(i)} | H = -1)} + \log \frac{P(d_{0:t}^{N(i)} = 0 | \xi_{0:t}^{(i)}, H = 1)}{P(d_{0:t}^{N(i)} = 0 | \xi_{0:t}^{(i)}, H = -1)} \end{aligned}$$

We therefore need to show that

$$\begin{aligned} 0 &= \log \frac{P(d_{0:t}^{N(i)} = 0 | \xi_{0:t}^{(i)}, H = 1)}{P(d_{0:t}^{N(i)} = 0 | \xi_{0:t}^{(i)}, H = -1)} \\ &= \log \frac{\sum_{d_{0:t}^{U(i)} \in \mathcal{C}_t^{U(i)}} P(d_{0:t}^{N(i)} = 0, d_{0:t}^{U(i)} | \xi_{0:t}^{(i)}, H = 1)}{\sum_{d_{0:t}^{U(i)} \in \mathcal{C}_t^{U(i)}} P(d_{0:t}^{N(i)} = 0, d_{0:t}^{U(i)} | \xi_{0:t}^{(i)}, H = -1)}. \end{aligned}$$

For a decision history of the unobserved agents $d_{0:t}^{U(i)}$, let $-d_{0:t}^{U(i)}$, be the opposite decision history. The negation changes the sign of each decision in the vector

6.1. ACCUMULATION OF EVIDENCE ON GENERAL NETWORKS

$d_{0:t}^{U(i)}$, and leaves non-decisions unaffected. As a special case, the vector of zeroes is equal to its negation. The key point is that $P(d_{0:t}^{N(i)} = 0, d_{0:t}^{U(i)} | \xi_{0:t}^{(i)}, H = 1) = P(d_{0:t}^{N(i)} = 0, -d_{0:t}^{U(i)} | \xi_{0:t}^{(i)}, H = -1)$ for all $d_{0:t}^{U(i)}$ and thus we get cancellation because every term in the numerator has a corresponding equal term in the denominator. We have $P(d_{0:t}^{N(i)} = 0, d_{0:t}^{U(i)} | \xi_{0:t}^{(i)}, H = 1) = P(d_{0:t}^{N(i)} = 0, -d_{0:t}^{U(i)} | \xi_{0:t}^{(i)}, H = -1)$ due to the symmetry of the evidence distributions and thresholds. This is because for symmetric distributions and thresholds

$$P(y_t^{(j)} = z | H = 1) = P(y_t^{(j)} = -z | H = -1).$$

Thus, given $H = 1$, any history of unobserved decision states $d_{0:t}^{U(i)}$ occurs with the same probability as the history $-d_{0:t}^{U(i)}$ given $H = -1$. $\square$

Thus, before decisions are made, agents independently accumulate evidence and social information is only informative when some agent makes a decision.

Remark. *As in the case of two agents, a decision by any agent will lead to a jump in the belief of all observing neighbors. Suppose an agent k observes both the deciding agent and one of its neighbors, call it j . Then agent k must take into account the fact that agent j has updated its belief due to the observed decision. Agent k now knows that the belief of agent j must have increased by a discrete amount in accordance with the observed decision. Thus agent k has gained knowledge about the belief of agent j , even if j does not immediately decide. The symmetry is therefore broken, and henceforth the situation is similar to the asymmetric case discussed in Section 6.1.3. This introduces additional drift into the evolution of the belief of agent k until either it or agent j make a decision. We thus concentrate on the dynamics up to the first decision, but discuss the general case in more*

detail in some of our examples.

6.2 Three-Agent Networks

To illustrate the behavior of the decision-making process on general networks we investigate some 3-agent network examples. If the agents are labelled, then there are 2^6 possible directed network topologies for three agents because there are six possible connections (two directed edge possible between any of the three pairs). However this number is reduced when we ignore agent labeling, and consider only networks without isolated agents. Thus we can eliminate those which are mapped to others via symmetries in the dihedral group D_3 , i.e., structures which are reflections and rotations of each other. There are thus 16 distinct directed networks with three agents (see [31]) and since 3 of those are not weakly connected (i.e., contain an agent that does not receive or send information, as in the isolated network, and the unidirectional or bidirectional coupling with an isolated third agent) this results in the 13 networks depicted in Figure 6.2.1.

6.2.1 NS1: The Fully Connected Network

We first investigate the three-agent, fully-connected network, where all agents communicate with each other (NS1 in Fig. 6.2.1). A fully connected network, or a clique, is one of the most tractable general networks because all agents are symmetrically connected. As every agent observes all others, the marginalization

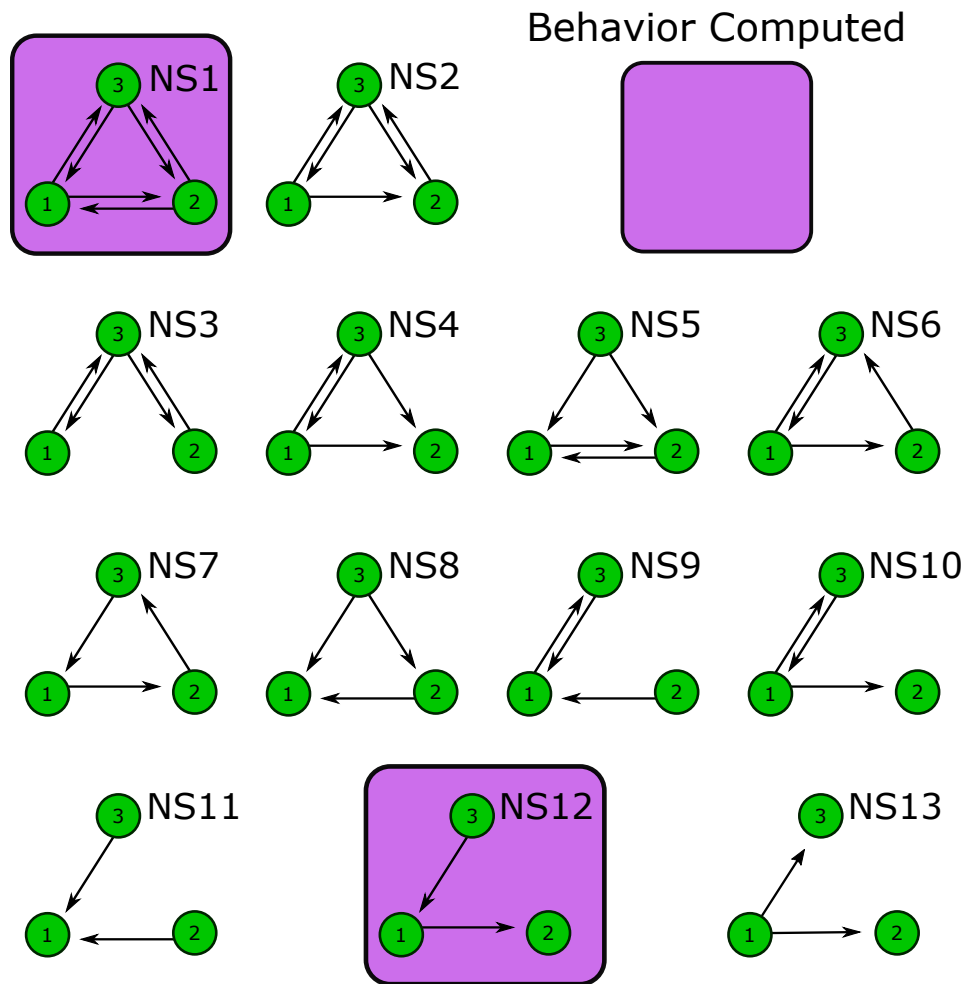


Fig. 6.2.1: We outline in pink the two networks whose behavior we analyze in detail. The other networks can be analyzed using a similar approach, or can be understood from the analysis of the two agent case (e.g. NS8 is essentially a two-agent network with an observer). We will refer to each network by the given label.

discussed in Section 6.1.2 is not required.

