

Diversity in Teams^{*}

Miaomiao Dong[†] Tatiana Mayskaya[‡] Vladimir Smirnov[§]
Olivia Taylor[¶] Andrew Wait^{||}

November 20, 2024

Abstract

We investigate how optimal cognitive diversity depends on the nature of production and the objective employed. When the output of the most productive worker becomes relatively more important, with a utilitarian objective, optimal diversity weakly increases. However, with a Rawlsian objective, which captures the preferences of a public service provider, a regulated firm, or a very risk averse manager, optimal diversity might decrease. Moreover, optimal diversity is weakly higher with the Rawlsian objective than with the utilitarian one, which suggests that high diversity in teams can be achieved by imposing a minimal satisfactory service requirement rather than micromanaging diversity directly.

Keywords: cognitive diversity, submodular, supermodular, teams, utilitarian, Rawlsian objective.

JEL classifications: L23, M14.

^{*}We are grateful to Mikhail Drugov, Toomas Hinnosaar, Anja Prummer, Mariya Teteryatnikova, Marta Troya-Martinez and Jörgen Weibull for insightful comments. We also thank the seminar participants at HSE, University of Sydney, the 22nd SAET conference 2023, Lancaster Game Theory Conference 2023, the 2023 NZMSG meeting, Asia-Pacific Industrial Organization Conference 2023, the 41st Australasian Economic Theory Workshop 2024 and the 22nd Annual International Industrial Organization Conference 2024.

[†]Pennsylvania State University, Department of Economics, e-mail: miaomiao.dong@psu.edu.

[‡]International College of Economics and Finance / Faculty of Economic Sciences, HSE University, Moscow, Russia, e-mail: tmayskaya@gmail.com.

[§]School of Economics, Social Sciences Building A02, University of Sydney, NSW 2006, Australia, e-mail: vladimir.smirnov@sydney.edu.au.

[¶]School of Economics, Social Sciences Building A02, University of Sydney, NSW 2006, Australia, e-mail: otay5294@sydney.edu.au.

^{||}School of Economics, Social Sciences Building A02, University of Sydney, NSW 2006, Australia, e-mail: andrew.wait@sydney.edu.au.

Contents

1	Introduction	3
2	Model	7
3	Utilitarian objective	11
4	Rawlsian objective	13
5	Comparison of optimal diversity under different objectives	18
6	Concluding comments	23
A	Proofs	28
A.1	Proof of Lemma 1	28
A.2	Proof of Theorem 1 (and Remark 1)	29
A.3	Proof of Proposition 1	30
A.4	Proof of Theorem 2 (and Remark 2)	31
A.5	Proof of Proposition 2	33
A.6	Proof of Theorem 3	35
A.7	Proof of Remark 4	36
A.8	Proof of Theorem 4	38
B	Generalized model	38
B.1	Setup	38
B.2	Utilitarian objective	40
B.3	Rawlsian objective	41
B.4	Comparison of optimal diversity under different objectives	42
B.5	Proofs	43
B.5.1	Example 1	43
B.5.2	Proof of Theorem 1*	45
B.5.3	Proof of Theorem 2* (and Remark 2*)	47
B.5.4	Proof of Theorem 3*	50

1 Introduction

Since the time of Adam Smith, economists have recognized the efficiency gains from specialization and the division of labor (Smith, 1776). Driven no doubt by the complexity of knowledge work (Homer-Dixon, 2000; Jones, 2009), teams are increasingly made up of members with a range of expertise (Guimera et al., 2005; Wuchty et al., 2007). However, the empirical research on the performance of diverse teams is mixed (Ozgen, 2021). While some studies have found diverse teams are more creative and innovative (Østergaard et al., 2011; Nathan and Lee, 2013; Parrotta et al., 2014), others suggest that diverse teams will have more conflicts or a lower level of coordination (Lyons, 2017). Examining academic papers in economics, Ductor and Prummer (2023) found no clear benefit to collaborate with an author of a different gender. For corporate boards, Frijns et al. (2016) found that more culturally diverse boards performed worse; in contrast, Kim and Starks (2016) argue that women bring a different set of skills to boards, enhancing firm value.¹ The documented differences in the impact of diversity suggest the need to understand factors that affect the value of diversity. In this paper, we focus on two of such factors — the objective function of a team and the production process.

Existing literature focuses on a utilitarian objective, which aims to maximize the total surplus. In many circumstances, however, the objective function of real teams is different. In the healthcare sector, for instance, the objective can include that every patient receives some minimum threshold of treatment. A regulated firm might be required to provide a minimum quality of service to every client.² Moreover, a manager of a for-profit firm might be extremely risk averse and, thus, will act to maximize the payoff for the worst-affected client. With these examples in mind, in addition to the utilitarian objective, we also analyze a Rawlsian objective, which maximizes the welfare of the worst-off client.

We show that optimal diversity of expertise in a team is weakly higher with the Rawlsian objective than it is with the utilitarian objective. This theoretical result leads to empirically relevant predictions. When a government or a regulator imposes a minimal satisfactory service requirement on a pro-profit firm, our model predicts that optimal diversity of specializations/services within the firm may increase. In contrast, the reverse may happen when a government-owned company is privatized. These predictions are consistent with for-profit hospitals focusing on low-cost (more profitable) patients (Cheng et al., 2015). A similar pat-

¹Others suggest that the link between diversity of corporate boards and financial performance is subject to endogeneity issues (Knyazeva et al., 2021).

