

An Informational Rationale for Viewpoint Neutrality in Education

Georgy Egorov*

Konstantin Sonin[†]

June 10, 2026

Abstract

Consider a society that faces uncertainty about a payoff-relevant state and wants to train students to make correct decisions. In educational institutions, students learn from their teachers, but they also get outside information, and later learn from peers. We show that privately Bayesian actions need not be optimal inputs into social learning: when students' actions reflect teacher-side information that is correlated across peers, observing many such actions can give this information excessive social weight. A social planner may therefore optimally reduce the precision of instruction, inducing students to rely more on outside information before their actions become signals for others, and students can end up better informed despite learning less from their teachers. The case for a precision cap is stronger when peer interaction is homophilous, because same-teacher information is more likely to survive aggregation, and weaker when outside information is itself systematically distorted. This provides an informational rationale for viewpoint neutrality as an institutional policy: it limits the social overrepresentation of correlated teacher-side information when students mostly learn from peers exposed to similar sources.

Keywords: viewpoint neutrality, social learning, information aggregation, homophily, echo chambers, education policy.

JEL Codes: D83, D85, I21, I28.

*Northwestern University and NBER; g-egorov@kellogg.northwestern.edu

[†]University of Chicago; ksonin@uchicago.edu

1 Introduction

In educational environments, teachers pass knowledge to students: they select material, organize arguments, and communicate their views. At the same time, students do not learn only from teachers. They learn from books, observation, casual empiricism, their own experience, and from their peers. Empirically, both teachers and peers matter. Teachers have large and persistent effects on student outcomes (Chetty, Friedman and Rockoff, 2014), and peer composition affects achievement in schools and universities (Sacerdote, 2001; Zimmerman, 2003; Carrell, Sacerdote and West, 2013). Classical accounts of education also emphasize inquiry and deliberation rather than the mere transmission of conclusions (Dewey, 1916; Hess and McAvoy, 2015). These observations raise a question usually framed in ethical, legal, or pedagogical terms, but which also has an informational dimension: how precisely should teachers transmit their own views?

A natural answer is that teachers should transmit their views as precisely as possible. If teachers are informed, and if students use information rationally, then more precise instruction seems to make students better informed. This logic is correct when a student acts in isolation, but it is incomplete when students later learn from one another. In that case, teachers and peers are not independent sources of information. What a student learns from her peers depends partly on what those peers previously learned from their own teachers. If different students inherit similar mistakes from their teachers, then learning from them does not provide as much new information as it appears to provide. Information that is useful for one student acting alone can therefore become redundant when it is later reflected in peers' words or actions. More precise transmission from teachers to students may then improve individual learning while making subsequent social learning less informative.

In this paper, we show that an institutionalized policy that makes teacher-student communication coarser can improve students' ultimate learning. The policy does not make

teachers less informed, nor does it require students to ignore them. Instead, it limits how precisely students inherit their teachers’ posterior views before their own decisions become inputs into peer learning. We interpret this policy as a formalization of *viewpoint neutrality*. In many educational debates, viewpoint neutrality is understood as a restriction on teachers using classroom authority to transmit their own conclusions on contested questions. Our model captures the informational content of such a restriction: as teachers cover different ideas over the course of instruction, students may learn from them without learning their posterior views precisely. The upside of such a policy is that, by making students less likely to reproduce the teacher’s conclusions with high precision, it reduces the extent to which teacher-side errors are later repeated by students and mistaken for independent confirmation. The rationale is therefore not that teachers lack expertise, or that their views are inherently biased. Indeed, the policy does not single out particular views: it governs the precision of transmission, not the content. Even valuable teacher information can be socially overrepresented when it circulates through peer learning.

We build a model in which society is uncertain about a payoff-relevant state, and students learn about this state from several sources. Teachers observe informative but noisy signals and form posterior views that students learn from, subject to a possible limit on the precision of transmission. Students also receive outside signals, coming from books, experimental observation, life experience, or other non-teacher sources, and combine teacher-side and outside information as Bayesians when choosing their actions. They then observe the actions of other students and choose their actions again. We show that, to increase the precision of students’ ultimate learning, it may be optimal to cap the precision of information transmitted from teachers to students. Such a cap makes students rely less on teacher-side information before their actions are observed by peers, reducing the extent to which correlated teacher-side information is carried into social learning. At the same time, it makes students’ actions reflect relatively more of their outside information, which is less correlated across students

and therefore more useful for aggregation. Indeed, under the optimal cap, students can end up better informed than their teachers after peer learning, even if the cap makes their first actions less precise.

We show that a tighter precision cap is optimal when peer interaction is more homophilous, when students observe more peer actions, and when the planner puts greater weight on decisions made after peer learning. Instructor-specific errors strengthen this force, because they are preserved among students of the same teacher rather than averaged out. When outside signals are relatively precise and idiosyncratic, a tighter cap improves aggregation by making students' actions reflect more independent information. In contrast, when outside signals are noisy, the case for a cap weakens, and when outside information is sufficiently distorted in a way common across students, the planner may instead prefer students to rely more on instructors. In a dynamic extension, the need for a precision cap eventually disappears as reliable knowledge accumulates and teacher-side errors shrink. The path, however, need not be monotone: viewpoint neutrality may become more important before it becomes obsolete.

These results suggest why the value of viewpoint neutrality may differ across fields. In engineering, mathematics, and the more settled parts of the natural sciences, central claims have typically been tested by proof, experiment, replication, and technical application, and disagreement among qualified instructors is correspondingly limited. Much of what is taught is therefore a summary of accumulated knowledge rather than personal interpretation. Making students rely less on it may push them toward noisier independent learning, where they are likely to rediscover mistakes that earlier generations have already corrected. The cost of a precision cap is therefore high, while the benefit from reducing redundant teacher-side information is limited. By contrast, in parts of the social sciences, law, public policy, and ethics, evidence is often less conclusive, interpretation is more central, and disagreement among qualified instructors is more persistent. Teachers' views are therefore more likely

to reflect contestable frameworks or broader worldviews. If students then mostly learn from peers exposed to similar frameworks, social learning preserves rather than diversifies the teacher-driven component of beliefs. In such environments, the informational case for viewpoint neutrality is correspondingly stronger.

Our paper contributes to several literatures. First, this work connects to studies of academic freedom, instruction on controversial topics, and viewpoint neutrality in particular. The classical “marketplace of ideas” tradition, associated with Mill’s defense of open discussion, argues that exposure to competing viewpoints helps separate truth from error (Mill, 1859). In modern higher education, the Kalven Report gives a canonical statement of institutional neutrality, arguing that universities should generally refrain from taking collective positions on political and social questions in order to preserve individual inquiry (Kalven Committee, 1967). Recent work revisits neutrality from legal, institutional, and academic-freedom perspectives (Post, 2024; American Association of University Professors, 2025; Soucek, 2026), with legal scholarship emphasizing the tension between educational authority and open inquiry (Annest, 2002; Salzman, 2022). The pedagogical literature also asks whether and when teachers should disclose their own views on controversial issues, and how classroom discussion should handle disagreement (Dewey, 1916; Kelly, 1986; Hess and McAvoy, 2015; Dunn, Sondel and Baggett, 2019). Our contribution is to provide a formal informational rationale for neutrality: even when teacher information is valuable for individual learning, it may become socially overrepresented when students later learn from peers with correlated teacher-side information.

Second, our work contributes to the literature on social learning, where individually rational behavior need not aggregate dispersed information efficiently. This arises in sequential learning models, where public histories can crowd out private signals (Banerjee, 1992; Bikhchandani, Hirshleifer and Welch, 1992) (see Bikhchandani, Hirshleifer, Tamuz

and Welch, 2024 for a recent review), and in simultaneous-action settings such as strategic voting (Feddersen and Pesendorfer, 1998). A related literature studies Bayesian learning in networks, including learning from neighbors’ information and actions (Acemoglu, Dahleh, Lobel and Ozdaglar, 2011), learning dynamics and experimentation in large networks (Board and Meyer-ter Vehn, 2021, 2024), and intergenerational learning about a changing state (Dasaratha, Golub and Hak, 2023). Our planner shapes the correlation structure of agents’ private information through the precision of teacher-student communication.

Third, our updating environment connects to DeGroot-style belief exchange (DeGroot, 1974; Golub and Jackson, 2010). In our main model, Bayesian updating from peer actions takes an averaging form, but this is not an assumed behavioral rule: students are Bayesian, and the posterior variance reflects the covariance among the observed actions. In richer networks Bayesian inference can be computationally hard (Hazla, Jadbabaie, Mossel and Rahimian, 2021); our Gaussian structure keeps the analysis tractable. The mechanism is also related to correlation neglect but does not rely on it: Eyster and Rabin (2010) and Levy and Razin (2015) study agents who overweight correlated information, whereas here students understand the information structure, and the inefficiency arises because actions compress private information and carry correlated teacher-side components.

Fourth, the paper relates to work on the social value and design of information. In Morris and Shin (2002), more precise public information can lower welfare because public signals receive disproportionate weight when agents coordinate. A related effect appears here: the teacher-side component of students’ information is common across students of the same instructor, so it can receive large social weight when their actions are later aggregated. The mechanism is different: there, the overweighting comes from a coordination motive; here, from aggregating correlated actions. In Bayesian persuasion and information design, a sender chooses an information structure to influence receivers’ beliefs and actions (Kamenica and Gentzkow, 2011; Bergemann and Morris, 2016). Our planner also chooses an information

structure, but the motive is corrective rather than persuasive: the goal is not to move beliefs toward a preferred state, but to make later social learning more informative. Related work studies influence and manipulation in social networks (Mostagir, Ozdaglar and Siderius, 2022); unlike these papers, the planner here shares students’ objective of accuracy and changes information transmission to correct a downstream aggregation problem.

Fifth, our paper connects to studies of homophily, peer effects, and education. Social ties are strongly homophilous (McPherson, Smith-Lovin and Cook, 2001), and economic models of friendship formation show how same-type bias generates segregated networks (Currarini, Jackson and Pin, 2009). Ideological segregation in news consumption and political information environments is documented by Gentzkow and Shapiro (2011) and Allcott and Gentzkow (2017), while Sunstein (2002, 2009) emphasizes that “echo chambers” can amplify prior views. In our model, homophily affects the aggregation problem by increasing teacher-side information redundancy. There is also a large empirical literature on peer effects in education: apart from the evidence cited earlier, Bursztyn and Jensen (2015) show that observability of effort to peers can discourage educational investment, consistent with mechanism studied by Austen-Smith and Fryer (2005), while Bursztyn, Egorov and Jensen (2019) distinguish peer cultures that reward ability from those that stigmatize effort. More recently, Tincani (forthcoming) uses variation from the Chilean earthquake to study peer effects and rank concerns in the classroom. In this work, we do not model classroom peer effects, but rather how students aggregate information that may be correlated across peers, because they share a teacher or because teachers themselves share a common bias.

The rest of the paper is organized as follows. Section 2 presents the theoretical framework. Section 3 analyzes the baseline model and characterizes the optimal precision cap. Section 4 considers extensions, including policies that shift students toward teacher-side information (dogmatism and infallibility), partial echo chambers, and dynamics. Section 5 concludes.

The Appendix contains all proofs and additional examples.

2 Theoretical Framework

In this section we build a model of individual and social learning.

2.1 Environment

The society is populated by a unit continuum of instructors $\mathbb{I}$, each of whom will teach a unit continuum of students $\mathbb{S}_i, i \in \mathbb{I}$. There is a state variable θ , and its distribution is common knowledge: $\theta \sim \mathcal{N}(\mu, \sigma_\theta^2)$; we denote its precision by $p_\theta = \sigma_\theta^{-2}$. Each instructor i observes a noisy signal of θ ,

$$s_i = \theta + b + \varepsilon_i, \quad (1)$$

where $b \sim \mathcal{N}(0, \sigma_b^2)$ is a common shock affecting all instructors and $\varepsilon_i \sim \mathcal{N}(0, \sigma_\varepsilon^2)$ is an idiosyncratic instructor error. The common shock b captures systematic bias in instructors' information environment: pooling the information of all instructors would average out the idiosyncratic errors but not this common shock. All error variables are independent unless explicitly stated otherwise.