Before any agent makes a decision, the only evidence comes from private observations. Therefore the first non-zero social evidence arrives when some agent j first makes a choice at time T , that is $\bar{d}_T^{(j)} = \pm 1$. Without loss of generality, assume that agent 3 reaches decision $H = 1$ at time T and the other agents have not decided, $\bar{d}_T^{(3)} = 1, d_T^{(1)} = d_T^{(2)} = 0$. The two other agents then know that $y_T^{(3)} = \theta$ and thus add $d_T^{(3)}\theta$ to their belief. For $i = 1, 2$ we have,

$$\begin{aligned}
 y_T^{(i)} &= \text{OB}_{0:T}^{(i)} + \log \frac{P(d_T^{(\neg i)} = 0, \bar{d}_T^{(3)} = 1 | \xi_{0:T}^{(i)}, H = 1)}{P(d_T^{(\neg i)} = 0, \bar{d}_T^{(3)} = 1 | \xi_{0:T}^{(i)}, H = -1)} \\
 &= \text{OB}_{0:T}^{(i)} + \log \frac{P(d_T^{(\neg i)} = 0 | \xi_{0:T}^{(i)}, H = 1)}{P(d_T^{(\neg i)} = 0 | \xi_{0:T}^{(i)}, H = -1)} + \log \frac{P(\bar{d}_T^{(3)} = 1 | \xi_{0:T}^{(i)}, H = 1)}{P(\bar{d}_T^{(3)} = 1 | \xi_{0:T}^{(i)}, H = -1)} \\
 &= \text{OB}_{0:T}^{(i)} + 0 + \log \frac{P(\bar{d}_T^{(3)} = 1 | \xi_{0:T}^{(i)}, H = 1)}{P(\bar{d}_T^{(3)} = 1 | \xi_{0:T}^{(i)}, H = -1)} \\
 &= \text{OB}_{0:T}^{(i)} + \theta.
 \end{aligned}$$

The updated belief could be sufficient to make a decision at this time. For simplicity we assume that all agents wait until all resulting decisions are made before continuing to gather further private information. As we will see, there could be a number of decision-making rounds before evidence accumulation continues.

After the decision of agent 3, there are three possible outcomes:

- (i) If, before accounting for the decision $\bar{d}_T^{(3)} = 1$, $y_T^{(i)} \geq 0$ for both remaining agents, $i = 1, 2$, then both agents immediately choose $H = 1$.
- (ii) If, before accounting for the decision $\bar{d}_T^{(3)} = 1$, $y_T^{(i)} \geq 0$ for only one agent,

6.2. THREE-AGENT NETWORKS

then the corresponding agent decides $H = 1$, and the other remains undecided.

- (iii) If, before accounting for the decision $\bar{d}_T^{(3)} = 1$, $y_T^{(i)} < 0$ for both remaining agents, $i = 1, 2$, then both remain undecided.

Evidence accumulation stops in case (i). We therefore only examine the latter cases. As in the case of bidirectional interactions discussed in the previous chapter, agents will update their beliefs and communicate their decisions repeatedly. Hence we define the computations during a sequence of substeps following the decision of agent 3:

$$\begin{aligned} y_{T,0}^{(i)} &= \text{OB}_{0:T}^{(i)} \\ y_{T,1}^{(i)} &= \text{OB}_{0:T}^{(i)} + \theta \\ y_{T,n}^{(i)} &= \text{OB}_{0:T}^{(i)} + \theta + \log \frac{P(d_{T,n-1}^{(-i)} | \bar{d}_T^{(3)} = 1, d_{T,0:n-1}^{(-i)}, d_{T,0:n-1}^{(i)}, \xi_{0:T}^{(i)}, H = 1)}{P(d_{T,n-1}^{(-i)} | \bar{d}_T^{(3)} = 1, d_{T,0:n-1}^{(-i)}, d_{T,0:n-1}^{(i)}, \xi_{0:T}^{(i)}, H = 1)}, \end{aligned}$$

where $d_{T,n}^{(-i)}$ is the decision state of the other remaining agent after it has updated its belief, $y_{T,n}^{(-i)}$.

Case (ii) - One Agent Undecided

Assume that $y_{T,0}^{(2)} \geq 0$ so that agent 2 decides right after observing agent 3's decision, i.e., $d_{T,1}^{(2)} = 1$. After observing the decisions of agent 3 and then agent 2, the

remaining agent 1 updates its belief:

$$\begin{aligned}
 y_{T,2}^{(1)} &= \text{OB}_{0:T}^{(1)} + \ln \frac{P(\bar{d}_{T,1}^{(2)} = 1, \bar{d}_{T,0}^{(3)} = 1 | H = 1)}{P(\bar{d}_{T,1}^{(2)} = 1, \bar{d}_{T,0}^{(3)} = 1 | H = -1)} \\
 &= \text{OB}_{0:T}^{(1)} + \ln \frac{P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = 1) P(\bar{d}_{T,0}^{(3)} = 1 | H = 1)}{P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = -1) P(\bar{d}_{T,0}^{(3)} = 1 | H = -1)} \\
 &= \text{OB}_{0:T}^{(1)} + \ln \frac{P(\bar{d}_{T,0}^{(3)} = 1 | H = 1)}{P(\bar{d}_{T,0}^{(3)} = 1 | H = -1)} + \ln \frac{P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = 1)}{P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = -1)} \\
 &= \text{OB}_{0:T}^{(1)} + \theta + \ln \frac{P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = 1)}{P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = -1)}.
 \end{aligned} \tag{6.2.1}$$

Note, here $P(\bar{d}_{T,0}^{(3)} = 1 | H = 1)$ implicitly means that *only* agent 3 has made a decision at that time. The private and social information can be split in the first line because of Proposition 8.

The last term can be rewritten because the belief of agent 2 must have been non-negative before observing the decision of agent 3,

$$P(\bar{d}_{T,1}^{(2)} = 1 | \bar{d}_{T,0}^{(3)} = 1, H = 1) = P(0 \leq y_T^{(2)} \leq \theta | H = 1).$$

Hence

$$y_{T,2}^{(1)} = \text{OB}_{0:T}^{(1)} + \theta + \ln \frac{P(0 \leq y_{T,0}^{(2)} \leq \theta | H = 1)}{P(0 \leq y_{T,0}^{(2)} \leq \theta | H = -1)}. \tag{6.2.2}$$

The last term in the sum is the update to the belief of agent 1 due to an observation of the decision of agent 2. As a consequence of Proposition 9 below, this update is smaller than θ , but can cause agent 1 to decide as well. If not it will continue with its private observation until it does. We estimate the evidence from the last term in Eq. (6.2.2) when we look at arbitrarily large cliques in Section 6.3.