²For example, Telstra's Universal Service Obligation (USO) states that they must “ensure standard telephone services (STS) and payphones are reasonably accessible to all people in Australia on an equitable basis, wherever they work or live.” Similarly, Australia Post has its community service obligations set out in Section 27 of the Australian Postal Corporation Act 1989, which aim to ensure a minimum service standard for all including those in remote and regional areas.

tern has been observed following the privatization of previously government-run bus services in New South Wales, Australia, where companies focused more on profits and densely populated areas rather than service quality for high-cost travellers on the fringe (Haylen, 2023).

The interconnection between the objective and optimal diversity provides an opportunity for managers and policy makers to exploit. If managers want greater diversity of services/skills in their organizations, they can achieve it by incentivizing their employees to value more the contentment of the least satisfied client. Similarly, policy makers can achieve greater diversity of services by adjusting the provider's objective to include a minimal satisfactory service requirement. The possibility to affect diversity through objective can be useful in many situations, as a principal typically has greater observability and control over broader objectives of a unit rather than details of its operations within.

Another factor that affects the value of diversity is the nature of production. Take the example of a team of chefs who own a restaurant.³ Diversity will be useful when chefs are coming up with the menu because it will help them generate a range of innovative ideas. However, once the recipe is decided and the grill is fired up, homogeneity will help the team coordinate more efficiently.⁴ In a different context, a synchronized swimming team with more similar skills would tend to perform better. Alternatively, in a math Olympiad teams with more diverse backgrounds will more likely solve all problems and win. These examples suggest that the nature of production (that is, the way individual contributions are aggregated) will influence the optimal level of diversity. Prat (2002) argues that teams with similar skills perform better if jobs are strategic complements (that is, team production is supermodular in the contributions of team members), while hiring people with different backgrounds is optimal if their jobs are strategic substitutes (that is, team production is submodular in the contributions of team members).

Of course, in reality, the degree of complementarity between tasks varies for different production processes. This raises the question as to how the results in Prat (2002) translate to more complex production environments; that is, as workers' tasks become less complementary, does optimal diversity in their skills increase? We show that this might not always be the case.

Taking firstly a utilitarian approach, we connect the results for pure substitutes and complements, allowing for the gamut of production possibilities between these two cases. As it turns out, as the relative importance of the most productive team member increases (production exhibits higher submodularity), the optimal level of diversity of the team also increases. That is, when the production function has greater weighting on the submodular component,

³This is inspired by an example outlined by Page (2008).

⁴Other examples like this abound. In a rock band, the songwriting and recording process can benefit from creative ideas drawn from different genres, but after the record has been completed, a tight live performance is aided by a similar musical approach.

optimal diversity increases. In contrast, when the production process is more supermodular, meaning that the relative importance of the contribution of the weakest team member increases, the optimal level of team diversity falls. This result implies that teams in which one individual is sufficient to complete a given task will have a higher level of optimal diversity than teams in which all members are critical. In doing so, we generalize the results of [Prat \(2002\)](#) for any degree of strategic complementarity between team members.

In contrast, we show that the Rawlsian objective may reverse the comparative statics. To illustrate our point, consider the following example. A hospital has two nurses who can perform CPR and its management has to decide where to locate each nurse along a lengthy corridor. A good CPR requires a quick initial response from a single nurse, followed by a quick response from a second nurse (who can take over to ensure that the first nurse is not fatigued and that high-quality chest compressions are delivered). If the hospital locates the nurses close to the ends of the corridor, then all patients have a similar chance of surviving: patients near the ends receive a very quick response from the first nurse but the second nurse will be very slow to come, while patients near the center receive relatively timely responses from both nurses. Now suppose that the nurses are equipped with a defibrillator. The addition of this technology means that the nurses become less complementary because a single nurse can have a greater impact when performing CPR. In this case, to ensure that the care of the worst-off patient is maximized, the hospital locates the nurses at $1/4$ and $3/4$ along the corridor. Indeed, with this positioning, a single nurse can timely reach any patient in their half of the corridor. This example shows that the optimal distance between the nurses can fall as they become less complementary. The distance between the nurses reflects the difference, or the level of diversity, in nurses' ability to perform CPR for a given patient. Thus, the example demonstrates that, contrary to the spirit of [Prat \(2002\)](#), optimal diversity may decrease as jobs become less complementary.

Our framework also allows us to investigate how optimal diversity changes as team members become more narrowly specialized. Within the utilitarian framework, our model predicts that teams with narrow specialists (such as an economist and a lawyer) will be less diverse than teams with generalists (such as MBA-trained managers). However, this relationship could be inverted with the Rawlsian objective. In addition, we show that an increase in diversity, caused by a transition from the utilitarian to Rawlsian objective, is higher for teams that are made up of less specialized members.

Related literature

The advantages of diversity when solving problems have been popularized by a series of articles and books by Lu Hong and Scott Page ([Hong and Page, 2001, 2004; Page, 2008](#)). This work, as well as the body of literature that has followed ([LiCalzi and Surucu, 2012](#)), focuses on the

benefits of diversity for tasks where only one agent is strictly necessary to solve a given problem (that is, submodular tasks). Taking this task type as a starting point, these papers study the benefits of diversity through knowledge spillovers between agents. However, there is limited recognition of the importance of the underlying task type to the results.