Instructors are Bayesian and know the distribution of b , though not its realization. Thus, instructor i 's posterior distribution of θ conditional on signal s_i is

$$\theta \mid s_i \sim \mathcal{N}(a_i, \sigma_a^2), \quad (2)$$

where

$$a_i = \frac{\sigma_\theta^2}{\sigma_\theta^2 + \sigma_b^2 + \sigma_\varepsilon^2} s_i + \left(1 - \frac{\sigma_\theta^2}{\sigma_\theta^2 + \sigma_b^2 + \sigma_\varepsilon^2}\right) \mu, \quad \sigma_a^2 = \frac{\sigma_\theta^2(\sigma_b^2 + \sigma_\varepsilon^2)}{\sigma_\theta^2 + \sigma_b^2 + \sigma_\varepsilon^2}.$$

2.2 Students and education

Each instructor i teaches a continuum of students $s \in \mathbb{S}_i$, and each such student s observes

$$y_{is} = a_i + \tau_{is}, \quad (3)$$

where $\tau_{is} \sim \mathcal{N}(0, \sigma_\tau^2)$ captures the residual imprecision in the student's observation of the instructor's posterior mean. We interpret this imprecision as induced by an institutional cap on how precisely instructors may transmit their posterior views. In what follows, we denote $T = \sigma_\tau^2$; a larger T corresponds to a tighter precision cap.¹

Apart from learning from the teacher, each student obtains an outside signal about the state of the world, perhaps by reading books or empirical observation:

$$z_{is} = \theta + c + \eta_{is}, \quad (4)$$

where $c \sim \mathcal{N}(0, \sigma_c^2)$ captures systematic distortion of the channel, while $\eta_{is} \sim \mathcal{N}(0, \sigma_\eta^2)$ is idiosyncratic noise. Equipped with signals y_{is} and z_{is} , the student forms the following posterior belief about θ :

$$\theta \mid y_{is}, z_{is} \sim \mathcal{N}(l_{is}, \sigma_l^2),$$

where σ_l^2 is a function of variances defined above, and l_{is} is a linear combination of μ, y_{is}, z_{is} .

At the end of this individual learning stage, each student chooses individual action u_{is} to minimize the quadratic error $\mathbb{E}[(u_{is} - \theta)^2]$; given the normal distributions in the problem, this action would be the student's posterior mean of θ . Practically, this action may correspond to passing exams or doing an internship or residency.

¹ Signal y_{is} in (3) may be interpreted as an aggregate of a large number K of signals $y_{isk} = a_i + \tau_{isk}$, $\tau_{isk} \sim \mathcal{N}(0, KT)$, transmitted to the student over the course of education; with this interpretation, $T = 0$ corresponds to communicating true posterior a_i every time, while viewpoint neutrality requires setting $T > 0$ so that signals from the instructor will span a range of possible values, allowing the student to learn from the teacher without getting a precise signal of the instructor's posterior a_i .

2.3 Social learning

After choosing their individual actions, each student observes the actions of $n > 0$ other students. We assume that with probability q the student is stuck in an echo chamber and observes actions of students of the same instructor, while with complementary probability $1 - q$ the n students are randomly taken from the entire society. These students (whose actions are observed by student s) will be denoted r_{sk} for $k = 1, \dots, n$, and the instructors of these students are i_{sk} (if the student s of instructor i is in an echo chamber, all $i_{sk} \equiv i$). Student s knows the identities of students and instructors, and in particular whether they coincide with her own, and uses this information as a Bayesian.²

In this notation, student s of instructor i observes n actions $u_{i_{sk}r_{sk}}$ from other students. The student then makes the following posterior update about θ :

$$\theta \mid l_{is}, \{u_{i_{sk}r_{sk}}\}_{k=1,\dots,n} \sim \mathcal{N}(m_{is}, \sigma_m^2). \quad (5)$$

After that, each student takes another action v_{is} with the goal of minimizing the quadratic error $\mathbb{E}[(v_{is} - \theta)^2]$. Because of symmetry of students (both in cases they are in an echo chamber and outside an echo chamber), this action will equal

$$v_{is} = m_{is} = \frac{l_{is} + \sum_{k=1}^n u_{i_{sk}r_{sk}}}{n + 1}. \quad (6)$$

Two comments are in order about the updating rule (5). First, it is important that students aggregate actions and not signals, and we believe it is natural. Observing the whole history of signals that other students got and that led them to choose their actions is hardly

²This structure is analytically simplest. If students did not know whether they are in an echo chamber, they would make inference about that from the empirical variance of the actions of other students that they observe, which is both analytically difficult and hardly realistic. An alternative setting, where students observe several actions of students of the same instructor and several actions of random students is considered in Section 4.

realistic; at best this is plausible for a small subset of close friends. For the model, however, this is a consequential assumption: if the students observed the other students’ signals, they would be able to perfectly deduce the effect of having a common instructor and learn from other students’ outside signals. In this case, echo chambers would still have an adverse effect on learning because students would not learn the signals from instructors other than their own, but viewpoint neutrality would not arise in the optimum.

Second, we could assume an alternative where a student observes others’ actions, but uses the original signals y_{is} and z_{is} rather than her own prior action to make inference about θ . In this case, the student would use this knowledge of what she learned from the teacher and what from an outside signal to help interpret other students’ actions; nevertheless, viewpoint neutrality can be optimal, as we show in Example A1 in the Appendix. The cost of going this route would be loss of tractability; furthermore, this way of social learning is likely less realistic than our baseline formulation.³

Figure 1 summarizes signals received and actions taken in the model.

2.4 Planner’s problem

The planner chooses the tightness of the precision cap, parameterized by $T = \sigma_\tau^2 \geq 0$, to minimize

$$A\mathbb{E}[(u_{is} - \theta)^2] + B\mathbb{E}[(v_{is} - \theta)^2]. \quad (7)$$

³For example, if students get numerous bits and pieces of information from the teacher (as described in Footnote 1) and from other sources, but then process it and remember the aggregate, then to condition on y_{is} and z_{is} the student would need to retain not only her final judgment about θ , but also a separable record of which parts of that judgment came from the instructor and which came from outside evidence. In contrast, (6) makes it clear that our preferred formulation (5) is equivalent to DeGroot averaging, which is simple and realistic (DeGroot, 1974; DeMarzo, Vayanos and Zwiebel, 2003; Golub and Jackson, 2010; Chandrasekhar, Larreguy and Xandri, 2020). The same rule arises if one assumes a two-selves structure in which a ‘student self’ studies and forms a judgment, and a ‘practitioner self’ subsequently uses that judgment in practice and uses it as a starting point when learning from peers.

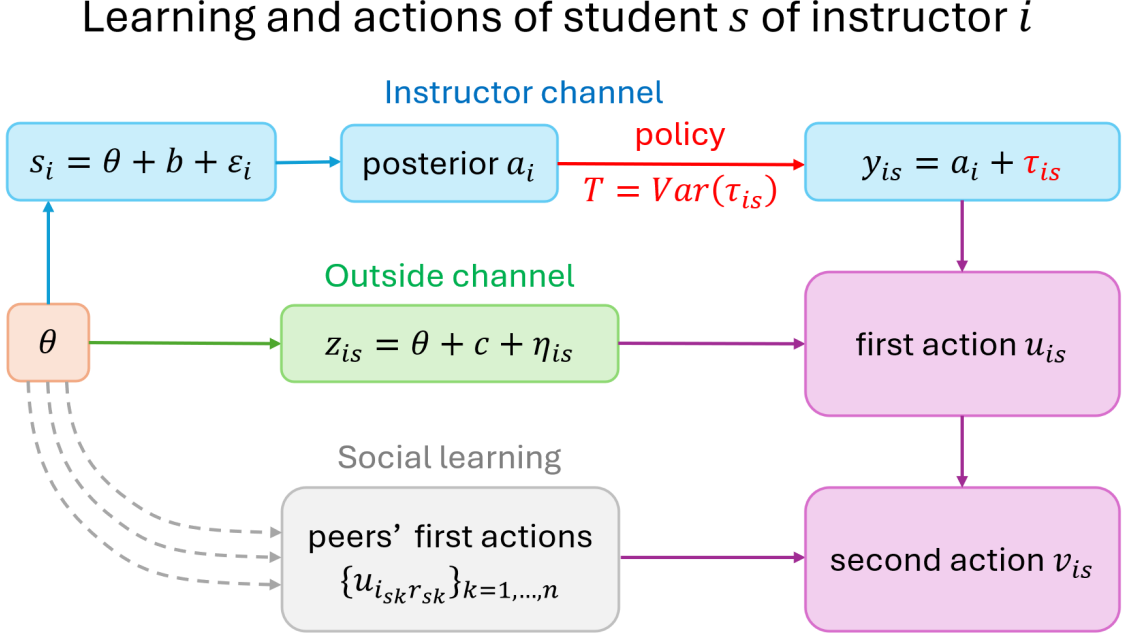


Figure 1: Learning and actions of student s of instructor i

In other words, the planner wants students to make actions that match the state of the world to the extent possible both before and after learning from their peers, and puts weights A and B on the actions before and after, respectively. This minimization problem seeks to capture the planner's interest in students choosing the right actions both after their individual training and after aggregating dispersed social information. For example, a medical student may start assisting doctors in surgeries immediately after concluding their studies, but then, after learning more and seeing others in action, they have a full career of their own practice. It is natural for weights A and B to be different and we maintain a flexible assumption. If the planner put full weight on the ultimate information that students have by the end of their careers (for example, when they start teaching a new generation), this would correspond to $A = 0$; we will see that this will make viewpoint neutrality even more likely to be optimal.

We assume that the only instrument the planner has is the precision cap on teacher-student communication. We parameterize this cap by T , the variance of the residual im-

precision with which students observe the instructor's posterior mean. Conceptually, $T = 0$ corresponds to no precision cap: the student observes a_i perfectly. By contrast, $T > 0$ corresponds to a binding cap on instructional precision: the student observes only $a_i + \tau_{is}$, and therefore cannot perfectly recover the instructor's posterior mean. Our results would not change conceptually if teaching were inherently imprecise, so that the lower bound were some $T_{\min} > 0$. For simplicity, we set $T_{\min} = 0$.

2.5 Timing and equilibrium concept

The timing is as follows.

1. The planner chooses the precision cap, parameterized by $T \geq 0$, which is then observed by all players.
2. Nature draws the state θ .
3. Each instructor i observes s_i given by (1) and forms a Bayesian posterior.
4. Each student $s \in \mathbb{S}_i$ assigned to instructor i observes signal y_{is} given by (3) and signal z_{is} given by (4).
5. Each student chooses action u_{is} .
6. Each student observes actions $\{u_{i_{sk}r_{sk}}\}$, $k = 1, \dots, n$, of n other students and aggregates this information according to (5).
7. Each student chooses action v_{is} .

An equilibrium is defined as a profile of strategies $(T, \{u_{is}(T, y_{is}, z_{is}), \{v_{is}(T, l_{is}, \{u_{i_{sk}r_{sk}}\})\})$ where T minimizes the planner's loss (7), u_{is} minimizes $\mathbb{E}[(u_{is} - \theta)^2]$ given its arguments T, y_{is}, z_{is} , and v_{is} minimizes $\mathbb{E}[(v_{is} - \theta)^2]$ given its arguments $T, l_{is}, \{u_{i_{sk}r_{sk}}\}$.

3 Analysis

We now solve the model. We will focus on the case in which $p_\theta = 0$, or equivalently $\sigma_\theta^2 \rightarrow \infty$. This assumption allows us to isolate the social-learning mechanism that is central to the paper without keeping track of the weight that teachers and then students put on the unconditional prior μ . We also find the improper-prior case natural, with individuals starting agnostic about the state of the world. Since for $p_\theta \geq 0$, all posterior distributions are normal and their means and variances are continuous in p_θ , all comparative statics results below hold for $p_\theta < \bar{p}$ with some $\bar{p} > 0$.

If $p_\theta = 0$, the posterior of instructor i after receiving signal s_i , (2) becomes

$$\theta \mid s_i \sim \mathcal{N}(s_i, \sigma_b^2 + \sigma_\varepsilon^2). \quad (8)$$

Consider student s who received signal y_{is} from instructor i and observes outside signal z_{is} . Since y_{is} is an unbiased (thanks to $p_\theta = 0$) signal of θ with variance

$$\sigma_y^2(T) \equiv \sigma_b^2 + \sigma_\varepsilon^2 + T = \sigma_b^2 + \sigma_\varepsilon^2 + \sigma_\tau^2$$

and z_{is} is an unbiased signal of θ with variance

$$\sigma_z^2 = \sigma_c^2 + \sigma_\eta^2,$$

then the posterior that the student forms is

$$\theta \mid y_{is}, z_{is} \sim \mathcal{N}(l_{is}, \sigma_l^2),$$

where

$$l_{is} = \frac{\sigma_z^2 y_{is} + \sigma_y^2(T) z_{is}}{\sigma_y^2(T) + \sigma_z^2}, \quad \sigma_l^2 = \frac{\sigma_y^2(T) \sigma_z^2}{\sigma_y^2(T) + \sigma_z^2}.$$

The first action of each student is therefore $u_{is} = l_{is}$ and the variance of each student's first action around θ is σ_l^2 .