Case (iii): Two Agents Undecided

If both agents remain undecided then this means both agents had negative beliefs prior to observing the decision of agent 3. Following the same computations as in Eqs. (6.2.1, 6.2.2) we have

$$y_{T,2}^{(i)} = \text{OB}_{0:T}^{(2)} + \theta + \ln \frac{P(-\theta \leq y_T^{(\neg i)} \leq 0 | H = 1)}{P(-\theta \leq y_T^{(\neg i)} \leq 0 | H = -1)}. \quad (6.2.3)$$

for $i = 1, 2$, where $\neg i$ refers to the other undecided agent. Since the last term must be negative, both agents communicate $d_{T,2}^{(i)} = 0$ upon making this update. Thus after observing the decision of agent 3, and observing that both are still undecided each agent can conclude that the other has gathered evidence in favor of $H = -1$, and adjust their belief accordingly.

Let us denote the evidence gained by learning that an agent has observed a decision in favor of $H = 1$ but still not decided (the last term in Eq. (6.2.3)) by:

$$R_-^{(\neg i)} := \ln \frac{P(-\theta \leq y_T^{(\neg i)} \leq 0 | H = 1)}{P(-\theta \leq y_T^{(\neg i)} \leq 0 | H = -1)}.$$

Due to symmetry, $R_-^{(1)} = R_-^{(2)}$ so we will suppress the superscript. Then the fact that $d_{T,0}^{(i)} = d_{T,1}^{(i)} = d_{T,2}^{(i)} = 0$ means

$$-\theta < y_{T,0}^{(i)} + \theta + R_- < \theta$$

so that

$$-2\theta - R_- < y_{T,0}^{(i)} < -R_-.$$

This together with the previous bounds means

$$\max\{-\theta, -2\theta - R_-\} < y_{T,0}^{(i)} < \min\{0, -R_-\} = 0.$$

Thus if $-\theta < -2\theta - R_-$, then this gives a finer bound on the private evidence $y_{T,0}^{(i)}$: Not only was it negative, but it was also bigger than $-2\theta - R_-$. Note that this process is similar to the equilibration computations done in the case of two bidirectionally coupled agents. Thus a non-decision $d_{T,2}^{(i)} = 0$ will no longer provide evidence to agent $\neg i$ if $-\theta \geq -2\theta - R_-$, i.e., if $R_- \geq -\theta$. The fact that $R_- \geq -\theta$ and the process equilibrates follows from the next proposition. Intuitively, knowing that an agent's belief is bounded $|y_{T,0}^{(i)}| < \theta$ cannot provide more than a change of θ in the belief.

Proposition 9. *Assume that agent j has not made a decision and has not accumulated any evidence from social information by time T . Let $-\theta < a \leq b < \theta$. Then*

$$\left| \ln \frac{P(a \leq y_T^{(j)} \leq b | H = 1)}{P(a \leq y_T^{(j)} \leq b | H = -1)} \right| < \theta.$$

That is, the evidence received from a bound on agent j 's evidence is itself bounded by θ .

Proof. If the belief is negative, then we need to show that

$$-\theta < \ln \left[\frac{P(y_T^{(j)} \in [a, b] | H = 1)}{P(y_T^{(j)} \in [a, b] | H = -1)} \right],$$

or, equivalently, that

$$P(y_T^{(j)} \in [a, b] | H = 1) > e^{-\theta} P(y_T^{(j)} \in [a, b] | H = -1)$$

Since $y_t^{(j)}$ consists only of evidence from private observations then the observations satisfy

$$\prod_{t=0}^T P_+(\xi_t^{(j)}) > e^{-\theta} \prod_{t=0}^T P_-(\xi_t^{(j)}),$$

6.2. THREE-AGENT NETWORKS

where $P_{\pm}(\xi_t^{(j)}) = P(\xi_t^{(j)} | H = \pm 1)$.

If observations are drawn from the space Ξ , then we will integrate over the subset of the product space consisting of valid observation sequences $V^T \subset \Xi^T$:

$$V^T = \{(\xi_0^{(j)}, \dots, \xi_T^{(j)}) \in \Xi^T : \text{OB}_{0:t}^{(j)} \in (-\theta, \theta), \text{ for } t \leq T, \text{OB}_{0:T}^{(j)} \in [a, b]\}.$$

Since no decision is made for valid observation sequences:

$$\left(\prod_{t=0}^T P_+(\xi_t^{(j)}) / \prod_{t=0}^T P_-(\xi_t^{(j)}) \right) > e^{-\theta}, \quad \forall \xi_{0:T} \in V^T.$$

Letting P_{true} be the evidence distribution from which the observations are actually drawn, we apply the conditional independence of the observations to get:

$$P_+(\xi_{0:T}^{(j)}) P_{\text{true}}(\xi_{0:T}^{(j)}) > P_-(\xi_{0:T}^{(j)}) P_{\text{true}}(\xi_{0:T}^{(j)}) e^{-\theta}, \quad \forall \xi_{0:T} \in V^T$$

so if we integrate over V^T

$$\int_{V^T} P_+(\xi_{1:i}^{(j)}) \xi_{0:T}^{(j)} > \int_{V^T} P_-(\xi_{1:i}^{(j)}) e^{-\theta} d\xi_{0:T}^{(j)}.$$

The left hand side is the probability that, given $H = 1$, agent j does not make a decision by time T and has log-likelihood evidence bounded inside $[a, b]$ and the right hand side is similar, we get

$$P(y_T^{(j)} \in [a, b] | H = 1) > e^{-\theta} P(y_T^{(j)} \in [a, b] | H = -1),$$

which is what we set out to prove. □

Now that sharing decision states is no longer informative, both agents can continue collecting observations. But, unlike before the initial decision, non-decisions

are now informative because, as discussed in the last remark in Section 6.1.3, both agents will have unequal priors over the two states (their beliefs will no longer be equal to zero after equilibration). The process thus continues on as in the bidirectional coupling with asymmetric decision thresholds.

6.2.2 NS11: The 3-Agent Unidirectional Line

Next we consider network NS11 of Fig. 6.2.1. After relabeling and redrawing we obtain the network depicted in Fig. 6.2.2. We use the new labeling in what follows.

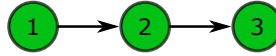


Fig. 6.2.2: Agents on a unidirectional Line

To begin, note that it suffices to examine the case when agent 2 makes a decision at time T , say $\bar{d}_T^{(2)} = 1$: If agent 3 makes a decision then it does not share the information with any other agent. If agent 1 makes a decision, then agent 2 reacts as in the case of two unidirectionally coupled agents discussed in Section 4.2. In that case we still have to analyze how agent 3 reacts after agent 2's eventual decision. Thus, the only case left to describe is how agent 3 reacts to $\bar{d}_T^{(2)} = 1$.

As we showed more generally in Section 6.1.3, there is no evidence from social information before agent 2 makes a decision. Thus we compute how agent 3 updates its belief after observing a decision of agent 2 by marginalizing over possible

decision histories of agent 1:

$$\begin{aligned}
 P(d_t^{(2)} = 1 | H = 1) &= P(d_t^{(2)} = 1, d_t^{(1)} = 0 | H = 1) \\
 &\quad + \sum_{s < t} P(d_t^{(2)} = 1, d_{s-1}^{(1)} = 0, d_s^{(1)} = -1 | H = 1) \\
 &\quad + \sum_{s < t} P(d_t^{(2)} = 1, d_{s-1}^{(1)} = 0, d_s^{(1)} = 1 | H = 1)
 \end{aligned}$$

The first term on the right hand side corresponds to the case when agent 2 decides before observing a decision of agent 1. Note that $P(d_t^{(2)} = 1, d_t^{(1)} = 0 | H = 1) = \alpha$, where α is the probability of making the correct decision as in the previous chapters. Thus if agent 3 knew agent 1 had not made a decision then it would update its belief by adding θ as in the unidirectional coupling. The other two terms on the right hand side correspond to the cases where agent 2 makes a decision after observing the decision of agent 1 at some prior time.