An alternative approach in both economics and social psychology such as [Steiner \(1972\)](#) emphasizes how optimal team composition is influenced by the way the skills of team members are aggregated. Central to this literature is that the degree of complementarity or substitutability between workers' skills determines optimal team structure. When workers' contributions are complementary, one agent performing better will raise the marginal product obtained from the other agent performing better. When contributions are substitutable, one agent performing better will lower the marginal product obtained from the other performing better.⁵

The literature on "skill dispersion" or "talent disparity" in teams focuses on whether workers of similar or different abilities should be matched together. [Rosen \(1981\)](#) suggests that when it is only the performance of a few "superstar" workers that matters (that is, in jobs with submodular production functions) there will be a dispersed skill distribution in the workforce. On the other hand, [Kremer \(1993\)](#) argues that in situations when there is an "O-ring" production function in that overall performance depends on the weakest link (supermodular production function), an optimal team will have workers of similar ability.

Whilst our approach has strong similarities to the literature on "talent disparity," we focus on cognitive diversity rather than ability diversity, meaning that workers are of equal ability, but their strengths lie in different "cognitive" areas.⁶ In this regard, the closest paper to ours, and a significant inspiration for us, is [Prat \(2002\)](#). He applies the "team theory" of [Marschak and Radner \(1972\)](#) to study whether team members should have homogeneous or heterogeneous backgrounds (in terms of their area of expertise). He finds that, when agents' actions are supermodular, homogeneity is optimal. Conversely, when agents' actions are submodular, heterogeneity is optimal.⁷ As noted above, we build on the work of [Prat \(2002\)](#) in two ways. First, using a distant measure of diversity and allowing the production function to vary in a continuous way from supermodular to submodular, we investigate how optimal diversity changes as the production function becomes less complementary. Second, in addition to the utilitarian

⁵For more discussion of supermodularity and submodularity see, for example, [Milgrom and Roberts \(1990\)](#). Most recently, [Glover and Kim \(2023\)](#) investigate how the degree of complementarity or substitutability between workers' skills determines the optimal contract for these workers. In contrast to us, they keep the team composition unchanged but allow for transfers.

⁶Cognitive diversity arises largely from experience or training. For example, an economist and an engineer have different skills and approach tasks differently given their training. It is of course possible that identity diversity will shape opportunities ([Page, 2008](#)).

⁷Similarly, [Jensen \(2020\)](#) finds that the impact of team diversity on outcomes depends on the magnitude of the production function's elasticity of substitution.

objective, we look at the Rawlsian objective, which yields qualitatively different results.

There is a parallel literature that focuses on how team composition influences incentives (Franco et al., 2011; Kaya and Vereshchagina, 2014; Bel et al., 2015; Glover and Kim, 2021). In these papers, the team-choice problem typically boils down to a trade-off between inducing greater effort in equilibrium in a diverse team or having a more productive team when members are more similar to each other. Central to these results is the assumption that there is always a productivity gain from homogeneity. In contrast, our paper here abstracts from incentive considerations. Moreover, we assume that diversity can increase productivity for some tasks, like problem-solving.⁸

The rest of the paper is organized as follows. Section 2 describes the setup. Sections 3 and 4 present the comparative statics results for the utilitarian and the Rawlsian objectives, respectively. Section 5 compares optimal diversity under the utilitarian and under Rawlsian objectives. Section 6 concludes.

2 Model

A principal (“she”) must hire a team of two agents (both “he”) to perform a task.

Individual output

Each agent $i = 1, 2$ has a particular cognitive type or specialty, modelled as a point on a unit interval, $t_i \in [0, 1]$. The task is also modelled as a point on a unit interval, $t \in [0, 1]$. The distance $d_i = |t - t_i|$ between the task and agent i ’s cognitive type determines his individual output on that task. Specifically, the closer the agent’s cognitive type to the location of the task, the higher the agent’s performance on that task. More precisely, agent i ’s output on a given task is represented by a function $v(d_i)$ with the following properties.

Assumption 1. *Function v is continuous, differentiable and strictly decreasing: $v'(d) < 0$ for all $d \in [0, 1]$.*

Assumption 2. *Function v is twice differentiable and strictly concave: $v''(d) < 0$ for all $d \in [0, 1]$.*

Figure 1 depicts agents’ individual outputs, described by function $v(d) = -d^2$, which satisfies Assumptions 1 and 2.

⁸We abstract from the general equilibrium issue of matching in teams in a market. Specifically, we assume that agents of any type are available at zero cost, with the only restriction on the number of hired agents. We also effectively assume zero communication costs between diverse agents, unlike in Dong and Mayskaya (2024).