Now consider social learning. First, consider student s who is not in an echo chamber. If so, she observes n actions $u_{i_{sk}r_{sk}}$ of other students. These students are equally informed, their actions have the same variance σ_l^2 , and by symmetry the covariances of actions of each pair of different students is also the same. Hence, each student's posterior mean (conditional on her own action and other students' actions) equals

$$m_{is} = \frac{u_{is} + \sum_{k=1}^n u_{i_{sk}r_{sk}}}{n+1}$$

These $n+1$ students are almost surely taught by different instructors, and thus the covariance between their first action errors is

$$\frac{(\sigma_z^2)^2 \sigma_b^2 + (\sigma_y^2(T))^2 \sigma_c^2}{(\sigma_y^2(T) + \sigma_z^2)^2}.$$

These covariances are present due to the systematic biases that all instructors share (b with variance σ_b^2) and all outside signals share (c with variance σ_c^2). Taking these covariances into account, we find that the variance of $m_{is} - \theta$ for a student not in an echo chamber equals

$$\sigma_{m,N}^2 = \frac{1}{n+1} \frac{\sigma_y^2(T) \sigma_z^2}{\sigma_y^2(T) + \sigma_z^2} + \frac{n}{n+1} \frac{(\sigma_z^2)^2 \sigma_b^2 + (\sigma_y^2(T))^2 \sigma_c^2}{(\sigma_y^2(T) + \sigma_z^2)^2}.$$

Since the student's posterior is normal with mean m_{is} and variance $\sigma_{m,N}^2$, the student takes second action $v_{is} = m_{is}$.

Now consider student s who found herself in an echo chamber. In this case, she observes

n actions $u_{ir_{sk}}$ of other students of the same instructor i . Again, these students are equally informed, and thus the student's posterior mean equals

$$m_{is} = \frac{u_{is} + \sum_{k=1}^n u_{ir_{sk}}}{n+1}.$$

The covariance between the first action errors of these students is

$$\frac{(\sigma_z^2)^2(\sigma_b^2 + \sigma_\varepsilon^2) + (\sigma_y^2(T))^2\sigma_c^2}{(\sigma_y^2(T) + \sigma_z^2)^2}.$$

The difference with the earlier expression is the term involving σ_ε^2 : students with the same instructor share not only the common shock b , but also the instructor-specific error ε_i . Taking these covariances into account, the variance of $m_{is} - \theta$ for a student in an echo chamber equals

$$\sigma_{m,E}^2 = \frac{1}{n+1} \frac{\sigma_y^2(T)\sigma_z^2}{\sigma_y^2(T) + \sigma_z^2} + \frac{n}{n+1} \frac{(\sigma_z^2)^2(\sigma_b^2 + \sigma_\varepsilon^2) + (\sigma_y^2(T))^2\sigma_c^2}{(\sigma_y^2(T) + \sigma_z^2)^2}.$$

Since the share of students in echo chambers is q , the planner's objective can be written as the following loss function:

$$\begin{aligned} \mathcal{L}(T) \equiv A\sigma_l^2 + B(q\sigma_{m,E}^2 + (1-q)\sigma_{m,N}^2) &= \left(A + \frac{B}{n+1}\right) \frac{\sigma_y^2(T)\sigma_z^2}{\sigma_y^2(T) + \sigma_z^2} \\ &+ \frac{nB}{n+1} \frac{(\sigma_z^2)^2(\sigma_b^2 + q\sigma_\varepsilon^2) + (\sigma_y^2(T))^2\sigma_c^2}{(\sigma_y^2(T) + \sigma_z^2)^2}. \end{aligned} \quad (9)$$

This expression makes the planner's trade-off clear. The first term corresponds to the direct loss from imprecision of a student's first action and, with weight $1/(n+1)$, second action. This term is unambiguously increasing in T , because a tighter precision cap makes the instructor-side signal less precise and raises the planner's quadratic loss from individual learning. The second term corresponds to the correlated components in a student's learning

of n other students' actions; it may be increasing or decreasing in T depending on the relative importance of learning from instructor and from the world. As a result the planner's loss from the first action is monotonically increasing in T , whereas the loss from the second action and the total loss may be nonmonotone, as illustrated in Figure 2. Formally, we prove the following Proposition.

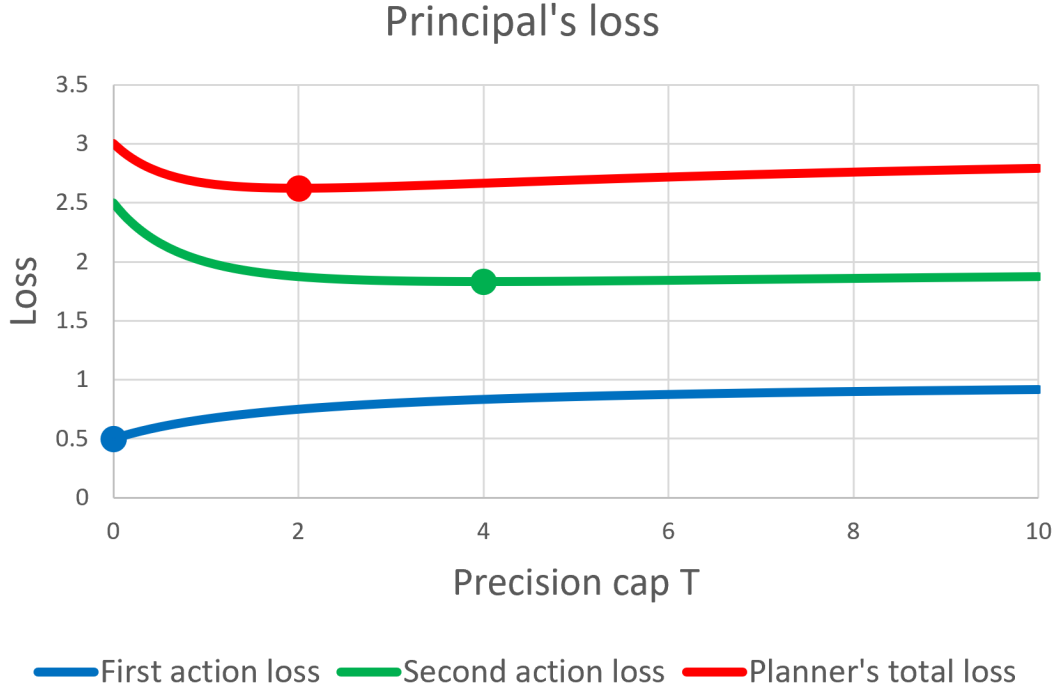


Figure 2: Planner's losses from first and second actions and total loss for $A = 1$, $B = 8$, $\sigma_b^2 = \sigma_c^2 = 0$, $\sigma_\varepsilon^2 = \sigma_\eta^2 = 1$, $n = 3$, $q = 1$. Circles denote minima of each graph.

Proposition 1. *Suppose $A > 0$, $B > 0$, and $\sigma_\varepsilon^2 > 0$. The planner's optimal solution T^* is unique. The optimal T^* is weakly increasing in q , n , and B/A (strictly increasing if $T^* > 0$). If σ_b^2 and σ_c^2 are sufficiently small (and in particular in the no-common-shock case $\sigma_b^2 = \sigma_c^2 = 0$), T^* is weakly increasing in σ_ε^2 and weakly decreasing in σ_η^2 , with strict inequalities whenever $T^* > 0$.*

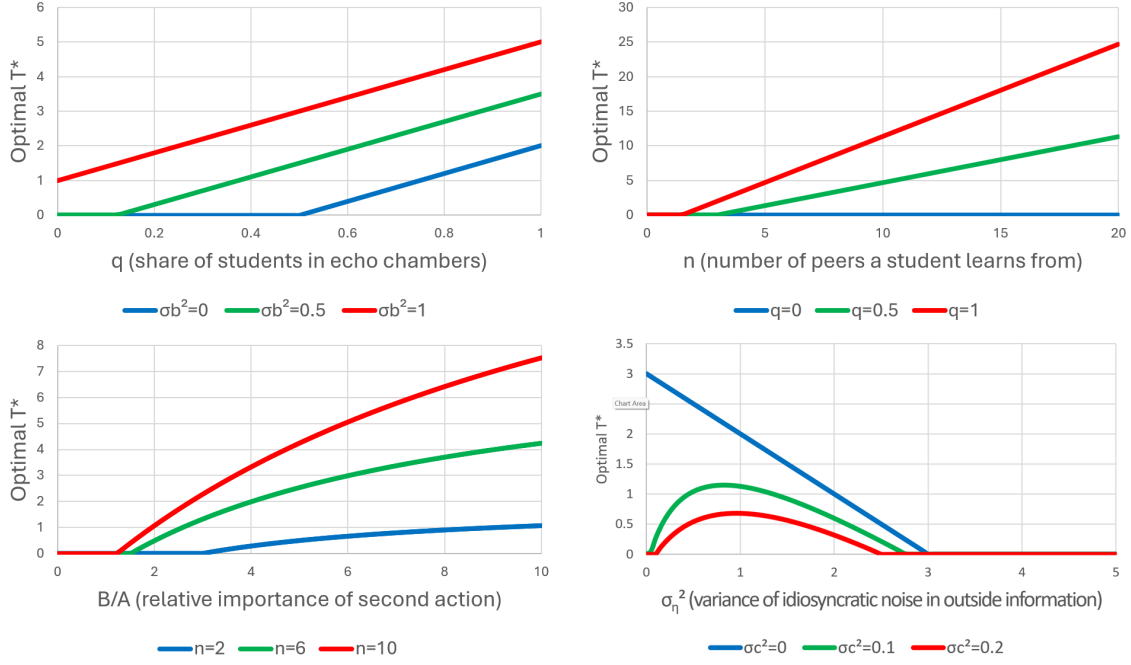


Figure 3: Comparative statics results. Variables, except the ones varied in each panel, are set to $A = 1$, $B = 8$, $\sigma_b^2 = \sigma_c^2 = 0$, $\sigma_\epsilon^2 = \sigma_\eta^2 = 1$, $n = 3$, $q = 1$.

Proof of Proposition 1. See Appendix.

The comparative statics are intuitive, and Figure 3 illustrates them. The planner chooses a higher T^* , and is more likely to choose $T^* > 0$, when the social-learning benefits of weakening the teaching channel are large relative to the individual-learning costs. First, T^* is increasing in q , the probability that a student is in an echo chamber. When echo chambers are more likely, more students learn from peers who share the same instructor-specific error, so reducing students' reliance on instructor-side information becomes more valuable. Second, T^* is increasing in n , the number of peer actions observed. When students observe more peer actions, correlated mistakes are repeated more often in the social-learning process, making redundancy more costly. Third, T^* is increasing in B/A : a tighter precision cap is more valuable when the planner puts greater weight on the second action, after social learning, relative to the first action.

Finally, when the common shocks are sufficiently small, T^* is increasing in σ_ϵ^2 and de-

creasing in σ_η^2 . A larger σ_ε^2 means that students taught by the same instructor have more strongly correlated first actions, because they inherit a larger instructor-specific error. The planner therefore has a stronger incentive to tighten the precision cap, that is, to increase T , making students rely less on the instructor-side signal, thereby reducing redundancy in social learning. By contrast, a larger σ_η^2 makes students' outside signals less precise. Weakening the teaching channel then becomes more costly, because students need to rely more heavily on their instructors to form accurate first actions. Thus, the planner has less reason to increase T when outside learning is noisier.

The analysis above makes clear that echo chambers hurt the precision of students' information after social learning, and Proposition 1 suggests that more echo chambers necessitate a higher T^* at the optimum. But would eliminating echo chambers imply that optimal $T^* = 0$? We have the following result.

Corollary 2. *Suppose $q = 0$. Then for σ_b^2 sufficiently low, the optimal $T^* = 0$. For σ_b^2 sufficiently high, we will have $T^* > 0$, provided that*

$$\frac{\sigma_c^2}{\sigma_\eta^2} < \frac{(2n-1)B - (n+1)A}{B + (n+1)A}. \quad (10)$$

Intuitively, if $\sigma_b^2 = 0$, then eliminating echo chambers removes the only source of teacher-side redundancy in social learning. Students no longer observe peers who share the same instructor-specific error ε_i , and there is no common instructor-side component b shared across students with different instructors. In this case, weakening teacher-student communication only makes individual learning worse and creates no offsetting benefit for social learning, so the optimal precision cap is non-binding: $T^* = 0$.

When $\sigma_b^2 > 0$, eliminating echo chambers need not be enough. Even students taught by different instructors inherit the common teacher-side shock b , so their actions remain correlated through instructor-side information. A binding precision cap can then be optimal,

but only if this common teacher-side component is sufficiently important and if shifting students toward outside signals improves social learning. This latter condition depends on the composition of outside-signal noise: the case for a binding cap is stronger when outside errors are mostly idiosyncratic, σ_η^2 , and weaker when they contain a large common component, σ_c^2 .

The results of Proposition 1 are particularly transparent in the case without common shocks.