We intuitively compute the evidence when agent 1 makes the decision $H = -1$ at time $s < t$. The decision state $d_s^{(1)} = -1$ provides $-\theta$ evidence, and if agent 2 did not immediately decide, it had to of had negative evidence at time t . Thus at time s , the evidence should be $E[y_s^{(2)} | y_s^{(2)} > 0] - \theta$. Then agent 2 updates its belief by $-\theta$, so if it eventually chooses $H = 1$ it has to collect θ evidence to override this update, plus an additional $\theta - E[y_t^{(2)} | y_s^{(2)} > 0]$ evidence to accumulate to the $H = 1$ threshold. Hence

$$\text{DEC}(d_t^{(2)} = 1, d_s^{(1)} = -1) = E[y_s^{(2)} | y_s^{(2)} > 0] - \theta + \theta + \theta - E[y_t^{(2)} | y_s^{(2)} > 0] = \theta.$$

Similarly

$$\text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = 1) = \theta + E[y_t^{(2)} | y_s^{(2)} > 0]$$

and for $s < t$

$$\text{DEC}(d_t^{(2)} = 1, d_s^{(1)} = 1) = \theta = \text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = 0).$$

The issue is that we want to compute $\text{DEC}(d_t^{(2)} = 1)$, but we cannot simply break up decision evidence as a sum of possible decision evidence terms because the marginalization has to be done in the numerator and denominator of $\text{DEC}(d_t^{(2)} = 1)$, so:

$$\begin{aligned} \text{DEC}(d_t^{(2)} = 1) &\neq \text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = -1) + \text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = 0) \\ &\quad + \text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = 1) \end{aligned}$$

If $\text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = c) = \theta$ for all possible choices $c \in \{-1, 0, 1\}$ then this would mean $\frac{P(d_t^{(2)}=1, d_t^{(1)}=c|H=1)}{P(d_t^{(2)}=1, d_t^{(1)}=c|H=-1)}$ is the same fraction, and thus $\text{DEC}(d_t^{(2)} = 1)$ would also be θ . However, we should not expect this since $\text{DEC}(d_t^{(2)} = 1, d_t^{(1)} = 1) = \theta + E[y_t^{(2)} | y_s^{(2)} > 0] > \theta$, using our intuitive computation.

Since we have

$$\frac{P(d_t^{(2)} = 1, d_s^{(1)} = -1 | H = 1)}{P(d_t^{(2)} = 1, d_s^{(1)} = -1 | H = -1)} = e^\theta = \frac{P(d_t^{(2)} = 1, d_t^{(1)} = 0 | H = 1)}{P(d_t^{(2)} = 1, d_t^{(1)} = 0 | H = -1)}$$

and

$$\frac{P(d_t^{(2)} = 1, d_s^{(1)} = -1 | H = 1)}{P(d_t^{(2)} = 1, d_s^{(1)} = -1 | H = -1)} = e^\theta$$

then the possibilities where agent 2 did not immediately agree with agent 1 can

be written as

$$\begin{aligned}
 & P(d_t^{(2)} = 1, d_t^{(1)} = 0 | H = 1) = e^\theta P(d_t^{(2)} = 1, d_t^{(1)} = 0 | H = -1) \\
 & + \sum_{s < t} P(d_t^{(2)} = 1, d_s^{(1)} = -1 | H = 1) + e^\theta \sum_{s < t} P(d_t^{(2)} = 1, d_s^{(1)} = -1 | H = -1) \\
 & + \sum_{s < t} P(d_t^{(2)} = 1, d_s^{(1)} = 1 | H = 1) + e^\theta \sum_{s < t} P(d_t^{(2)} = 1, d_s^{(1)} = 1 | H = -1)
 \end{aligned}$$

so we define A to be the left hand side. Then $\text{DEC}(d_t^{(2)} = 1)$ is

$$\log \left(\frac{P(d_t^{(2)} = 1, d_t^{(1)} = 1 | H = 1) + A}{P(d_t^{(2)} = 1, d_t^{(1)} = 1 | H = -1) + e^{-\theta} A} \right).$$

Then $P(d_t^{(2)} = 1, d_t^{(1)} = 1 | H = 1)$ is the product of agent 2 having positive evidence and agent 1 choosing the correct choice at time t :

$$P(y_t^{(2)} \in [0, \theta) | H = 1) P(d_t^{(1)} = 1 | H = 1).$$

Let's assume

$$P(d_t^{(1)} = 1 | H = 1) = e^\theta P(d_t^{(1)} = 1 | H = -1)$$

and

$$P(y_t^{(2)} \in [0, \theta) | H = 1) = e^{R_+} P(y_t^{(2)} \in [0, \theta) | H = -1),$$

where

$$R_+ = \log \frac{P(y_t^{(2)} \in [0, \theta) | H = 1)}{P(y_t^{(2)} \in [0, \theta) | H = -1)}.$$

Then $0 \leq R_+ \leq \theta$. If θ is small we can approximate $e^\theta \approx e^{R_+}$ in which case we get

$$\text{DEC}(d_t^{(2)} = 1) \approx \theta$$

but is this is a somewhat special case, so all we are certain is that

$$\text{DEC}(d_t^{(2)} = 1) \geq \theta.$$

Thus, as in the two-agent interactions a decision provides an update at least equal to the threshold it corresponds to. However, with using the explicit distributions and decision times we cannot say by how much the update may exceed that threshold.

6.2.3 Other Networks

Notice that the analysis of the other 3 agent networks follow from the clique and line examples. In general, the process can be broken into 3 steps, as follows (we assume that the indices i, j, k all refer to different agents):

1. Agents accumulate private information until one, say agent i , decides. If no other agents observe i , then the remaining two agents behave as discussed in Chapters 4 and 5.
 - (a) In NS10, if agent 2 decides, then this becomes a bidirectional coupling example.
 - (b) In NS8, if agent 1 decides, then this becomes a unidirectional coupling example.
 - (c) In NS11, if agent 1 decides, then this becomes two uncoupled agents.

6.2. THREE-AGENT NETWORKS

2. Suppose agent j observes the decision of agent i , but does not have k as a neighbor. If i has k as a neighbor, then j must marginalize over the possible decision histories of k as in the unidirectional line.
 - (a) In NS3, if agent 3 decides, then agent 2 marginalizes over the possible decision histories of agent 1, and agent 1 marginalizes over the possible decision histories of agent 2.
 - (b) In NS2, if agent 3 decides, then agent 1 marginalizes over the possible decision histories of agent 2.
3. If an agent i has both other agents as neighbors, then it will behave as an agent in a clique. If it observes agent j deciding first, it will increase its belief by θ . If agent j 's decision followed that of agent k , then the belief update will be smaller, as discussed in the case of a clique. A similar process takes place when two agents observe the the third agent, and have each other as neighbors.
 - (a) In NS2, if agent 3 decides, then agent 2 does not have to marginalize over the possible decisions of agent 1.
 - (b) In NS5, if agent 3 decides, then agents 1 and 2 undergo a belief update as in the 3-agent clique.