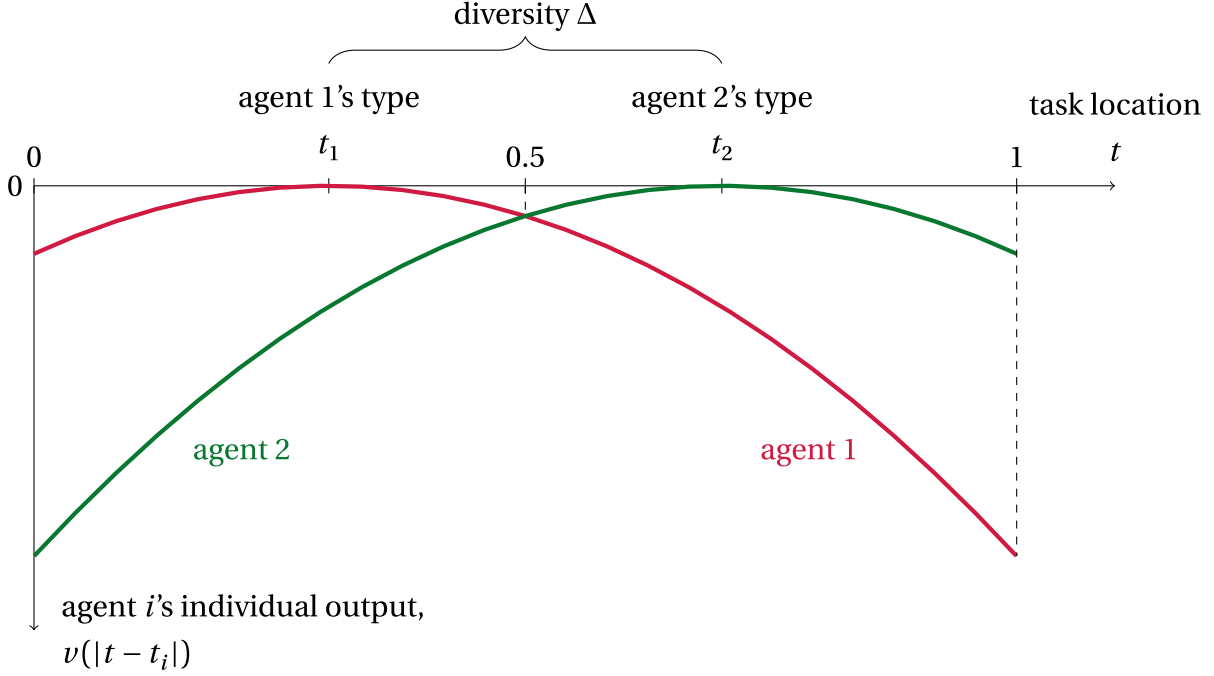


Figure 1: Agents' individual outputs, $v(d) = -d^2$, $t_1 = 0.3$, $t_2 = 0.7$

The decreasing function (Assumption 1) captures that the agent's individual output is maximized when the task matches exactly his area of specialty ($d = 0$), while the assumption of concavity (Assumption 2) ensures that there is a unique interior optimal diversity under the Rawlsian objective defined below.

Diversity

Without loss of generality, let $t_1 \leq t_2$. We define **diversity** Δ as the distance between the agents' types, $\Delta = t_2 - t_1$. Since the agents' types belong to $[0, 1]$, Δ can take any value from $[0, 1]$. For example, in Figure 1, $\Delta = 0.4$.

Timeline

First, the principal chooses t_1 and t_2 ; in this way, the principal chooses the team's level of diversity Δ . Second, the required task t materializes according to the uniform distribution on the support $[0, 1]$.⁹ Third, agents work on the task and team output is realized.¹⁰

⁹Most of our results are robust to the uniform distribution assumption. The results sensitive to this assumption are highlighted in footnotes 13 and 16.

¹⁰The agents are non-strategic players in the game.

Team output

We make the following assumption about the team output.¹¹

Assumption 3. *The team output is given by*

$$\omega(x_1, x_2) = \beta \max\{x_1, x_2\} + (1 - \beta) \min\{x_1, x_2\}, \quad (1)$$

where $x_i = v(|t - t_i|)$ is agent i 's individual output.

The team output (1) features two task types (the terminology is borrowed from [Steiner \(1972\)](#)).

- With a **disjunctive task**, the output depends only on the contribution of the more productive or stronger agent. In this task, agents' contributions are strategic *substitutes*. In other words, this task is *submodular* with respect to the contributions of the agents.¹² For example, only one student in a group working on an assignment needs to find the solution for the whole study group to be successful.
- With a **conjunctive task**, the output depends only on the contribution of the less productive or weaker agent. In this task, agents' contributions are strategic *complements*. In other words, this task is *supermodular* with respect to the contributions of the agents. This is an example of the o-ring production in [Kremer \(1993\)](#); if one individual fails, the product is faulty.

Parameter $\beta \in [0, 1]$ captures the relative importance of the conjunctive and disjunctive tasks. As β grows from 0 to 1, the task changes continuously between conjunctive (supermodular production function) to disjunctive (submodular production function). We can think that a low- β task t requires significant input from both agents. The team output for such task is close to the minimum output of the agents, see case $\beta = 0.2$ in [Figure 2](#). In contrast, a successful completion of a high- β task depends largely on the output of the best suited agent. The team output for such task is close to the maximum output of the agents, see case $\beta = 0.9$ in [Figure 2](#).

If $\beta = 0.5$, then the production is both supermodular and submodular in that output depends only on the sum of the agents' contributions. Following [Steiner \(1972\)](#), we refer to this case as an **additive task**. In this task, the agents' contributions are strategically independent. An example could be a team of people fishing, each with their own rod. If they are all far enough apart from one another so that there is no externality between them, the team output is the sum of their individual catches.

¹¹We generalize Assumption 3 in [Appendix B](#).

¹²We elaborate more on the connection between substitutability/complementarity and sub/supermodularity in [Appendix B](#) where we generalize Assumption 3.