Example 1. Suppose $\sigma_b^2 = \sigma_c^2 = 0$. In this case, the signals of students of different instructors are conditionally independent, as are outside signals of all students. This simplifies the problem and the intuitions. The planner's objective becomes

$$\mathcal{L}(T) = \left(A + \frac{B}{n+1} \right) \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + \frac{qnB}{n+1} \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2}. \quad (11)$$

Differentiating and simplifying, we get

$$T^* = \max \left\{ \frac{(2qn-1)B - (n+1)A}{B + (n+1)A} \sigma_\varepsilon^2 - \sigma_\eta^2, 0 \right\}. \quad (12)$$

Equivalently, in the no-common-shock case, $T^* > 0$ if and only if

$$\frac{(2qn-1)\frac{B}{A} - (n+1)}{\frac{B}{A} + (n+1)} > \frac{\sigma_\eta^2}{\sigma_\varepsilon^2}.$$

Thus, if echo chambers are sufficiently unlikely, or if the second action is sufficiently unimportant relative to the first action, the planner chooses $T^* = 0$. Positive T is optimal only when the social-learning gains from reducing correlated first actions are large enough to justify the loss in individual learning.

We saw that in equilibrium, the planner may impose a binding precision cap on teacher-

student communication in order to improve social learning. As a result, students may not observe their instructors' posterior means perfectly. A natural question to ask is whether the next generation (students) is necessarily more informed than the previous generation (instructors). If the planner chose $T = 0$, the students would clearly have more precise information than instructors, by virtue of learning their instructor's signal precisely plus getting an outside signal and using Bayesian updating. With $T > 0$, this is not necessarily the case because outside signals may not fully compensate for the cap-induced loss of precision in teacher-student communication. In fact, Example A2 in the Appendix shows how the planner may optimally impose a binding precision cap, $T^* > 0$, that makes students' first actions noisier than the actions their instructors would choose. Nevertheless, we show that after social learning, each student has more precise information about the state of the world θ than their instructors.

Proposition 3. *Under the planner's optimal choice T^* , $\sigma_{m,N}^2 < \sigma_b^2 + \sigma_\varepsilon^2$ and $\sigma_{m,E}^2 < \sigma_b^2 + \sigma_\varepsilon^2$.*

In other words, both students not in an echo chamber and in an echo chamber end up with more precise information after social learning than their instructors. Intuitively, if $T^* > 0$, the planner accepts a loss in individual learning only because the cap improves the informational value of subsequent social learning, so students should be more informed in the end. However, students are heterogeneous, and this logic only shows that some of the students should have lower ex-post variance. This implies that students not in echo chambers will learn more than instructors, but does not necessarily say the same about students in echo chambers. The idea of the proof is to look at the extreme case where all students are in echo chambers. In this case, T^* is even higher by Proposition 1, and the same argument as before suggests that in this case, students in echo chambers learn more, because all students are the same. The proof then bridges the gap between the two cases.

Despite learning more after the social learning stage, the students' information may be

noisier after the individual learning stage, as Example A2 shows.

4 Extensions

4.1 Dogmatism and infallibility

In the main model, the planner chooses $T^* > 0$ when the correlation in students' first actions introduced by instructor-side information is large enough to hurt social learning. It is worth noting that it is never optimal to set $T = \infty$; in other words, “ignore the instructor” is never the right policy. The reason is intuitive: if students ignored their instructors altogether, this correlation would be absent, and introducing some information transmission from instructors to students would improve the planner's objective. Thus, the optimal T^* is always finite: students always put some weight on teacher-side information, although it may be optimal to reduce that weight by imposing a binding precision cap.

By changing T , the planner balances two effects: the precision of learning from instructors versus the weights that students put on instructor and outside signals. In these terms, choosing $T^* = 0$ means that the planner opted to take the first effect (precision of teaching) to its theoretical maximum. One may wonder, however, if it may still be advantageous to have students put an even larger weight on the instructors. It turns out that the answer is affirmative: a planner would want students to put a larger weight on instructors when the information they receive through outside signals is sufficiently damaging for social learning. In the model, this happens when σ_c^2 is high, so outside signals have a systematic distortion. Realistically, this is likely when the information students obtain from outside sources is expected to be biased.

To keep things as simple as possible, assume $\sigma_b^2 = 0$, so instructors do not have common shocks. We consider two policies that planners could use to make students put higher weights

on instructors. The first is *dogmatism*; by that, we mean requiring students to learn from instructors rather than studying on their own or using empirical observations. Within the model, we capture this by a marginal increase in σ_η^2 , holding other parameters fixed. We say that dogmatism is locally desirable if the planner's utility, $-\mathcal{L}$, is locally increasing in σ_η^2 (note that full dogmatism $\sigma_\eta^2 = \infty$ cannot be optimal for the same reason as $T^* = \infty$ is never optimal: if students do not get outside signals, having some students get them and take them into account would have a positive first-order effect on the planner's objective).

The second policy we discuss is *infallibility*, by which we mean treating instructors as more reliable than they actually are. On the margin, this means ensuring that students assume that the instructors' signals have lower variance and therefore putting higher weight on these signals and lower weight on outside signals without actually changing the signals. From the planner's perspective, this would distort individual learning predictably, and the second actions of students would still be the averages of the first actions they observed from other students and themselves. We say that infallibility is locally desirable if the planner's utility is locally increasing in the weight of the teacher-side signal, when evaluated at the Bayesian weight that we modeled so far.

Proposition 4. *Suppose $A > 0$, $B > 0$, $\sigma_b^2 = 0$, and $q \in (0, 1)$. Then there exist $\bar{\sigma}_{dog}^2$ and $\bar{\sigma}_{inf}^2$ with $0 < \bar{\sigma}_{inf}^2 < \bar{\sigma}_{dog}^2 \leq \infty$ such that infallibility is locally desirable if and only if $\sigma_c^2 > \bar{\sigma}_{inf}^2$, whereas dogmatism is locally desirable if and only if $\sigma_c^2 > \bar{\sigma}_{dog}^2$. Infallibility and dogmatism can only be locally desirable if $T^* = 0$, and $\bar{\sigma}_{dog}^2$ and $\bar{\sigma}_{inf}^2$ are increasing in q , so they are more likely to be locally desirable when echo chambers are rare.*

The proposition shows that dogmatism and infallibility operate only after the precision-cap margin has been exhausted. If $T^* > 0$, the planner wants students to put less weight on teacher-side information. In that region, it cannot be beneficial to make students put even more weight on teachers, and it also cannot be beneficial to weaken outside evidence.

Once $T^* = 0$, however, teacher-student communication is already as precise as possible. Dogmatism and infallibility emerge as ways to shift students toward teachers by further distorting students' information processing, either by making them treat teachers as more reliable than they are, or by making them less responsive to outside evidence.

The force behind both distortions is the same. They can be useful only when outside evidence contains a sufficiently large common component, σ_c^2 . If outside evidence is mostly idiosyncratic, reducing reliance on it only destroys useful information and worsens both individual and social learning. But if the outside signal contains a large systematic distortion, then students who rely on outside evidence make correlated mistakes. In that case, shifting weight toward teachers can improve social learning. The role of echo chambers is especially stark. When all students are in echo chambers, instructor-specific errors are not diversified through peer interaction, so making students rely even more on instructors is not useful. By contrast, when q is low, instructor-specific errors are diversified across peers, while common outside errors are not. This is the environment in which the planner may want students to put additional weight on teachers.

Finally, infallibility is easier to justify than dogmatism. Infallibility changes the weights students put on teacher-side and outside information while holding the true information structure fixed. Dogmatism also shifts students away from outside evidence, but it does so by making that evidence less informative. Thus, in reducing reliance on the common outside component, it also destroys useful idiosyncratic outside information. For this reason, infallibility is preferred to dogmatism, and in particular for σ_c^2 sufficiently high infallibility always becomes locally desirable, while the same is not necessarily true for dogmatism.

4.2 Partial echo chambers

The previous benchmark considered two polar cases: a student either observes only peers taught by different instructors, or only peers taught by the same instructor. We now allow for the intermediate case. Each student observes the actions of n other students. Of these, n_1 are students of the same instructor, while $n_2 = n - n_1$ are students of different instructors.

For simplicity, we will focus on the no-common-shocks case: $p_\theta = 0$ and $\sigma_b^2 = \sigma_c^2 = 0$. There are two natural ways to proceed. First, one can assume that social learning takes the form of simple averaging, as in DeGroot-style models of social learning. This averaging is similar to social learning in the main model; however, in the main model this coincided with Bayesian inference (as long as the student could only condition on their posterior after individual learning but not separate signals) because all peers the student learned from were equally informed, whereas here this departs from Bayesian updating, because actions of students from different instructors carry more information. The second way to proceed is to assume a Bayesian rule; this would imply, in particular, that since the correlation between actions of students of the same instructor is higher than that of students of different instructors (which under the assumption $\sigma_b^2 = 0$ is zero), the weights put on the actions of different instructors should be higher.

Under the DeGroot-style approach, students choose the second action equal to

$$v_{is} = \frac{u_{is} + \sum_{k=1}^n u_{isk} r_{sk}}{n+1}.$$

The planner, however, evaluates this action under the true covariance structure. Since the average contains $n+1$ actions and n_1+1 of them come from students taught by the same instructor, the planner's objective is

$$\mathcal{L}_D(T) = \left(A + \frac{B}{n+1} \right) \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + B \frac{n_1(n_1+1)}{(n+1)^2} \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2}.$$

Note the similarity to (11): the expressions become identical if $n_1 = n$ and $q = 1$. As such, we have a closed-form solution

$$T_D^* = \max \left\{ \left[\frac{2Bn_1(n_1 + 1)}{(n + 1)(A(n + 1) + B)} - 1 \right] \sigma_\varepsilon^2 - \sigma_\eta^2, 0 \right\}.$$

This expression nests the two polar cases. If $n_1 = 0$, then $T_D^* = 0$. If $n_1 = n$, the expression reduces to the all-same-instructor benchmark (12).

Alternatively, consider students using Bayesian weights. In this case, students recognize that same-instructor actions are correlated and therefore put lower per-action weight on the same-instructor block than on different-instructor actions. These weights now depend on T , which complicates the analysis and may make the minimization problem non-convex. As a result, the planner's objective need not yield a unique interior critical point (see Example A3). Nevertheless, comparative statics are similar, understood in the monotone-set sense of Milgrom and Shannon (1994).

Proposition 5. *Consider the no-common-shocks benchmark with $p_\theta = 0$, $\sigma_b^2 = \sigma_c^2 = 0$, and $T \geq 0$. Suppose each student observes n_1 same-instructor peer actions and n_2 different-instructor peer actions, with $n = n_1 + n_2$.*

Under DeGroot social learning, the optimal precision cap is unique and is weakly increasing in the relative weight B/A , weakly increasing in n_1 holding n fixed, weakly increasing in σ_ε^2 , and weakly decreasing in σ_η^2 . If $A = 0$, so that the planner only values the post-social-learning action, then it is also weakly increasing in n holding the ratio n_1/n_2 fixed.

Under Bayesian social learning, the optimal precision cap need not be unique, but the same comparative statics hold in the monotone-set sense of Milgrom and Shannon (1994). That is, increases in the relative weight B/A and increases in same-instructor exposure (increase of n_1/n_2 , holding n fixed) shift the set of optimal choices of T upward in the strong set order, as do an increase in σ_ε^2 , a decrease in σ_η^2 , and, provided that $A = 0$, an increase in n holding

n_1/n_2 fixed.

The proposition shows that the main logic of the benchmark generalizes. A precision cap is useful only because same-instructor peer actions are redundant. When students mostly learn from different-instructor peers, their peer actions are conditionally independent, and garbling the teacher’s communication only worsens first-stage learning. By contrast, when students learn from peers taught by the same instructor, their actions contain a common instructor error. A tighter precision cap weakens the dependence of students’ actions on that common error and makes later social learning less redundant.

The DeGroot formulation gives a smooth version of this result. Since students mechanically average all observed actions, the planner’s second-stage loss is simply the variance of that average under the true covariance matrix. The problem reduces to minimizing a weighted sum of two terms: one increasing in T , coming from the loss of individual precision, and one decreasing in T , coming from the reduction in same-instructor covariance. This yields the closed-form expression above. The fully Bayesian case does not yield a simple closed-form solution, but the comparative statics results hold.

4.3 Dynamics of learning

The static analysis above can be extended to study societal learning over several generations. Suppose that time is discrete and indexed by $t \geq 0$. In each generation t , every instructor teaches a continuum of students, as before. After students learn from their instructors, receive outside signals, and observe peers’ first actions, one representative student of each instructor is randomly selected to become an instructor in generation $t+1$. These instructors are assigned a new generation of students.