6.3 Large Cliques: Simulations and Asymptotics

We next consider a clique of N agents who all have the same symmetric decision boundaries. We will show that the decision process has the following steps:

1. An agent makes a decision first. By symmetry, any agent is equally likely to make the first decision.
2. A group of other agents makes the same decision.
3. The remaining undecided agents compare the number of agents that decided in step 2 to the number that did not and update their beliefs using this distribution.
4. As we will see with simulations, after step 3, most agents will have decided with high probability. The remaining agents undergo an equilibration process similar to that in the case of two bidirectionally coupled agents, where they communicate their decision states until it is no longer informative. Any agents that have not decided by this point will continue integrating private information, but are likely to start with unequal priors over the two states. As this appears to happen infrequently, and symmetry has been broken, we do not analyze this case further.

Without loss of generality, we assume that agent 1 accumulates enough evidence to decide on $H = 1$. This means $y_t^{(1)} = \theta$, and thus every other agent updates their log-likelihood evidence

$$y_{t,1}^{(i)} = y_{t,0}^{(i)} + \theta, \forall i \neq 1.$$

As before agents will stop making observations and communicate their decision states until there is no more information to be gained. Observing $d_{t,0}^{(1)} = 1$ will cause any agent whose belief satisfies $y_{t,0}^{(i)} \geq 0$ to make the same decision. We call these the **agreeing agents**. The agents whose beliefs satisfy $y_{t,0}^{(i)} < 0$ do not make a decision, and we call them the **disagreeing agents**. Upon observing the decision of agent 1, the belief of any disagreeing agent satisfies $0 < y_{t,1}^{(i)} < \theta$. We give a schematic in Figure 6.3.1.

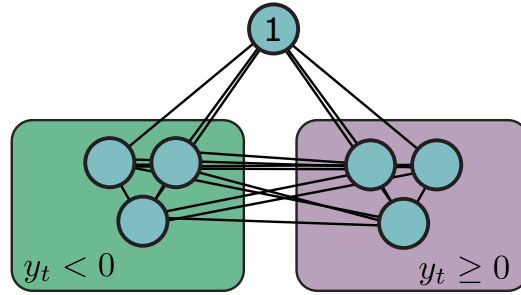


Fig. 6.3.1: The first decider, agent 1 is on top. The agreeing agents are on the right in the purple group and the disagreeing agents are on the left in the green group.

Intuitively, each agreeing agent now gives additional evidence for $H = 1$, while each disagreeing agent gives evidence for $H = -1$. By symmetry, each type of agent should provide the same magnitude of evidence one way or the other. Therefore, if more than half of the remaining agents follow the first and decide in favor of $H = 1$, this provides additional evidence for that choice. On the other hand, if the majority of agents stay silent, the evidence favors $H = -1$. We compute the corresponding belief update in Section 6.3.1.

6.3.1 Agreement Information

Before any agent makes a choice, the communicated decision states are uninformative. Once agent 1 chooses $H = 1$ at time T , a wave of agreeing agents follows. Let $N = a + d + 1$ where a is the number of agreeing agents who had positive evidence at time t , and d is the number of disagreeing agents who had negative evidence at time t . We let A be the set of agreeing agents and D be the set of disagreeing agents.

We want to know what a disagreeing agent $j \in D$ (so $d_{T,1}^{(j)} = 0$) learns by observing the distribution of A and D . To compute the evidence update for agent j first note:

$$\begin{aligned} P(d_{T,0}^{(1)} = 1, d_{T,1}^{(i)}, d_{T,1}^{(k)} | H) &= P(d_{T,1}^{(i)}, d_{T,1}^{(k)} | d_{T,0}^{(1)} = 1, H) P(d_{T,0}^{(1)} = 1 | H = 1) \\ &= P(d_{T,1}^{(i)} | d_{T,0}^{(1)} = 1, H) P(d_{T,1}^{(k)} | d_{T,0}^{(1)} = 1, H) P(d_{T,0}^{(1)} = 1 | H = 1), \end{aligned}$$

by a rule of conditional probability and the independence of the observations of the agents. From this we get

$$y_{T,1}^{(j)} = y_{T,0}^{(j)} + \log \frac{P(d_{T,0}^{(1)} = 1 | H = 1)}{P(d_{T,0}^{(1)} = 1 | H = -1)}$$

and

$$\begin{aligned} y_{T,2}^{(j)} &= y_T^{(j)} + \log \frac{P(d_{T,0}^{(1)} = 1 | H = 1)}{P(d_{T,0}^{(1)} = 1 | H = -1)} \\ &\quad + \sum_{i \in A} \log \frac{P(d_{T,1}^{(i)} = 1 | d_{T,0}^{(1)} = 1, H = 1)}{P(d_{T,1}^{(i)} = 1 | d_{T,0}^{(1)} = 1, H = -1)} \\ &\quad + \sum_{k \in D} \log \frac{P(d_{T,1}^{(k)} = 0 | d_{T,0}^{(1)} = 1, H = 1)}{P(d_{T,1}^{(k)} = 0 | d_{T,0}^{(1)} = 1, H = -1)}. \end{aligned}$$

This simplifies to:

$$y_{T,2}^{(j)} = y_{T,0}^{(j)} + \theta + a \log \frac{P(0 \leq y_{T,0}^{(i)} < \theta | H = 1)}{P(0 \leq y_{T,0}^{(i)} < \theta | H = -1)} \\ + d \log \frac{P(-\theta < y_{T,0}^{(k)} < 0 | H = 1)}{P(-\theta < y_{T,0}^{(k)} < 0 | H = -1)}$$

by independence.

Note

$$P(0 \leq y_{T,0}^{(i)} < \theta | H = 1) = P(0 \leq y_{T,0}^{(i)} | y_{T,0}^{(i)} \in \Theta, H = 1) P(y_{T,0}^{(i)} \in \Theta | H = 1)$$

with $\Theta = (-\theta, \theta)$. So we define $S_{\pm,T} = P(y_{T,0}^{(i)} \in \Theta | H = \pm 1)$ as well as

$$R_{\pm,T} = P(0 \leq y_{T,0}^{(i)} | y_{T,0}^{(i)} \in \Theta, H = \pm 1), L_{\pm,T} = P(y_{T,0}^{(i)} < 0 | y_{T,0}^{(i)} \in \Theta, H = \pm 1).$$

Then we suppress the index which denotes the dependence on the initial decision time T and compute:

$$y_{T,2}^{(j)} - y_{T,0}^{(j)} - \theta = a \log \frac{P(0 \leq y_{T,0}^{(i)} < \theta | H = 1)}{P(0 \leq y_{T,0}^{(i)} < \theta | H = -1)} + d \log \frac{P(-\theta < y_{T,0}^{(k)} < 0 | H = 1)}{P(-\theta < y_{T,0}^{(k)} < 0 | H = -1)} \\ = a \log \frac{R_+ S_+}{R_- S_-} + d \log \frac{L_+ S_+}{L_- S_-} \\ = (a + d) \log \frac{S_+}{S_-} + a \log \frac{R_+}{R_-} + d \log \frac{L_+}{L_-}$$

The first term here is zero when we assume the distributions and boundaries are symmetric. We also have $R_+ = L_-$ because the probability that the evidence is accumulating in the correct direction given the state is independent of the state when the evidence distributions are symmetric, and similarly, $R_- = L_+$. Thus

$$y_{T,2}^{(j)} = y_{T,0}^{(j)} + \theta + (a - d) \log \frac{R_{+,T}}{R_{-,T}}.$$

6.3. LARGE CLIQUES: SIMULATIONS AND ASYMPTOTICS

If $(a - d) \log \frac{R_{+,T}}{R_{+,T}} < -2\theta$ (when d is sufficiently larger than a) we have a somewhat counterintuitive situation: If insufficiently many agents **fail** to follow the decision of agent 1, then in the next step every agent in D chooses $H = -1$!