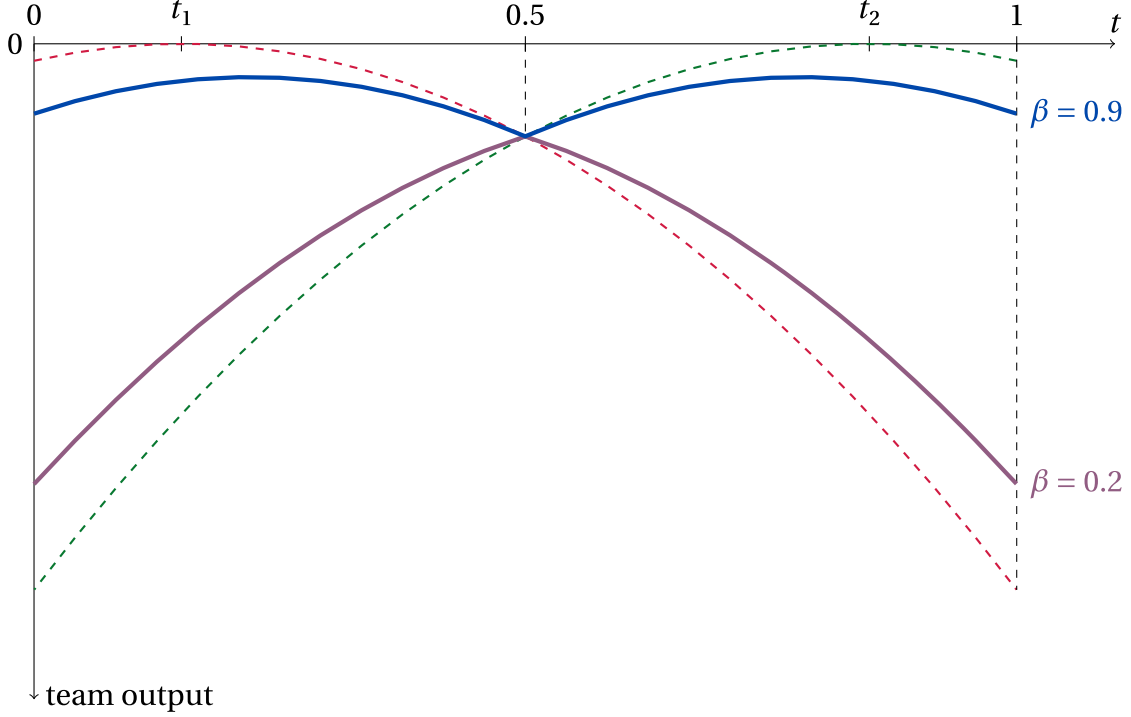


Figure 2: Team output (1) for different β , $v(d) = -d^2$, $t_1 = 0.15$, $t_2 = 0.85$

Principal's objective

The principal chooses t_1 and t_2 to maximize a certain objective. We consider two different objectives: the **utilitarian objective**, defined as the expected value of the team output,

$$U = \int_0^1 \omega(v(|t - t_1|), v(|t - t_2|)) dt; \quad (2)$$

and the **Rawlsian objective**, defined as the minimal value of the team output among all tasks,

$$R = \min_{t \in [0,1]} \omega(v(|t - t_1|), v(|t - t_2|)). \quad (3)$$

Lemma 1 ensures that without loss of generality we can focus on the cognitive types which are symmetric around 0.5 (as in Figure 1).

Lemma 1. *With Assumptions 1 and 3, for any fixed diversity level Δ , $t_1 = (1 - \Delta)/2$ and $t_2 = (1 + \Delta)/2$ maximize both the utilitarian and the Rawlsian objectives.*

Proof. See Appendix A.1. □

Thus, we can rewrite the principal's objectives (2) and (3) as functions of Δ , denoted as $U(\Delta)$ and $R(\Delta)$, respectively. We refer to **optimal diversity** Δ_U^* and Δ_R^* as the values of Δ that maximize the utilitarian and the Rawlsian objectives, respectively.

Our goal is to investigate how optimal diversity changes with β and compare the results for different objectives. In addition to that, we also look at the special case when the agents' individual output function $v(d)$ is given by $v(d) = -d^\alpha$, and investigate how optimal diversity changes with α .

3 Utilitarian objective

Comparative statics with respect to β

Theorem 1 shows how optimal diversity changes with the weight on the contribution of the stronger team member (β) under the utilitarian objective. If the team output heavily relies on the contribution of both team members ($\beta \leq 0.5$), a homogeneous workforce is optimal ($\Delta_U^* = 0$). However, as team output depends more on the stronger contributor (as β grows from 0.5 to 1), an increasingly diverse workforce should be employed.

Theorem 1. *With Assumptions 1 and 3, under the utilitarian objective, optimal diversity is unique for all $\beta \in [0, 1]$ and equal to $\Delta_U^* = 0$ for $\beta \in [0, 0.5]$; for $\beta \in [0.5, 1]$, optimal diversity $\Delta_U^*(\beta)$ strictly increases from 0 to 0.5.*

Proof. See Appendix A.2. □

To explain the intuition of the results in Theorem 1, first notice that (a) the weaker agent's expected output is always falling with diversity, while (b) that of the stronger agent is increasing in diversity for all $\Delta \in (0, 0.5)$ and maximized at $\Delta = 0.5$. Statement (a) is trivial since the weaker agent's output is falling with diversity for every task $t \in [0, 1]$. Statement (b) is less obvious. The output from the stronger agent is increasing with diversity for tasks close to the corners, $t \in [0, t_1) \cup (t_2, 1]$, and decreasing for tasks close to the center, $t \in [t_1, t_2]$. As diversity increases, the losses from $t \in [t_1, t_2]$ grow and the gains from $t \in [0, t_1) \cup (t_2, 1]$ fall, and they are balanced at $\Delta = 0.5$ when the sets of tasks that benefit from diversity, $[0, t_1) \cup (t_2, 1]$, and those that suffer from diversity, $[t_1, t_2]$, have equal measure.¹³

It is intuitive that the minimal diversity $\Delta_U^* = 0$ is optimal for $\beta \in [0, 0.5]$, that is, if the team output depends more on the performance of the weaker agent than on the stronger contributor. Any gains in contribution from the stronger agent are more than matched by the losses from the weaker agent, given the higher weighting on the weaker agent in the production function. As output from the weaker agent is always falling with diversity (statement (a)), it is optimal to select $\Delta_U^* = 0$. This result echoes Prat (2002) who finds that if production is supermod-

¹³Under the utilitarian objective, the value of optimal diversity at $\beta = 1$ depends on the distribution of the required task t . This value may be higher or lower than 0.5 if the distribution of t is different from uniform.

ular ($\beta \leq 0.5$) one possible optimal team is homogenous, which corresponds to $\Delta_U^* = 0$ in our model. We strengthen his result by proving that optimal diversity is unique in our setting.