Since each instructor i in generation $t \geq 1$ is a former student, she starts period t with

posterior mean

$$a_{it} = \theta + b_t + \varepsilon_{it},$$

where b_t is the component common to all instructors in generation t , and ε_{it} is the instructor-specific component. Unlike instructors in period 0, these instructors' values are not chosen by Nature but rather determined in the previous period through learning. We denote the variances of these components by σ_{bt}^2 and $\sigma_{\varepsilon t}^2$, respectively. Student s of this instructor will observe $y_{ist} = a_{it} + \tau_{ist}$, with variance $\sigma_{\tau t}^2 = T_t$. This variance T_t is the choice of period t planner, who, we assume, is short-lived and cares about actions taken by students in period t . As before, students get outside signals $z_{ist} = \theta + c_t + \eta_{ist}$, with c_t being the common component and η_{ist} being idiosyncratic. If $\{c_t\}$ are the same for all periods, we will call outside shocks persistent, otherwise we will call them transient.

To keep the dynamic comparison simple, we focus on the two polar peer-learning structures. Either $q = 0$, so students learn only from peers taught by different instructors, or $q = 1$, so students learn only from peers taught by the same instructor. In these two cases, all students in a generation have the same ex ante distribution of final beliefs, and hence the next generation's teachers are drawn from a normal distribution as in the baseline model. In this way, all instructors within a generation have the same precision of posterior beliefs, which makes the analysis manageable.

In the no-common-shocks benchmark, the dynamic analysis is especially simple. Suppose first that $\sigma_{b0}^2 = \sigma_b^2 = 0$ and $\sigma_c^2 = 0$. Then students' final beliefs in period 0 do not contain a component that is common across all future instructors. Hence $\sigma_{bt}^2 = 0$ for all t , and the dynamic state is one-dimensional: only $\sigma_{\varepsilon t}^2$ matters.

When $q = 0$, students learn from different-instructor peers, so peer actions are conditionally independent and the optimal precision cap is always 0. When $q = 1$, students learn only inside echo chambers, so a nontrivial precision cap may be optimal in early genera-

tions. However, students' final beliefs must be more precise than the instructor beliefs they inherited. Hence $\sigma_{\varepsilon,t+1}^2 < \sigma_{\varepsilon t}^2$. As a result, the optimal T_t^* is weakly decreasing over time, by Proposition 1. Furthermore, when $\sigma_{\varepsilon t}^2$ is sufficiently small, T_t^* becomes zero, and it necessarily happens in a finite number of periods. In other words, the dynamics of T^* is monotonic under either value of q in this benchmark case.

With common shocks, the dynamics is no longer captured by a one-dimensional variable. Students' final errors contain both a common component and an instructor-specific component, and these become the common and instructor-specific components of the next generation's teacher beliefs. Hence the relevant state is $(\sigma_{bt}^2, \sigma_{\varepsilon t}^2)$. A decline in total variance need not imply that both components decline. A transient outside common shock may be partly incorporated into students' final beliefs and therefore become part of the next generation's common teacher-side component.

The persistence of the outside common shock is crucial. If the outside common shock is persistent across generations, then full learning occurs with probability zero. Indeed, aggregation eliminates idiosyncratic errors, but the initial distortions by common shocks b_0 and c_0 will remain, and in the long run the posterior of every student will converge not to the true value of θ but rather to its value distorted by these initial shocks.

If, instead, outside common shocks are transient and independent across generations, then long-run learning is restored, because these outside shocks together contain precise information about θ . To see the logic, note first that $T = 0$ is always feasible. Under $T = 0$, each generation combines inherited teacher beliefs with fresh outside information. Although the outside common shock is shared by all students within a generation, it is independent across generations and therefore provides fresh information over time. Thus total teacher-side variance falls under the no-cap policy.⁴ The optimal policy cannot do worse in terms

⁴Note that if the minimal T were not zero (e.g., $T \geq T_{\min} > 0$), then there would be no convergence to true θ . The influence of older outside shocks on generation t 's knowledge would decrease exponentially over time, and thus the learning process would not converge to a point. This effect is similar to learning about

of learning: if the planner chooses $T_t^* > 0$, it must be because the precision cap improves final social learning enough to compensate for the loss in first-stage learning. Hence the total variance of next generation’s teacher beliefs under the optimal cap is no larger than under the no-cap policy. Therefore, with transient outside common shocks, total teacher-side variance converges to zero under the optimal policy, even though the separate components σ_{bt}^2 and $\sigma_{\varepsilon t}^2$ may move non-monotonically along the transition. In particular, this may imply a nonmonotone path of T_t^* , as Example A4 shows.

The following proposition summarizes the dynamic implications.

Proposition 6. *In the no-common-shocks benchmark, if $\sigma_{b0}^2 = 0$ and $\sigma_{\varepsilon}^2 = 0$, then there is full learning in the limit: the posterior variance of beliefs of students of generation t converges to zero. Moreover, the optimal precision cap becomes weakly less binding over time. In particular, $T_t^* = 0$ for all t when $q = 0$, while for $q = 1$ the sequence T_t^* is weakly decreasing and reaches zero in finite time.*

With common shocks, if the common component of outside shocks is persistent, then full learning does not occur. If outside common shocks are transient and independent across generations, then full learning is achieved and the optimal precision cap T_t^ reaches zero in finite time; however, its path may be non-monotone.*

The proposition highlights that the static tradeoff becomes weaker over time when learning is successful. In the no-common-shocks benchmark, each generation improves the information of the next one, so the correlation problem that motivates a precision cap eventually disappears. With common shocks, long-run learning depends on persistence: persistent common shocks prevent full learning, while transient shocks are eventually averaged out. The transition, however, need not be monotone, because early generations may reduce instructor-specific dispersion while temporarily importing common outside distortions into the teacher

an evolving state of the world, where older signals are remembered but become increasingly obsolete, as in Dasaratha, Golub and Hak (2023).

population.

Finally, one can compare the short-lived planners considered above with a planner who can commit to a full path $\{T_t\}_{t \geq 0}$ and discounts students' actions in different periods with discount factor β^t for $\beta \in (0, 1)$. A committed planner internalizes that students' final beliefs in period t determine the distribution of instructors in period $t + 1$. Therefore, relative to a short-lived planner, commitment increases the value of final learning today. In the no-common-shocks case, the committed planner's Euler equation is equivalent to replacing the current-period weight B on final actions with a larger effective weight that also includes the continuation value of improving future teachers. Thus a committed planner may impose a tighter precision cap in early generations than a short-lived planner. However, as teacher beliefs become more accurate, the same force that drives the short-lived cap to zero also applies under commitment: eventually, the benefit of garbling teacher communication disappears, and the optimal cap becomes nonbinding.

5 Conclusion

This paper develops a theory of viewpoint neutrality as an informational policy. The central idea is simple: a teacher's information may be useful for a student learning on her own, but if many students receive similar teacher-side information, their actions become correlated signals for one another. Precise instruction can then make peer learning less informative, since students' actions reflect less independent information. A planner may therefore optimally limit the precision of teacher-student communication, not because teacher information lacks value, but because precise transmission can make that information socially overrepresented in later learning.

The planner does not know which teacher views are correct, only that teacher-side information may be correlated across students. The response is thus not to select the right

viewpoint, but to limit how sharply teacher-side views enter peer networks before students learn from each other. The rationale is consequentialist: neutrality is valuable as an instrument to reduce the overrepresentation of teacher-side information in homophilous social learning. Viewpoint neutrality is therefore not a claim that teachers lack expertise, nor is it a demand that all views be treated as equally plausible. It is a restraint on the social amplification of teacher-side information.

This logic suggests why neutrality matters more in some fields than others, and more at some times than others. What distinguishes fields is not whether instructors agree, but whether agreement reflects independent verification. Where claims are settled by proof, experiment, and replication, a common instructor view usually reflects what has survived independent checks, and limiting its transmission mainly discards information. Where evidence is thinner and interpretation carries more weight, as in much of the social sciences, law, and ethics, though nowhere exclusively, a shared view may instead reflect a common framework, and homophilous peer learning can entrench it rather than test it. Our dynamic analysis suggests that this comparison is not fixed over time. As a body of knowledge matures, instructor views become more reliable and the rationale for viewpoint neutrality diminishes, though the path need not be monotonic, and neutrality can matter more before it matters less.

Several directions for future research show promise. First, teachers' bias may be partly known to the planner, who would then want the learning process to counter it; how this motive interacts with the aggregation force studied here is an open question. Second, because it is ultimately the teacher's decision what to say in the classroom, implementation may require a combination of incentives, professional norms, and rules, and understanding how such instruments should be designed is an important problem. Finally, not all parameters of the model are easy to observe, which raises the question of how to design transparent, easy-to-check criteria for when neutrality should apply. Answering these questions would

help make viewpoint neutrality a more disciplined and effective policy.

References

- Acemoglu, Daron, Munther A. Dahleh, Ilan Lobel and Asuman Ozdaglar. 2011. “Bayesian Learning in Social Networks.” *The Review of Economic Studies* 78(4):1201–1236.
- Allcott, Hunt and Matthew Gentzkow. 2017. “Social Media and Fake News in the 2016 Election.” *Journal of Economic Perspectives* 31(2):211–236.
- American Association of University Professors. 2025. “On Institutional Neutrality.” Policy statement.
- Annest, Janna J. 2002. “Only the News That’s Fit to Print: The Effect of Hazelwood on the First Amendment Viewpoint-Neutrality Requirement in Public School-Sponsored Forums.” *Washington Law Review* 77:1227–1260.
- Austen-Smith, David and Roland G. Fryer. 2005. “An Economic Analysis of “Acting White”.” *The Quarterly Journal of Economics* 120(2):551–583.
- Banerjee, Abhijit V. 1992. “A Simple Model of Herd Behavior.” *The Quarterly Journal of Economics* 107(3):797–817.
- Bergemann, Dirk and Stephen Morris. 2016. “Bayes Correlated Equilibrium and the Comparison of Information Structures in Games.” *Theoretical Economics* 11(2):487–522.
- Bikhchandani, Sushil, David Hirshleifer and Ivo Welch. 1992. “A Theory of Fads, Fashion, Custom, and Cultural Change as Informational Cascades.” *Journal of Political Economy* 100(5):992–1026.

- Bikhchandani, Sushil, David Hirshleifer, Omer Tamuz and Ivo Welch. 2024. “Information Cascades and Social Learning.” *Journal of Economic Literature* 62(3):1040–1093.
- Board, Simon and Moritz Meyer-ter Vehn. 2021. “Learning Dynamics in Social Networks.” *Econometrica* 89(6):2601–2635.
- Board, Simon and Moritz Meyer-ter Vehn. 2024. “Experimentation in Networks.” *American Economic Review* 114(9):2940–2980.
- Bursztyn, Leonardo, Georgy Egorov and Robert Jensen. 2019. “Cool to Be Smart or Smart to Be Cool? Understanding Peer Pressure in Education.” *The Review of Economic Studies* 86(4):1487–1526.
- Bursztyn, Leonardo and Robert Jensen. 2015. “How Does Peer Pressure Affect Educational Investments?” *The Quarterly Journal of Economics* 130(3):1329–1367.
- Carrell, Scott E., Bruce I. Sacerdote and James E. West. 2013. “From Natural Variation to Optimal Policy? The Importance of Endogenous Peer Group Formation.” *Econometrica* 81(3):855–882.
- Chandrasekhar, Arun G., Horacio Larreguy and Juan Pablo Xandri. 2020. “Testing Models of Social Learning on Networks: Evidence from Two Experiments.” *Econometrica* 88(1):1–32.
- Chetty, Raj, John N. Friedman and Jonah E. Rockoff. 2014. “Measuring the Impacts of Teachers II: Teacher Value-Added and Student Outcomes in Adulthood.” *American Economic Review* 104(9):2633–2679.
- Currarini, Sergio, Matthew O. Jackson and Paolo Pin. 2009. “An Economic Model of Friendship: Homophily, Minorities, and Segregation.” *Econometrica* 77(4):1003–1045.

- Dasaratha, Krishna, Benjamin Golub and Nir Hak. 2023. “Learning from Neighbours about a Changing State.” *Review of Economic Studies* 90(5):2326–2369.
- DeGroot, Morris H. 1974. “Reaching a Consensus.” *Journal of the American Statistical Association* 69(345):118–121.
- DeMarzo, Peter M., Dimitri Vayanos and Jeffrey Zwiebel. 2003. “Persuasion Bias, Social Influence, and Unidimensional Opinions.” *Quarterly Journal of Economics* 118(3):909–968.
- Dewey, John. 1916. *Democracy and Education: An Introduction to the Philosophy of Education*. New York: The Macmillan Company.
- Dunn, Alyssa Hadley, Beth Sondel and Hannah Carson Baggett. 2019. ““I Don’t Want to Come Off as Pushing an Agenda”: How Contexts Shaped Teachers’ Pedagogy in the Days after the 2016 U.S. Presidential Election.” *American Educational Research Journal* 56(2):444–476.
- Eyster, Erik and Matthew Rabin. 2010. “Naive Herding in Rich-Information Settings.” *American Economic Journal: Microeconomics* 2(4):221–243.
- Feddersen, Timothy and Wolfgang Pesendorfer. 1998. “Convicting the Innocent: The Inferiority of Unanimous Jury Verdicts under Strategic Voting.” *American Political Science Review* 92(1):23–35.
- Gentzkow, Matthew and Jesse M. Shapiro. 2011. “Ideological Segregation Online and Offline.” *The Quarterly Journal of Economics* 126(4):1799–1839.
- Golub, Benjamin and Matthew O. Jackson. 2010. “Naive Learning in Social Networks and the Wisdom of Crowds.” *American Economic Journal: Microeconomics* 2(1):112–149.
- Hazla, Jan, Ali Jadbabaie, Elchanan Mossel and M. Amin Rahimian. 2021. “Bayesian Decision Making in Groups Is Hard.” *Operations Research* 69(2):632–654.