We thus refine the process described in the beginning of this section.

1. Agent 1 decides $H = 1$ and the positive evidence agents choose $H = 1$.
2. The disagreeing agents all jump by $\theta + (a - d) \log \frac{R_+}{R_-}$.
3. This may cause some disagreeing agents to make a choice:
 - (i) if $(a - d) \log \frac{R_+}{R_-} \geq \theta$ all agents choose $H = 1$;
 - (ii) if $(a - d) \log \frac{R_+}{R_-} < -2\theta$ all agents in D choose $H = -1$, and thus we have $a + 1$ agents choosing $H = 1$ versus d agents choosing $H = -1$;
 - (iii) if $-\theta \leq (a - d) \log \frac{R_+}{R_-} \leq 0$, for instance, when $a = d$, then no agent in D will make a decision and the process might equilibrate here;
 - (iv) if $-2\theta < (a - d) \log \frac{R_+}{R_-} < -\theta$ then some of the agents in D may choose $H = -1$ and this will give the remaining agents a bound on where their LLRs must be (similar to the $n = 3$ case);
 - (v) if $0 < (a - d) \log \frac{R_+}{R_-} < \theta$ then some of the agents in D may choose $H = 1$ and this will give the remaining agents a bound on where their LLRs must be (similar to the $n = 3$ case).

In cases (i) and (ii) we are done. For cases (iii)-(v) we must carry out an analysis similar to that in section 6.2.1.

Asymptotics

Let α_N be the probability that the first agent to decide in an N -agent clique makes the correct decision. We conjecture that as the clique size $N \rightarrow \infty$, the probability that the majority of the agents make the correct decision goes to 1; furthermore, all agents will make the correct decision with probability α_N .

In Figures 6.3.2, 6.3.3, and 6.3.4 we plot information about the behavior of cliques of size 10, 100, and 1000, respectively. We assume the correct state is $H = 1$. As expected, the more agents in the clique, the earlier a first agent will make a decision. Also, the probability that a given agent will have the correct belief (meaning here that its belief is positive) is independent of the network size.

We also plot the probability the exactly k agents have the correct belief for various values of $k < N$ to show that a majority of agents will have the correct belief and thus observing the ratio of a to d will be informative. We see that most of the agents will make the correct choice immediately after the first agent decides.

Thus two things can happen. If the first agent is wrong, the majority of agents will have positive evidence and this will cause the expected evidence to be large enough for those agents to make the correct decision. Alternatively, that agent is correct, and all agents will agree with it and a unanimous decision will be reached. Intuitively, this is one difference between cliques and unidirectional lines: we expect unanimous decisions on unidirectional lines to be rare.

We need to further explore to confirm that the probability of having a majority of agents with correct belief overwhelms how early decisions occur. As a start, we

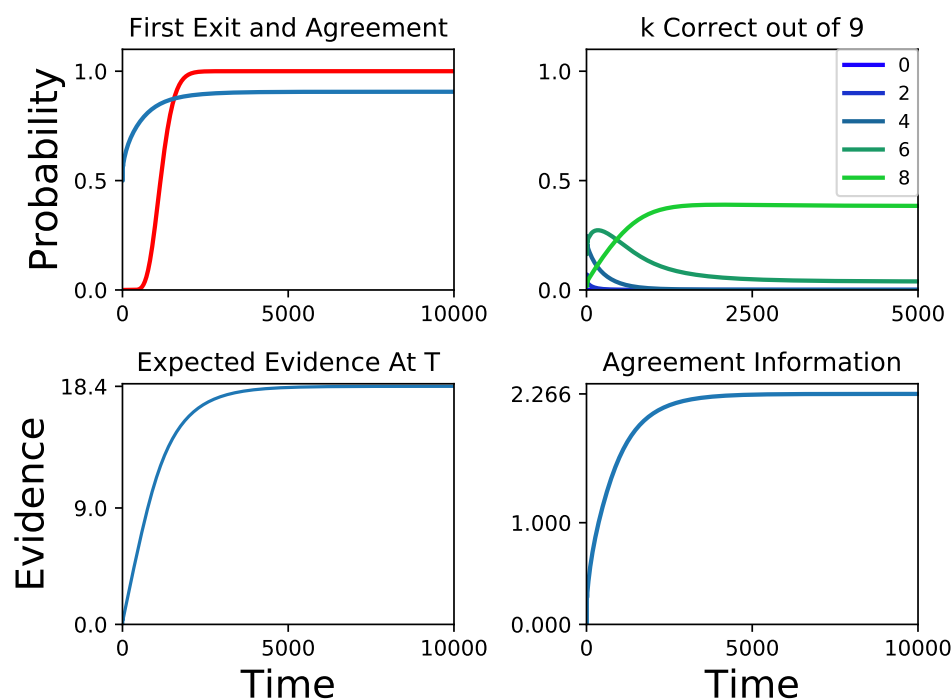


Fig. 6.3.2: Simulations for a fully-connected network with 10 agents for $H = 1$. Top left: probability some agent decides (red) by time t and and probability a given agent has positive belief (blue). Top right: the probability exactly k agents have positive belief. Bottom left: the expected amount of evidence gained after seeing the distribution of agreeing and disagreeing agents. Bottom right: how much evidence is gained per agreeing agent.

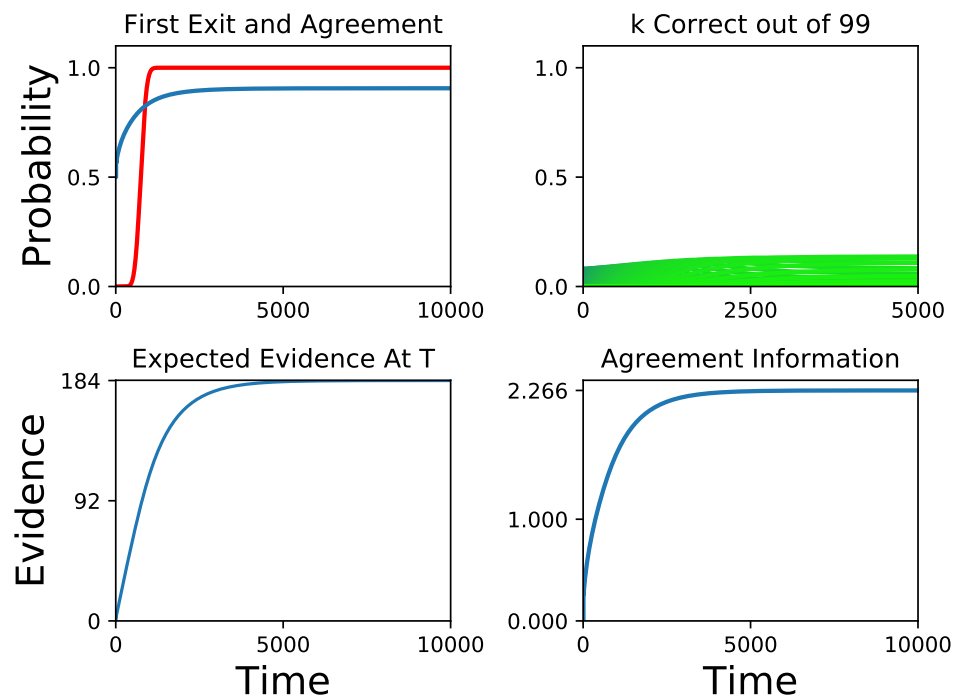


Fig. 6.3.3: The plots are the same as in Figure 6.3.2, here done for 100 agents. The only difference is, in the top right, we plot the probability k agents are correct for all k larger than half the total number of agents. This gives us an idea of how likely the majority of agents agree.