As the submodular component becomes relatively more important — that is, if $\beta \in (0.5, 1]$ — it is optimal to have a diverse team, $\Delta_U^* > 0$. Moreover, according to Theorem 1, optimal diversity Δ_U^* increases from 0 to 0.5 as β increases from 0.5 to 1. Intuitively, with increasing weight on the maximum output, team success depends more on the contribution of the stronger agent. Optimal diversity is a weighted average of the diversity level that maximizes the contribution of the weaker agent (supermodular component), i.e., $\Delta = 0$ by statement (a), and the one that maximizes the contribution of the stronger agent (submodular component), i.e., $\Delta = 0.5$ by statement (b). As the weight on the submodular component, β , increases from 0.5 to 1, optimal diversity increases from $\Delta_U^* = 0$ to $\Delta_U^* = 0.5$. This result extends the result in Prat (2002) for submodular production functions ($\beta > 0.5$ in our setting), according to which at least one optimal team is diverse (corresponding to $\Delta_U^* > 0$ in our model). In contrast, in our setting, optimal team composition is unique, which allows us to derive comparative statics for different submodular production functions — that is, for different $\beta > 0.5$.

The logic for $\beta \in (0.5, 1]$ is illustrated more formally through Remark 1.

Remark 1. In Appendix A.2, in the proof of Theorem 1, we show that when $\beta \in [0.5, 1]$, optimal diversity uniquely solves

$$\beta \left(v \left(\frac{\Delta_U^*}{2} \right) - v \left(\frac{1 - \Delta_U^*}{2} \right) \right) = (1 - \beta) \left(v \left(\frac{\Delta_U^*}{2} \right) - v \left(\frac{1 + \Delta_U^*}{2} \right) \right). \quad (4)$$

The difference $v \left(\frac{\Delta_U^*}{2} \right) - v \left(\frac{1 - \Delta_U^*}{2} \right)$ on the left-hand side of (4) is the marginal benefit from increasing diversity when the task is disjunctive ($\beta = 1$). It is equal to the difference in the performance of the stronger agent when the task is at the center ($t = 0.5$) and when the task is at the corner ($t = 0$ or $t = 1$). Similarly, the difference $v \left(\frac{\Delta_U^*}{2} \right) - v \left(\frac{1 + \Delta_U^*}{2} \right)$ on the right-hand side of (4) is the marginal benefit from decreasing diversity when the task is conjunctive ($\beta = 0$). Thus, equation (4) equates marginal benefits from increasing and decreasing diversity, with respective weights. A higher weight on the disjunctive task increases the weight on the marginal benefit from increasing diversity, thus raising optimal diversity.

Comparative statics with respect to α

Figure 3 depicts optimal diversity for several functions of the form $v(d) = -d^\alpha$. For all functions, when the contribution of the worse performing team member is relatively important — $\beta \in [0, 0.5]$ — there is no diversity within the team. On the other hand, for each production technology shown, as the weight β on the output of the stronger agent increases from 0.5 to 1, optimal diversity increases to $\Delta^* = 0.5$.

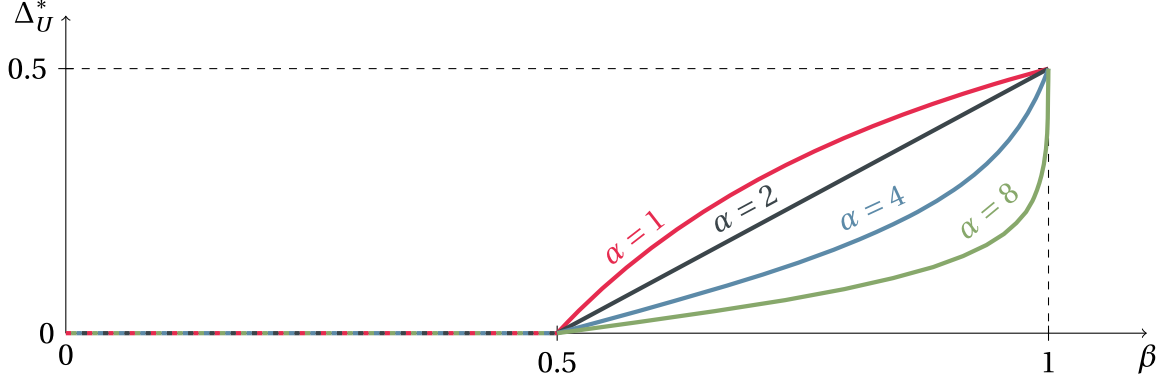


Figure 3: Optimal diversity under the utilitarian objective, $v(d) = -d^\alpha$ for $\alpha = 1, 2, 4, 8$.

According to Figure 3, optimal diversity has lower values in the region $\beta \in (0.5, 1)$ for more concave individual output functions, that is, for functions with higher α . Proposition 1 formally proves this observation for any $\alpha > 0$.