- Hess, Diana E. and Paula McAvoy. 2015. *The Political Classroom: Evidence and Ethics in Democratic Education*. New York: Routledge.
- Kalven Committee. 1967. Report on the University's Role in Political and Social Action. Technical report University of Chicago. Commonly known as the Kalven Report.
- Kamenica, Emir and Matthew Gentzkow. 2011. "Bayesian Persuasion." *American Economic Review* 101(6):2590–2615.
- Kelly, Thomas E. 1986. "Discussing Controversial Issues: Four Perspectives on the Teacher's Role." *Theory and Research in Social Education* 14(2):113–138.
- Levy, Gilat and Ronny Razin. 2015. "Correlation Neglect, Voting Behavior, and Information Aggregation." *American Economic Review* 105(4):1634–1645.
- McPherson, Miller, Lynn Smith-Lovin and James M. Cook. 2001. "Birds of a Feather: Homophily in Social Networks." *Annual Review of Sociology* 27:415–444.
- Milgrom, Paul and Chris Shannon. 1994. "Monotone Comparative Statics." *Econometrica* 62(1):157–180.
- Mill, John Stuart. 1859. *On Liberty*. London: John W. Parker and Son.
- Morris, Stephen and Hyun Song Shin. 2002. "Social Value of Public Information." *American Economic Review* 92(5):1521–1534.
- Mostagir, Mohamed, Asuman Ozdaglar and James Siderius. 2022. "When Is Society Susceptible to Manipulation?" *Management Science* 68(10):7153–7175.
- Post, Robert C. 2024. "The Kalven Report, Institutional Neutrality, and Academic Freedom." Yale Law School, Public Law Research Paper.

- Sacerdote, Bruce. 2001. “Peer Effects with Random Assignment: Results for Dartmouth Roommates.” *The Quarterly Journal of Economics* 116(2):681–704.
- Salzman, Dylan. 2022. “The Constitutionality of Orthodoxy: First Amendment Implications of Laws Restricting Critical Race Theory in Public Schools.” *University of Chicago Law Review* 89(4):1069–1112.
- Soucek, Brian. 2026. *The Opinionated University: Academic Freedom, Diversity, and the Myth of Neutrality in American Higher Education*. Chicago: University of Chicago Press.
- Sunstein, Cass R. 2002. “The Law of Group Polarization.” *Journal of Political Philosophy* 10(2):175–195.
- Sunstein, Cass R. 2009. *Going to Extremes: How Like Minds Unite and Divide*. New York: Oxford University Press.
- Tincani, Michela. forthcoming. “Peer Effects and Rank Concerns in the Classroom.” *Journal of Political Economy* .
- Zimmerman, David J. 2003. “Peer Effects in Academic Outcomes: Evidence from a Natural Experiment.” *The Review of Economics and Statistics* 85(1):9–23.

Appendix

Proof of Proposition 1. Denote $K = B + (n + 1)A$. Differentiating (9) with respect to T gives

$$\frac{\partial \mathcal{L}(T)}{\partial T} = \frac{\sigma_z^2}{(\sigma_y^2(T) + \sigma_z^2)^3} \left(\frac{K}{n+1} \sigma_z^2 (\sigma_y^2(T) + \sigma_z^2) + \frac{2nB}{n+1} (\sigma_c^2 \sigma_y^2(T) - (\sigma_b^2 + q\sigma_\varepsilon^2) \sigma_z^2) \right).$$

The sign of the derivative is the same as the sign of the second factor. This expression is affine in $\sigma_y^2(T)$, and its coefficient at $\sigma_y^2(T)$ is

$$\frac{K}{n+1} \sigma_z^2 + \frac{2nB}{n+1} \sigma_c^2 > 0.$$

Therefore, the derivative crosses zero at most once and if so, it changes the sign from negative to positive. If it is nonnegative at $T = 0$, the objective is minimized at $T^* = 0$. Otherwise, there is a unique interior point at which the derivative is zero, and this point is the unique minimum. Hence T^* is unique.

The unconstrained solution to the first-order condition is

$$T_u = \frac{\sigma_z^2 [2nB(\sigma_b^2 + q\sigma_\varepsilon^2) - K\sigma_z^2]}{K\sigma_z^2 + 2nB\sigma_c^2} - \sigma_b^2 - \sigma_\varepsilon^2. \quad (\text{A1})$$

Thus, $T^* = \max\{T_u, 0\}$, and it suffices to study comparative statics of T_u with respect to the variables of interest.

First,

$$\frac{\partial T_u}{\partial q} = \frac{2nB\sigma_\varepsilon^2 \sigma_z^2}{K\sigma_z^2 + 2nB\sigma_c^2} \geq 0,$$

with strict inequality if $\sigma_\varepsilon^2 > 0$. Hence T^* is weakly increasing in q , and strictly increasing at an interior solution whenever $\sigma_\varepsilon^2 > 0$.

Second,

$$\frac{\partial T_u}{\partial n} = \frac{2B(A+B)(\sigma_z^2)^2(\sigma_b^2 + q\sigma_\varepsilon^2 + \sigma_c^2)}{(K\sigma_z^2 + 2nB\sigma_c^2)^2} \geq 0.$$

Differentiating with respect to A yields

$$\frac{\partial T_u}{\partial A} = -\frac{2nB(n+1)(\sigma_z^2)^2(\sigma_b^2 + q\sigma_\varepsilon^2 + \sigma_c^2)}{(K\sigma_z^2 + 2nB\sigma_c^2)^2} \leq 0;$$

the derivative with respect to B is similar but has the opposite sign. Thus, T^* is decreasing in A and increasing in B .

It remains to discuss comparative statics with respect to σ_ε^2 and σ_η^2 . Differentiating the unconstrained solution with respect to σ_ε^2 , holding the other parameters fixed, gives

$$\frac{\partial T_u}{\partial \sigma_\varepsilon^2} = \frac{2qnB\sigma_z^2}{K\sigma_z^2 + 2nB\sigma_c^2} - 1.$$

Thus, outside the no-common-shock case, the sign of this comparative static need not be globally positive. Increasing instructor-specific noise both creates more redundancy inside echo chambers and makes the instructor-side signal less useful.

Differentiating with respect to σ_η^2 , recalling that $\sigma_z^2 = \sigma_c^2 + \sigma_\eta^2$, gives

$$\frac{\partial T_u}{\partial \sigma_\eta^2} = \frac{4n^2B^2\sigma_c^2(\sigma_b^2 + q\sigma_\varepsilon^2) - 4nBK\sigma_c^2\sigma_z^2 - K^2(\sigma_z^2)^2}{(K\sigma_z^2 + 2nB\sigma_c^2)^2}.$$

This expression also need not have a global sign for arbitrary common shocks. A noisier outside signal makes outside learning less useful, but when outside information contains a common component, changing σ_η^2 also changes the relative importance of the common and idiosyncratic parts of outside information.

When $\sigma_b^2 = \sigma_c^2 = 0$, however, the unconstrained solution becomes

$$T_u = \left(\frac{2qnB}{K} - 1 \right) \sigma_\varepsilon^2 - \sigma_\eta^2.$$

Therefore,

$$T^* = \max \left\{ \left(\frac{2qnB}{K} - 1 \right) \sigma_\varepsilon^2 - \sigma_\eta^2, 0 \right\}.$$

This expression is weakly increasing in σ_ε^2 and weakly decreasing in σ_η^2 . Moreover, if $T^* > 0$, then necessarily

$$\frac{2qnB}{K} - 1 > 0,$$

so the effect of σ_ε^2 is strictly positive, while the effect of σ_η^2 is strictly negative.

Finally, because T_u , $\partial T_u / \partial \sigma_\varepsilon^2$, and $\partial T_u / \partial \sigma_\eta^2$ are continuous in σ_b^2 and σ_c^2 , these signs extend locally to sufficiently small common shocks whenever the solution is not exactly at the threshold $T_u = 0$. At the threshold, $T^* = 0$, so the weak comparative statics follow from continuity of the constrained solution $T^* = \max\{T_u, 0\}$. Thus, if σ_b^2 and σ_c^2 are sufficiently small, T^* is weakly increasing in σ_ε^2 and weakly decreasing in σ_η^2 , with strict inequalities whenever $T^* > 0$. This proves the proposition. ■

Proof of Corollary 2. Rewriting (A1) for $q = 0$ yields

$$T_u = \frac{\sigma_z^2 [2nB\sigma_b^2 - K\sigma_z^2]}{K\sigma_z^2 + 2nB\sigma_c^2} - \sigma_b^2 - \sigma_\varepsilon^2.$$

For σ_b^2 sufficiently small, this expression is negative, so optimal $T^* = 0$. Furthermore, the coefficient at σ_b^2 equals

$$\frac{2nB\sigma_\eta^2 - K(\sigma_c^2 + \sigma_\eta^2)}{K(\sigma_c^2 + \sigma_\eta^2) + 2nB\sigma_c^2}.$$

It is therefore positive whenever the condition (10) holds. If so, then for sufficiently high σ_b^2 , $T_u > 0$, and therefore $T^* > 0$. This completes the proof. ■

Proof of Proposition 3. We will use notation from the proof of Proposition 1. Note that if $\sigma_z^2 = 0$, then each student observes the state perfectly through the outside signal, and the result is immediate. We therefore consider the case $\sigma_z^2 > 0$.

Echo-chamber students have weakly higher second-stage variance than students outside echo chambers:

$$\sigma_{m,E}^2(T) - \sigma_{m,N}^2(T) = \frac{n}{n+1} \frac{(\sigma_z^2)^2 \sigma_\varepsilon^2}{(\sigma_y^2(T) + \sigma_z^2)^2} \geq 0.$$

It therefore suffices to prove that $\sigma_{m,E}^2(T^*) < \sigma_b^2 + \sigma_\varepsilon^2$.

We have

$$\sigma_{m,E}^2(T) = \frac{1}{n+1} \frac{\sigma_y^2(T) \sigma_z^2}{\sigma_y^2(T) + \sigma_z^2} + \frac{n}{n+1} \frac{(\sigma_z^2)^2 (\sigma_b^2 + \sigma_\varepsilon^2) + (\sigma_y^2(T))^2 \sigma_c^2}{(\sigma_y^2(T) + \sigma_z^2)^2}.$$

Differentiating with respect to T , which enters the formula only through $\sigma_y^2(T)$, gives

$$\frac{d\sigma_{m,E}^2(T)}{dT} = \frac{\sigma_z^2}{(n+1) (\sigma_y^2(T) + \sigma_z^2)^3} \left(\sigma_z^2 (\sigma_y^2(T) + \sigma_z^2) + 2n (\sigma_c^2 \sigma_y^2(T) - (\sigma_b^2 + \sigma_\varepsilon^2) \sigma_z^2) \right).$$

The second factor is affine and strictly increasing in $\sigma_y^2(T)$, and hence in T . Thus $\sigma_{m,E}^2(T)$ decreases up to its constrained minimizer and increases afterward. Solving the linear equation and using $\sigma_y^2(T) = \sigma_b^2 + \sigma_\varepsilon^2 + T$, we obtain the following expression for this constrained minimizer:

$$\bar{T} = \max \left\{ \frac{\sigma_z^2 (2n(\sigma_b^2 + \sigma_\varepsilon^2) - \sigma_z^2)}{\sigma_z^2 + 2n\sigma_c^2} - \sigma_b^2 - \sigma_\varepsilon^2, 0 \right\}.$$

Consequently, $\sigma_{m,E}^2(T)$ is weakly decreasing on $[0, \bar{T}]$.