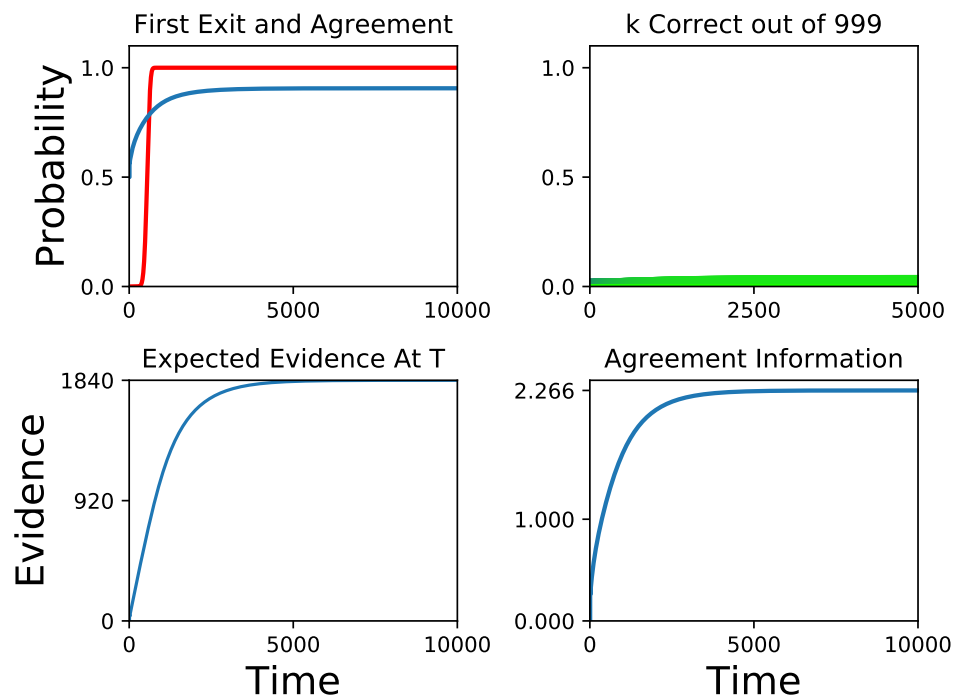


Fig. 6.3.4: The plots are the same as in Figure 6.3.2, here done for 1000 agents. The only difference is, in the top right, we plot the probability k agents are correct for all k larger than half the total number of agents. This gives us an idea of how likely the majority of agents agree.

6.3. LARGE CLIQUES: SIMULATIONS AND ASYMPTOTICS

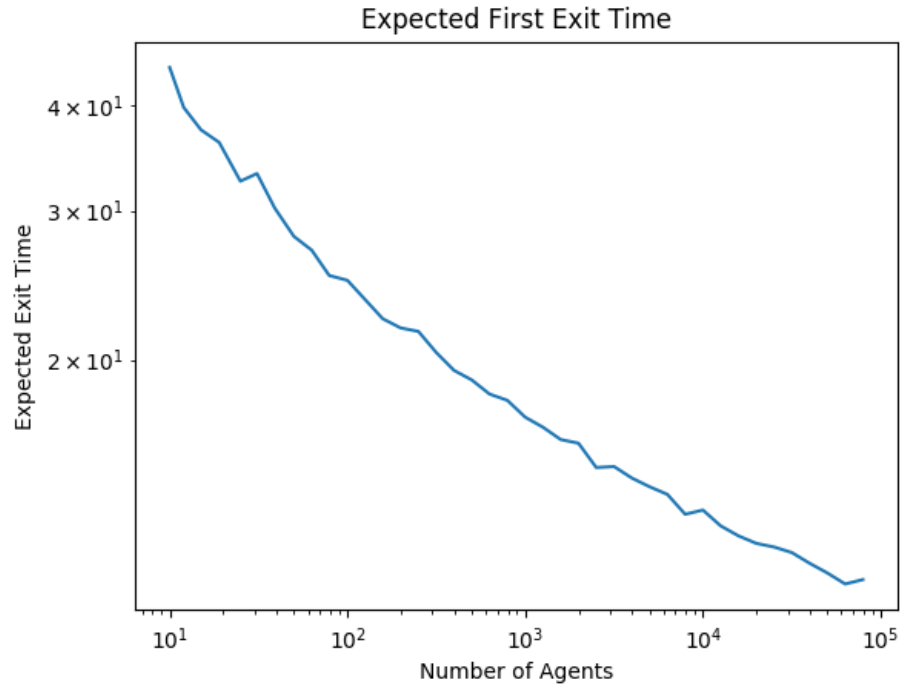


Fig. 6.3.5: The expected first passage time for a given sized clique. Each point was done by averaging the first exit times over 100 simulations.

compute the rate at which the time that the first agent makes a decision goes to zero as the clique size increases in Figure 6.3.5.

Bibliography

- [1] Daron Acemoglu, Munther A Dahleh, Ilan Lobel, and Asuman Ozdaglar. Bayesian learning in social networks. *Rev. Econ. Stud.*, 78(4):1201–1236, 2011.
- [2] Robert J Aumann. Agreeing to disagree. *The annals of statistics*, pages 1236–1239, 1976.
- [3] Bahador Bahrami, Karsten Olsen, Peter E Latham, Andreas Roepstorff, Geraint Rees, and Chris D Frith. Optimally interacting minds. *Science*, 329(5995):1081–1085, 2010.
- [4] Venkatesh Bala and Sanjeev Goyal. Learning from neighbours. *Rev. Econ. Stud.*, 65(3):595–621, 1998.
- [5] Abhijit V Banerjee. A simple model of herd behavior. *Q. J. Econ.*, pages 797–817, 1992.
- [6] Jeffrey M Beck, Wei Ji Ma, Xaq Pitkow, Peter E Latham, and Alexandre Pouget. Not noisy, just wrong: the role of suboptimal inference in behavioral variability. *Neuron*, 74(1):30–39, 2012.
- [7] Krishna Bharat and George A Mihaila. When experts agree: using non-affiliated experts to rank popular topics. In *Proceedings of the 10th international conference on World Wide Web*, pages 597–602. ACM, 2001.
- [8] Manisha Bhardwaj, Samuel Carroll, Wei Ji Ma, and Krešimir Josić. Visual decisions in the presence of measurement and stimulus correlations. *Neural Comput.*, 27(11):2318–2353, 2015.