Proposition 1. *With Assumption 3, under the utilitarian objective, for individual output function $v(d) = -d^\alpha$, for any $\beta \in (0.5, 1)$, optimal diversity Δ_U^* is decreasing in $\alpha > 0$.*

Proof. See Appendix A.3. □

The result in Proposition 1 is interesting because it illuminates the role of specialization. We could think of the agents as being more narrowly specialized with increasing concavity of their individual output function. Intuitively, greater concavity implies that the value the agents add to a task falls at a greater rate with the distance from their area of specialty, t_i . This means that they have a greater relative advantage at t_i . Moving the agents apart now causes a greater loss from the contribution of the weaker agent. So, the agents with more concave $v(d)$ cannot afford to move apart at the same rate as their less concave counterparts. Hence, there will be a lower level of optimal diversity in teams with specialists (e.g., surgeons, legal specialists, academic researchers) than in teams with generalists (e.g., management consulting, academic deans).

4 Rawlsian objective

Some organizations are concerned about the quality of the service delivered to *all* clients. This sort of objective could be imposed by a regulator or as part of the service contract. To explore these issues, we now turn our attention to the Rawlsian objective.

Comparative statics with respect to β

Theorem 2 shows how optimal diversity changes with the weight on the contribution of the stronger team member (β) under the Rawlsian objective. For $\beta \leq 0.5$, the Rawlsian objective leads to the same optimal diversity as the utilitarian objective — $\Delta_R^* = \Delta_U^* = 0$. For $\beta > 0.5$, in contrast to the utilitarian objective, the Rawlsian objective delivers qualitatively different comparative statics. Specifically, as β grows from 0.5 to 1, optimal diversity changes non-monotonically.¹⁴

Theorem 2. *With Assumptions 1 to 3, under the Rawlsian objective, optimal diversity is unique for all $\beta \in [0, 1]$, equal to $\Delta_R^* = 0$ for $\beta \in [0, 0.5]$ and equal to $\Delta_R^* = 0.5$ for $\beta = 1$; for $\beta \in [0.5, 1]$, there exists $\beta^* \in (0.5, 1)$ such that optimal diversity $\Delta_R^*(\beta)$ strictly increases for $\beta \in (0.5, \beta^*)$ and strictly decreases for $\beta \in (\beta^*, 1)$.*

Proof. See Appendix A.4. □

The rest of the section discusses the intuition behind the non-monotonic relationship featured in Theorem 2. To help frame this discussion, the following remark highlights that the worst-case team output is achieved either on the edges or at the center.

Remark 2. In Appendix A.4, in the proof of Theorem 2, we show that the Rawlsian objective can be written as

$$R = \min \left\{ \beta v \left(\frac{1-\Delta}{2} \right) + (1-\beta) v \left(\frac{1+\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right\}, \quad (5)$$

where the first term is the team output for the corner tasks ($t = 0$ and $t = 1$) and the second term is the team output for the task at the center ($t = 1/2$).

Indeed, as illustrated in Figure 2, for low β , the team output achieves its minimum at $t = 0$ and $t = 1$. Intuitively, low β requires input from both agents and task $t = 0$ ($t = 1$) is the most challenging for agent 2 (agent 1). According to Figure 2, for high β , if diversity $\Delta = t_2 - t_1$ is sufficiently large, the team output achieves its minimum at $t = 1/2$. This observation relies on the concavity assumption (Assumption 2). Intuitively, larger diversity lowers the team output for the middle tasks, $t \in [t_1, t_2]$. When v is concave, among all middle tasks $t \in [t_1, t_2]$, the minimum team output is achieved at the center $t = 1/2$.

Optimal diversity maximizes (5), the minimum between the team output for the corner tasks and for the center task. In Figure 4, the blue curve illustrates the team output for the center task, $v(\frac{\Delta}{2})$. The red curves show the team output for the corner tasks for different β , $\beta v(\frac{1-\Delta}{2}) + (1-\beta)v(\frac{1+\Delta}{2})$. As indicated by the declining blue curve, the team output for the center task is maximized with zero diversity; indeed, for task $t = 0.5$, the best output is when

¹⁴In Appendix B.3, we discuss sufficient conditions under which the Rawlsian objective, within a generalized model, can produce monotone optimal diversity function $\Delta_R^*(\beta)$.

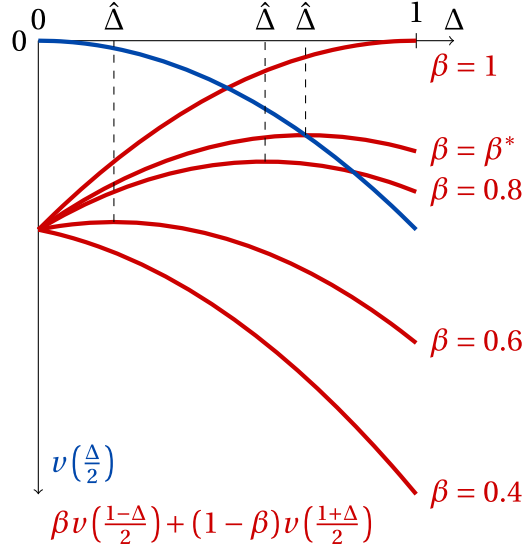


Figure 4: Team output, $v(d) = -d^2$. The red graphs correspond to the corner tasks ($t = 0$ and $t = 1$). The blue graph corresponds to the center task ($t = 1/2$).

$t_1 = t_2 = 0.5$. Thus, any possible incentive for the principal to choose a positive level of agent diversity comes from a consideration for the tasks at the corners.

For a task located at either periphery (0 or 1), an increase in diversity means that the agent who is closer to the task gets even closer, while the other agent moves further away. Thus, the individual output increases for the stronger agent and decreases for the weaker agent. The combined effect of an increase in diversity on the team output depends on weight β on the stronger agent output.