We next show that the planner's chosen cap satisfies $T^* \leq \bar{T}$. From the proof of Proposition 1, the planner's unconstrained solution (A1) satisfies

$$T_u = \frac{\sigma_z^2 [2nB(\sigma_b^2 + q\sigma_\varepsilon^2) - K\sigma_z^2]}{K\sigma_z^2 + 2nB\sigma_c^2} - \sigma_b^2 - \sigma_\varepsilon^2.$$

The fraction on the right-hand side is increasing in $\sigma_b^2 + q\sigma_\varepsilon^2$ and decreasing in K . Since $\sigma_b^2 + q\sigma_\varepsilon^2 \leq \sigma_b^2 + \sigma_\varepsilon^2$ and $K \geq B$, we have

$$T_u \leq \frac{\sigma_z^2 (2n(\sigma_b^2 + \sigma_\varepsilon^2) - \sigma_z^2)}{\sigma_z^2 + 2n\sigma_c^2} - \sigma_b^2 - \sigma_\varepsilon^2.$$

Therefore

$$T^* = \max\{T_u, 0\} \leq \max\left\{\frac{\sigma_z^2 (2n(\sigma_b^2 + \sigma_\varepsilon^2) - \sigma_z^2)}{\sigma_z^2 + 2n\sigma_c^2} - \sigma_b^2 - \sigma_\varepsilon^2, 0\right\} = \bar{T}.$$

It remains to show that the echo-chamber student's variance is already below the instructor's variance at $T = 0$. At $T = 0$, $\sigma_y^2(0) = \sigma_b^2 + \sigma_\varepsilon^2$, and direct substitution gives

$$\sigma_y^2(0) - \sigma_{m,E}^2(0) = \frac{(\sigma_b^2 + \sigma_\varepsilon^2)^2 [(n+1)(\sigma_b^2 + \sigma_\varepsilon^2) + (n+1)\sigma_c^2 + (2n+1)\sigma_\eta^2]}{(n+1)(\sigma_b^2 + \sigma_\varepsilon^2 + \sigma_z^2)^2} > 0.$$

Thus, $\sigma_{m,E}^2(0) < \sigma_b^2 + \sigma_\varepsilon^2$. Since $T^* \leq \bar{T}$ and $\sigma_{m,E}^2(T)$ is weakly decreasing on $[0, \bar{T}]$, we have $\sigma_{m,E}^2(T^*) \leq \sigma_{m,E}^2(0) < \sigma_b^2 + \sigma_\varepsilon^2$, and thus the same holds for $\sigma_{m,N}^2(T^*)$. This completes the proof. ■

Proof of Proposition 4. A policy is locally desirable if it locally decreases $\mathcal{L}$ at $T = 0$. If students put weight ω on the teacher-side signal, then, given $\sigma_b^2 = 0$, the planner's loss is

$$\begin{aligned} \mathcal{L}(\omega, T) &= A [\omega^2(\sigma_\varepsilon^2 + T) + (1 - \omega)^2(\sigma_c^2 + \sigma_\eta^2)] \\ &\quad + B \left[\omega^2 \frac{\sigma_\varepsilon^2 + T + qn\sigma_\varepsilon^2}{n+1} + (1 - \omega)^2 \left(\sigma_c^2 + \frac{\sigma_\eta^2}{n+1} \right) \right]. \end{aligned}$$

The Bayesian weight on the teacher-side signal is

$$\omega^B(T) = \frac{\sigma_c^2 + \sigma_\eta^2}{\sigma_\varepsilon^2 + T + \sigma_c^2 + \sigma_\eta^2}.$$

We first show that when $T^* > 0$, neither policy is locally desirable. At the Bayesian weight,

$$\left. \frac{\partial \mathcal{L}}{\partial \omega} \right|_{\omega=\omega^B(T)} = \frac{2Bn}{(n+1)(\sigma_\varepsilon^2 + T + \sigma_c^2 + \sigma_\eta^2)} [q\sigma_\varepsilon^2(\sigma_c^2 + \sigma_\eta^2) - (\sigma_\varepsilon^2 + T)\sigma_c^2].$$

Note that,

$$\begin{aligned} \frac{d\mathcal{L}(\omega^B(T), T)}{dT} &= \frac{\sigma_c^2 + \sigma_\eta^2}{(n+1)(\sigma_\varepsilon^2 + T + \sigma_c^2 + \sigma_\eta^2)^3} \\ &\quad \times \left[A(n+1)(\sigma_c^2 + \sigma_\eta^2)(\sigma_\varepsilon^2 + T + \sigma_c^2 + \sigma_\eta^2) \right. \\ &\quad \left. + B((\sigma_c^2 + \sigma_\eta^2)(\sigma_\varepsilon^2 + T + \sigma_c^2 + \sigma_\eta^2) \right. \\ &\quad \left. + 2n((\sigma_\varepsilon^2 + T)\sigma_c^2 - q\sigma_\varepsilon^2(\sigma_c^2 + \sigma_\eta^2))) \right]. \end{aligned}$$

If $T^* > 0$, this last expression equals 0 at T^* . Since the terms in the first three lines are positive, we must have $(\sigma_\varepsilon^2 + T^*)\sigma_c^2 - q\sigma_\varepsilon^2(\sigma_c^2 + \sigma_\eta^2) < 0$. But this implies that $\left. \frac{\partial \mathcal{L}}{\partial \omega} \right|_{\omega=\omega^B(T^*)} > 0$. Thus, increasing the teacher-side weight increases the loss, so infallibility is not locally desirable. Dogmatism is not locally desirable either, as it both raises $\omega^B(T)$, as infallibility does, and also increases σ_η^2 , which also increases the loss directly (holding $\omega^B(T)$ fixed). Thus both policies can only be locally desirable if $T^* = 0$.

It remains to characterize the conditions for local desirability of the two policies at $T = 0$. For infallibility, setting $T = 0$ and rearranging the expression for $\left. \frac{\partial \mathcal{L}}{\partial \omega} \right|_{\omega=\omega^B(T^*)}$, we get that infallibility is locally desirable if and only if $q\sigma_\eta^2 - (1-q)\sigma_c^2 < 0$, or, equivalently, if $\sigma_c^2 > \bar{\sigma}_{\text{inf}}^2$, as stated.

For dogmatism, setting $T = 0$ and differentiating $\mathcal{L}(\omega^B(T), T)$ with respect to σ_η^2 yields

$$\begin{aligned} \frac{d\mathcal{L}(\omega^B(0), 0)}{d\sigma_\eta^2} &= \frac{\sigma_\varepsilon^4}{(n+1)(\sigma_\varepsilon^2 + \sigma_c^2 + \sigma_\eta^2)^3} \left[A(n+1)(\sigma_c^2 + \sigma_\varepsilon^2 + \sigma_\eta^2) \right. \\ &\quad \left. + B((1 - 2n(1-q))\sigma_c^2 + \sigma_\varepsilon^2 + (1 + 2qn)\sigma_\eta^2) \right]. \end{aligned}$$

The sign of this expression is determined by the term in brackets. Note that it is positive at $\sigma_c^2 = 0$, and also nondecreasing in σ_c^2 if $B(2n(1-q) - 1) - A(n+1) \leq 0$. In this case, therefore, dogmatism is not locally optimal for any σ_c^2 , so $\bar{\sigma}_{\text{dog}}^2 = \infty$. Otherwise, dogmatism is locally desirable if and only if $\sigma_c^2 > \bar{\sigma}_{\text{dog}}^2$, for

$$\bar{\sigma}_{\text{dog}}^2 = \frac{A(n+1)(\sigma_\varepsilon^2 + \sigma_\eta^2) + B(\sigma_\varepsilon^2 + (1+2qn)\sigma_\eta^2)}{B(2n(1-q) - 1) - A(n+1)}.$$

Lastly, notice that whenever $\bar{\sigma}_{\text{dog}}^2 < \infty$,

$$\bar{\sigma}_{\text{dog}}^2 - \bar{\sigma}_{\text{inf}}^2 = \frac{(A(n+1) + B)((1-q)\sigma_\varepsilon^2 + \sigma_\eta^2)}{(1-q)(B(2n(1-q) - 1) - A(n+1))} > 0$$

which implies $\bar{\sigma}_{\text{inf}}^2 < \bar{\sigma}_{\text{dog}}^2$.

Finally,

$$\frac{\partial \bar{\sigma}_{\text{inf}}^2}{\partial q} = \frac{\sigma_\eta^2}{(1-q)^2} > 0,$$

and, whenever $\bar{\sigma}_{\text{dog}}^2 < \infty$,

$$\frac{\partial \bar{\sigma}_{\text{dog}}^2}{\partial q} = \frac{2Bn((A(n+1) + B)\sigma_\varepsilon^2 + 2Bn\sigma_\eta^2)}{(B(2n(1-q) - 1) - A(n+1))^2} > 0.$$

This shows that both $\bar{\sigma}_{\text{inf}}^2$ and $\bar{\sigma}_{\text{dog}}^2$ are increasing in q . This completes the proof. ■

Proof of Proposition 5. We first prove the result for DeGroot social learning. For $p_\theta = 0$ and $\sigma_b^2 = \sigma_c^2 = 0$, each student's first action has variance $\frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2}$. The covariance between first actions of students of different instructors is zero, whereas for students of the same instructor it equals

$$\frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2}.$$

Under DeGroot learning, the posterior belief, and thus the second action, is the average of

the student's own first action and the n peer actions. Since this average contains $n + 1$ actions, of which $n_1 + 1$ come from students of the same instructor, the planner's loss is

$$\mathcal{L}_D(T) = \left(A + \frac{B}{n+1} \right) \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + B \frac{n_1(n_1 + 1)}{(n+1)^2} \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2}.$$

Differentiating with respect to T ,

$$\mathcal{L}'_D(T) = \left(A + \frac{B}{n+1} \right) \frac{\sigma_\eta^4}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2} - 2B \frac{n_1(n_1 + 1)}{(n+1)^2} \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^3}.$$

The unconstrained optimum is therefore

$$T_D^u = \left[\frac{2Bn_1(n_1 + 1)}{(n+1)(A(n+1) + B)} - 1 \right] \sigma_\varepsilon^2 - \sigma_\eta^2,$$

with $T_D^* = \max\{T_D^u, 0\}$. Thus, T_D^* is weakly increasing in B/A , in n_1 , in σ_ε^2 , because $T_D^u > 0$ would imply that the coefficient at σ_ε^2 is positive, and weakly decreasing in σ_η^2 . Lastly, if $A = 0$, then holding n_1/n fixed it is also weakly increasing in n .

We now turn to the Bayesian case. The Bayesian posterior variance after observing the student's own action, n_1 same-instructor peer actions, and n_2 different-instructor peer actions is

$$\frac{\frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} \left[\frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + n_1 \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2} \right]}{(n+1) \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + n_1 n_2 \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2}}.$$

This follows by inverting the covariance matrix of observed actions: the student's own action and the n_1 same-instructor peer actions form a block of size $n_1 + 1$, while the n_2 actions of students with different instructors are independent of that block and of each other. Therefore

the planner's objective is

$$\mathcal{L}_B(T) = A \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + B \frac{\frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} \left[\frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + n_1 \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2} \right]}{(n+1) \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T)}{\sigma_\varepsilon^2 + T + \sigma_\eta^2} + n_1 n_2 \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T + \sigma_\eta^2)^2}}.$$

If we denote $D = \sigma_\varepsilon^2 + T$. and differentiate $\mathcal{L}_B(T)$ by T then, after multiplying by a positive denominator, we will see that the sign of $\mathcal{L}'_B(T)$ is the sign of the fourth degree polynomial

$$\begin{aligned} P(D) = & (A(n+1)^2 + B(n+1)) D^4 \\ & + [2(A(n+1)^2 + B(n+1)) \sigma_\eta^2 - 2Bn_1(n_1+1)\sigma_\varepsilon^2] D^3 \\ & + \left[(A(n+1)^2 + B(n+1)) \sigma_\eta^4 \right. \\ & \quad \left. + 2n_1\sigma_\varepsilon^2\sigma_\eta^2(A(n+1)n_2 + B(n_2 - n_1 - 1)) \right] D^2 \\ & + 2(A(n+1) + B) n_1 n_2 \sigma_\varepsilon^2 \sigma_\eta^4 D \\ & + n_1^2 n_2 (An_2 + B) \sigma_\varepsilon^4 \sigma_\eta^4. \end{aligned}$$

Note that the zeroes of $\mathcal{L}'_B(T)$ are zeroes of $P(D)$ with $D \geq \sigma_\varepsilon^2$. Notice that the coefficients at D^4 , D , and the constant term are nonnegative, with the coefficient at D^4 being strictly positive, and only the coefficients at D^3 and D^2 can be negative. Hence, the sequence of coefficients of P has at most two sign changes, and by Descartes' rule of signs, P has at most two positive roots. Therefore, $\mathcal{L}'_B(T)$ has at most two zeros on $T \geq 0$. Moreover, since the coefficient at D^4 is positive, at the largest zero the derivative changes its sign from negative to positive, hence it is a minimum. We therefore have two possibilities: either a unique local minimum, or two local minima at $T = 0$ and the larger of the roots of $\mathcal{L}'_B(T)$.