BIBLIOGRAPHY

- [9] Sushil Bikhchandani, David Hirshleifer, and Ivo Welch. A theory of fads, fashion, custom, and cultural change as informational cascades. *J. Polit. Econ.*, pages 992–1026, 1992.
- [10] Sushil Bikhchandani, David Hirshleifer, and Ivo Welch. Learning from the behavior of others: Conformity, fads, and informational cascades. *J. Econ. Perspect.*, 12(3):151–170, 1998.
- [11] Patrick Billingsley. *Probability and measure*. John Wiley & Sons, 2008.
- [12] Rafal Bogacz, Eric Brown, Jeff Moehlis, Philip Holmes, and Jonathan D Cohen. The physics of optimal decision making: a formal analysis of models of performance in two-alternative forced-choice tasks. *Psychological review*, 113(4):700, 2006.
- [13] Rafal Bogacz, Eric Brown, Jeff Moehlis, Philip Holmes, and Jonathan D. Cohen. The physics of optimal decision making: A formal analysis of models of performance in two-alternative forced-choice tasks. *Psychological Review*, 113(4):700–765, 2006.
- [14] Béla Bollobás. *Random Graphs*. Number 73 in Cambridge Studies in Advanced Mathematics. Cambridge University Press, 2001.
- [15] Bingni W Brunton, Matthew M Botvinick, and Carlos D Brody. Rats and humans can optimally accumulate evidence for decision-making. *Science*, 340(6128):95–98, 2013.
- [16] Reginald J Caginalp and Brent Doiron. Decision dynamics in groups with interacting members. *SIAM Journal on Applied Dynamical Systems*, 16(3):1543–1562, 2017.
- [17] Marquis de Condorcet. Essay on the application of analysis to the probability of majority decisions. *Paris: Imprimerie Royale*, 1785. Reprinted in Condorcet: Selected Writings, Keith Michael Baker, ed, 1976.
- [18] Iain D Couzin, Jens Krause, Nigel R Franks, and Simon A Levin. Effective leadership and decision-making in animal groups on the move. *Nature*, 433(7025):513–516, 2005.
- [19] Morris H DeGroot. Reaching a consensus. *JASA*, 69(345):118–121, 1974.
- [20] Morris H DeGroot. *Optimal statistical decisions*, volume 82. John Wiley & Sons, 2005.

BIBLIOGRAPHY

- [21] Sophie Deneve. Bayesian spiking neurons i: inference. *Neural computation*, 20(1):91–117, 2008.
- [22] Joseph L Doob and Joseph L Doob. *Stochastic processes*, volume 7. Wiley New York, 1953.
- [23] David Easley and Jon Kleinberg. *Networks, crowds, and markets*, volume 1. Cambridge Univ Press, 2010.
- [24] Benjamin Enke and Florian Zimmermann. Correlation neglect in belief formation. *CESifo Working Paper Series No. 4483*, 2013.
- [25] Rafael M Frongillo, Grant Schoenebeck, and Omer Tamuz. Social learning in a changing world. pages 146–157, 2011.
- [26] Douglas Gale and Shachar Kariv. Bayesian learning in social networks. *GEB*, 45(2):329–346, 2003.
- [27] John D. Geanakoplos and Heraklis M. Polemarchakis. We can’t disagree forever. *Journal of Economic Theory*, 28(1):192–200, 1982.
- [28] Joshua I Gold and Michael N Shadlen. The neural basis of decision making. *Annu. Rev. Neurosci.*, 30:535–574, 2007.
- [29] Benjamin Golub and Evan D Sadler. Learning in social networks. 2017.
- [30] S.a b Goyal. *Connections: An introduction to the economics of networks*. 2012.
- [31] Frank Harary. *Graph theory*. 1969. Addison-Wesley, Reading, MA.
- [32] Steven M Kay. *Fundamentals of statistical signal processing, volume I: estimation theory*. Prentice Hall, 1993.
- [33] Margot Kimura and Jeff Moehlis. Group decision-making models for sequential tasks. *SIAM Review*, 54(1):121–138, 2012.
- [34] Ron Kohavi and Foster Provost. Glossary of terms. *Machine Learning*, 30(2-3):271–274, 1998.
- [35] J Komlós. On the determinant of random matrices. *Stud. Sci. Math. Hung.*, 3(4):387–399, 1968.
- [36] Gilat Levy and Ronny Razin. Correlation neglect, voting behavior, and information aggregation. *The American Economic Review*, 105(4):1634–1645, 2015.

BIBLIOGRAPHY

- [37] Wei Ji Ma, Jeffrey M Beck, Peter E Latham, and Alexandre Pouget. Bayesian inference with probabilistic population codes. *Nature neuroscience*, 9(11):1432–1438, 2006.
- [38] Tyler McMillen and Philip Holmes. The dynamics of choice among multiple alternatives. *Journal of Mathematical Psychology*, 50(1):30–57, 2006.
- [39] Rubén Moreno-Bote, Jeffrey Beck, Ingmar Kanitscheider, Xaq Pitkow, Peter Latham, and Alexandre Pouget. Information-limiting correlations. *Nat. Neurosci.*, 17(10):1410–1417, 2014.
- [40] Elchanan Mossel, Allan Sly, and Omer Tamuz. Asymptotic learning on bayesian social networks. *Probab. Theory Related Fields*, 158(1-2):127–157, 2014.
- [41] Elchanan Mossel and Omer Tamuz. Opinion exchange dynamics. *arXiv preprint arXiv:1401.4770*, 2014.
- [42] Elchanan Mossel and Omer Tamuz. Opinion exchange dynamics. *arXiv preprint arXiv:1401.4770*, 2014.
- [43] Elchanan Mossel and Omer Tamuz. Efficient bayesian learning in social networks with gaussian estimators. In *Communication, Control, and Computing (Allerton), 2016 54th Annual Allerton Conference on* (pp. 425-432). IEEE., 2016.
- [44] Manuel Mueller-Frank. A general framework for rational learning in social networks. *Theoretical Economics*, 8(1):1–40, 2013.
- [45] Manuel Mueller-Frank. Does one bayesian make a difference? *Journal of Economic Theory*, 154:423–452, 2014.
- [46] Pietro Ortoleva and Erik Snowberg. Overconfidence in political behavior. *The American Economic Review*, 105(2):504–535, 2015.
- [47] Yuval Peres. Brownian Motion. *Colloids and Surfaces A Physicochemical and Engineering Aspects*, 106:230601, 2008.
- [48] Ioannis Poulakakis, Luca Scardovi, and Naomi Ehrich Leonard. Node classification in networks of stochastic evidence accumulators. *arXiv preprint arXiv:1210.4235*, 2012.
- [49] Jeffrey S Rosenthal. *A first look at rigorous probability theory*. World Scientific Publishing Co Inc, 2006.

BIBLIOGRAPHY

- [50] Lones Smith and Peter Sorensen. Pathological Outcomes of Observational Learning. *Econometrica*, 68(2):371–398, 2000.
- [51] Vaibhav Srivastava and Naomi Ehrich Leonard. Collective decision-making in ideal networks: The speed-accuracy tradeoff. *IEEE Transactions on Control of Network Systems*, 1(1):121–132, 2014.
- [52] Vaibhav Srivastava and Naomi Ehrich Leonard. Collective decision-making in ideal networks: The speed-accuracy tradeoff. *IEEE Transactions on Control of Network Systems*, 1(1):121–132, 2014.
- [53] Howard M Taylor and Samuel Karlin. *An introduction to stochastic modeling*. Academic press, 2014.
- [54] Alan Veliz-Cuba, Zachary P Kilpatrick, and Kresimir Josic. Stochastic models of evidence accumulation in changing environments. *SIAM Review*, 58(2):264–289, 2016.
- [55] Abraham Wald and Jacob Wolfowitz. Optimum character of the sequential probability ratio test. *The Annals of Mathematical Statistics*, pages 326–339, 1948.
- [56] Stanley Wasserman and Katherine Faust. *Social network analysis: Methods and applications*, volume 8. Cambridge university press, 1994.
- [57] Ivo Welch. Sequential sales, learning, and cascades. *J. Finance*, 47(2):695–732, 1992.
- [58] David Williams. *Probability with martingales*. Cambridge university press, 1991.