Let D^* denote the larger positive root of $P(D)$, when this root exists and corresponds to a strict local minimum, so that $P_D(D^*) > 0$. We first consider the comparative static with respect to B/A (the other results are proved similarly). Normalize the planner's objective

by A :

$$\frac{\mathcal{L}_B(T)}{A} = \frac{\sigma_\eta^2 D}{D + \sigma_\eta^2} + \frac{B}{A} V_B(D),$$

where $V_B(D)$ is the Bayesian posterior variance after social learning. The first-order condition at D^* is

$$\frac{\sigma_\eta^4}{(D^* + \sigma_\eta^2)^2} + \frac{B}{A} V'_B(D^*) = 0.$$

Since D^* is a strict local minimum, the derivative of the left-hand side with respect to D is positive. The first-order condition implies $V'_B(D^*) < 0$, which implies that its left-hand side is decreasing in B/A . The implicit function theorem now implies that D^* is increasing in B/A .

It remains to check that if for some B/A , the planner's objective is minimized both at $T = 0$ and at $D^* - \sigma_\varepsilon^2$, then an increase in B/A switches the optimum to the latter. Notice that at such a tie,

$$\frac{\sigma_\eta^2 D^*}{D^* + \sigma_\eta^2} + \frac{B}{A} V_B(D^*) = \frac{\sigma_\eta^2 \sigma_\varepsilon^2}{\sigma_\varepsilon^2 + \sigma_\eta^2} + \frac{B}{A} V_B(\sigma_\varepsilon^2).$$

Since the first term of the left-hand side is strictly increasing in D^* , it must be that $V_B(D^*) < V_B(\sigma_\varepsilon^2)$. By the envelope theorem, the derivative of the value difference between the positive local optimum and the corner one with respect to B/A is

$$V_B(D^*) - V_B(\sigma_\varepsilon^2) < 0.$$

An increase in B/A up from the value with a tie breaks the tie in favor of the positive local optimum. But we proved that this positive local optimum is increasing in B/A ; this establishes that the planner's optimal set is increasing in B/A in the strong set order.

The remaining comparative statics are shown analogously. This completes the proof. ■

Proof of Proposition 6. Let V_t denote the total variance of the teacher-side belief

inherited by generation t . Thus, in the no-common-shocks benchmark, $V_t = \sigma_{\varepsilon t}^2$, and with common shocks, $V_t = \sigma_{bt}^2 + \sigma_{\varepsilon t}^2$.

Consider first the no-common-shocks benchmark, with $\sigma_{b0}^2 = 0$ and $\sigma_c^2 = 0$. Since no common component is present initially and outside signals have no common component, no common component is ever created. Hence $\sigma_{bt}^2 = 0$ for all t , and the state of the dynamic system is fully captured by one variable $V_t = \sigma_{\varepsilon t}^2$.

In each generation, $T = 0$ is feasible. Under $T = 0$, students observe the inherited teacher-side information without additional imprecision and also receive fresh outside information. Hence their final posterior variance is strictly below V_t . Moreover, the optimal policy cannot induce a higher final posterior variance than the no-cap policy: if it did, setting $T = 0$ would improve first-stage learning and also improve final learning. Therefore $V_{t+1} < V_t$ whenever $V_t > 0$.

To show convergence, note that under $T = 0$, each generation combines inherited teacher-side information with fresh outside signal with variance σ_η^2 . Even ignoring peer learning, the posterior variance after individual learning is $\frac{V_t \sigma_\eta^2}{V_t + \sigma_\eta^2}$. Peer learning can only add information, so the final variance inherited by the next generation satisfies $V_{t+1} \leq \frac{V_t \sigma_\eta^2}{V_t + \sigma_\eta^2}$. This is equivalent to

$$\frac{1}{V_{t+1}} \geq \frac{1}{V_t} + \frac{1}{\sigma_\eta^2},$$

which implies $V_t \rightarrow 0$. Since the optimal policy yields final variance no larger than choosing $T = 0$, full learning holds under the optimal sequence $\{T_t^*\}$.

Now consider the path of the precision cap in the no-common-shocks benchmark. If $q = 0$, students learn from different-instructor peers, so peer actions are conditionally independent. The static formula therefore implies $T_t^* = 0$ for every t . If $q = 1$, the static formula applies in each generation after replacing σ_ε^2 by $V_t = \sigma_{\varepsilon t}^2$. Since the optimal cap is weakly increasing in σ_ε^2 in the static problem, and since V_t is decreasing over time, T_t^* is weakly decreasing.

Moreover, the static formula implies that $T_t^* = 0$ whenever V_t is sufficiently small. Since $V_t \rightarrow 0$, this occurs after finitely many generations.

Now consider common shocks. If the outside common component is persistent across generations, then each generation observes outside information distorted by the same unobserved component. Aggregation can eliminate idiosyncratic errors, but it cannot separate θ from a persistent common distortion that is never separately observed. Hence the limiting posterior variance is bounded away from zero, and full learning does not occur.

Finally, suppose outside common shocks are transient and independent across generations. Under $T = 0$, each generation combines inherited teacher beliefs with fresh outside information. Although the outside common shock is shared by students within a generation, it is independent across generations and therefore provides new information over time. The same recursive argument as above implies $V_t = \sigma_{bt}^2 + \sigma_{\varepsilon t}^2 \rightarrow 0$, provided that $T = 0$ is always chosen. Since $T = 0$ is feasible, the optimal policy yields a weakly lower posterior variance. Therefore full learning is achieved under the optimal policy.

It remains to show that T_t^* reaches zero in finite time. From the proof of Proposition 1, we know that the unconstrained optimum is given by (A1). Thus, if $\sigma_b^2 + \sigma_\varepsilon^2$ is sufficiently small, T_u is negative, and optimal $T^* = 0$. Applying this argument to the dynamic setting, since $\sigma_{bt}^2 + \sigma_{\varepsilon t}^2 \rightarrow 0$, there is a finite t starting from which is $T_t^* = \max\{T_u, 0\} = 0$. The fact that the path of T_t^* need not be monotone in this case is shown in Example A4. This completes the proof. ■

Example A1. This example shows that a binding precision cap can be optimal even if students, when learning from peers, use their teacher-side and outside signals to make a Bayesian inference (i.e., if they are fully Bayesian with perfect recall of the individual learning stage).

Suppose $p_\theta = 0$, $A = 0$, $B = 1$, $q = 0$, and $n = 10$. Thus the planner cares only about

the post-social-learning action, and each student observes the first actions of ten students taught by different instructors. Let

$$\sigma_b^2 = \frac{1}{2}, \quad \sigma_\varepsilon^2 = \frac{1}{5}, \quad \sigma_c^2 = 0, \quad \sigma_\eta^2 = 1.$$

For a given T , write

$$x = \sigma_b^2 + \sigma_\varepsilon^2 + T = T + \frac{7}{10}, \quad z = \sigma_\eta^2 = 1.$$

A peer's first action is the Bayesian posterior mean after observing her teacher-side signal y_k and outside signal z_k . Hence

$$u_k = \alpha y_k + (1 - \alpha) z_k, \quad \alpha = \frac{1}{T + \frac{17}{10}}.$$

Now consider student s . Unlike in the compressed-judgment benchmark, assume that at the social-learning stage she retains both of her own signals (y_s, z_s) . She observes the ten peer actions $u_1, \dots, u_{10}$. By exchangeability, their average is a sufficient statistic for the peer-action component of her information, so her second-stage information is

$$(y_s, z_s, \bar{u}), \quad \bar{u} = \frac{1}{10} \sum_{k=1}^{10} u_k.$$

The covariance matrix of the error vector $(y_s - \theta, z_s - \theta, \bar{u} - \theta)$ is

$$\Sigma(T) = \begin{pmatrix} T + \frac{7}{10} & 0 & \frac{5}{10T+17} \\ 0 & 1 & 0 \\ \frac{5}{10T+17} & 0 & \frac{100T^2+240T+569}{10(10T+17)^2} \end{pmatrix}.$$

Therefore the posterior variance of θ given $(y_s, z_s, \bar{u})$ is

$$V(T) = \text{Var}(\theta \mid y_s, z_s, \bar{u}) = (\mathbf{1}'\Sigma(T)^{-1}\mathbf{1})^{-1}.$$

A direct calculation gives

$$V(T) = \frac{1000T^3 + 3100T^2 + 7370T + 1483}{11000T^3 + 45100T^2 + 52470T + 10403}.$$

Differentiating,

$$V'(T) = \frac{100(110000T^4 - 572000T^3 - 1874600T^2 - 692680T - 11429)}{(11000T^3 + 45100T^2 + 52470T + 10403)^2}.$$

In particular, $V'(0) < 0$. Thus, starting from no cap, a small increase in T strictly reduces the student's posterior variance after social learning. The derivative has a unique positive root, $T^* \simeq 7.5635$, and this root is the global minimizer on $T \geq 0$. Moreover,

$$V(0) = \frac{1483}{10403} \simeq 0.1426, \quad V(T^*) \simeq 0.0861, \quad \lim_{T \rightarrow \infty} V(T) = \frac{1}{11} \simeq 0.0909.$$

Thus the optimal positive cap reduces the posterior variance by about

$$\frac{V(0) - V(T^*)}{V(0)} \simeq 39.6\%.$$

This shows that $T^* > 0$ may be optimal even when the student uses (y_s, z_s) as independent signals in the peer-learning stage.

Example A2. This example shows that, under the planner's optimal precision cap, students' first actions may be less precise than instructors' information, even though students become more informed than instructors after social learning.

Consider the special case with no common shocks: $\sigma_b^2 = \sigma_c^2 = 0$. Let

$$q = 1, \quad A = 0, \quad B = 1, \quad n = 3, \quad \sigma_\varepsilon^2 = 1, \quad \sigma_\eta^2 = 2.$$

Thus, all students are in echo chambers, the planner cares only about the second action, and the instructor's posterior variance is $\sigma_\varepsilon^2 = 1$.

Using the expression for the planner's optimum, we obtain $T^* = 3$. The variance of a student's first action is therefore

$$\sigma_l^2 = \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T^*)}{\sigma_\varepsilon^2 + T^* + \sigma_\eta^2} = \frac{2(1 + 3)}{1 + 3 + 2} = \frac{4}{3}.$$

Hence the student's first action is less precise than the instructor's information: $\sigma_l^2 = \frac{4}{3} > 1 = \sigma_\varepsilon^2$.

However, after social learning, the student's posterior variance is

$$\sigma_{m,E}^2 = \frac{1}{n+1} \frac{\sigma_\eta^2(\sigma_\varepsilon^2 + T^*)}{\sigma_\varepsilon^2 + T^* + \sigma_\eta^2} + \frac{n}{n+1} \frac{\sigma_\eta^4 \sigma_\varepsilon^2}{(\sigma_\varepsilon^2 + T^* + \sigma_\eta^2)^2}.$$

Substituting the same parameter values gives $\sigma_{m,E}^2 = \frac{5}{12}$. This means $\sigma_{m,E}^2 = \frac{5}{12} < 1 = \sigma_\varepsilon^2$. In other words, even though in this example students' first actions are noisier than instructors' information, they become more informed than instructors after social learning.

Example A3. Take

$$\sigma_\varepsilon^2 = \sigma_\eta^2 = 1, \quad n_1 = 3, \quad n_2 = 1, \quad B = 1, \quad A \approx 0.06320858.$$

In this case, the Bayesian objective has two global minimizers: $T = 0$ and $T \approx 0.8490$. There is also an interior critical point at approximately $T \simeq 0.2151$, which is a local maximum.

Example A4. Consider the case $q = 0$, so students learn from different-instructor peers, and suppose outside common shocks are transient and i.i.d. across generations. Let

$$\sigma_{b0}^2 = 0, \quad \sigma_{\varepsilon 0}^2 = \frac{1}{2}, \quad \sigma_c^2 = \frac{1}{2}, \quad \sigma_\eta^2 = \frac{1}{5},$$

with

$$n = 30, \quad A = \frac{1}{1000}, \quad B = 1.$$

Under the optimal policy, the precision cap is initially nonbinding, then becomes increasingly binding for two generations, and only after that starts to relax:

$$T_0^* = 0, \quad T_1^* \simeq 0.0088, \quad T_2^* \simeq 0.0103, \quad T_3^* \simeq 0.0070,$$

$$T_4^* \simeq 0.0045, \quad T_5^* \simeq 0.0025, \quad T_6^* \simeq 0.0008, \quad T_7^* = 0.$$

At the same time, total teacher-side variance falls monotonically. It starts at $\sigma_{b0}^2 + \sigma_{\varepsilon 0}^2 = 0.5000$, and in generations 1 through 5 it is approximately 0.0934, 0.0747, 0.0652, 0.0579, 0.0521. The non-monotonicity of T_t^* comes from the changing composition of teacher error. In the first generation, social learning substantially reduces instructor-specific dispersion but imports some common outside distortion. This temporarily raises the value of a binding precision cap, even though total teacher-side error is falling. Once the common component begins to fall as well, the optimal cap relaxes and eventually becomes nonbinding.