If $\beta \leq 0.5$, then the output of the weaker agent gets more weight. In this case, as diversity increases, the combined return falls, making the team performance on the corner tasks worse off. As shown in Figure 4 for $\beta = 0.4$, the red curve falls with diversity. Thus, when the agents' types are complements ($\beta \leq 0.5$), there is no conflict between the corners and the center, so optimal diversity is zero.

If $\beta > 0.5$, then the output of the stronger agent gets more weight. In this case, the team output for a corner task will be maximized at some positive diversity — in Figure 4, each red curve for $\beta > 0.5$ is maximized at some $\hat{\Delta} > 0$. As the tasks become less complementary (i.e., β increases), diversity $\hat{\Delta}$ that maximizes the team output for corner tasks starts increasing because the output of the weaker agent becomes less relevant for the team output and, thus, the agents should be located closer to the corners to increase the output of the stronger agent.

For values of β slightly greater than 0.5, such as $\beta = 0.6$ and $\beta = 0.8$ illustrated in Figure 4, when diversity is equal to $\hat{\Delta}$, the team output at the corners is still lower than in the center. Consequently, the principal, trying to maximize the worst-case output, will choose the level of

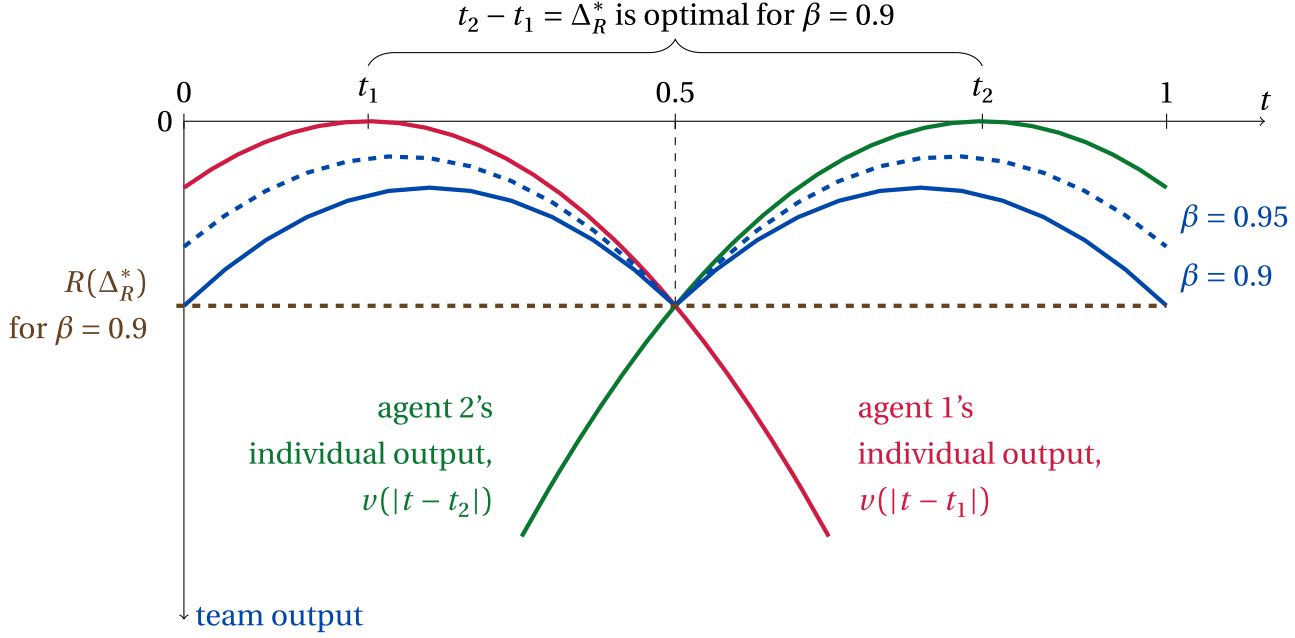


Figure 5: Team output (1) for different β , $v(d) = -d^2$. The agents' types t_1 and t_2 are chosen optimally for $\beta = 0.9$.

diversity which is the best for the tasks at the extremes and, therefore, opt for $\Delta_R^* = \hat{\Delta}$. In this case, as the tasks become less complementary (i.e., as β increases), the principal increases diversity to increase the performance of the closer agent for the corner tasks.

As β increases beyond some $\beta = \beta^*$, when diversity is equal to $\hat{\Delta}$, the team output at the corners is higher than in the center. Thus, the principal faces the corners/center trade-off: optimal diversity equalizes the team output of the agents at the corner tasks and at the center task. In these cases, the principal will choose Δ_R^* where the red and blue lines intersect.

Figure 5 illustrates that for $\beta > \beta^*$, to balance the corners/center team output, the principal decreases diversity as β increases. In the figure, diversity $\Delta = t_2 - t_1$ is fixed at the optimal level for $\beta = 0.9$. The blue solid curve corresponds to the team output for $\beta = 0.9$. In line with the discussion above, the team output achieves its minimum both at the corner tasks and at the center task. Now suppose that weight β on the stronger agent increases to 0.95. Then, the team output raises to the blue dashed curve, so that it becomes closer to the individual output of the stronger agent, which is the maximum of the red and green graphs. Figure 5 demonstrates that as β changes from 0.9 to 0.95, the team output at the corner tasks increases, while the team output at the center task does not change. To reestablish the balance between the corner tasks and the center task, the principal has to increase the team output at the center task. For that, she decreases diversity.

Let us apply the above logic to the illustrative example from the introduction. The addition

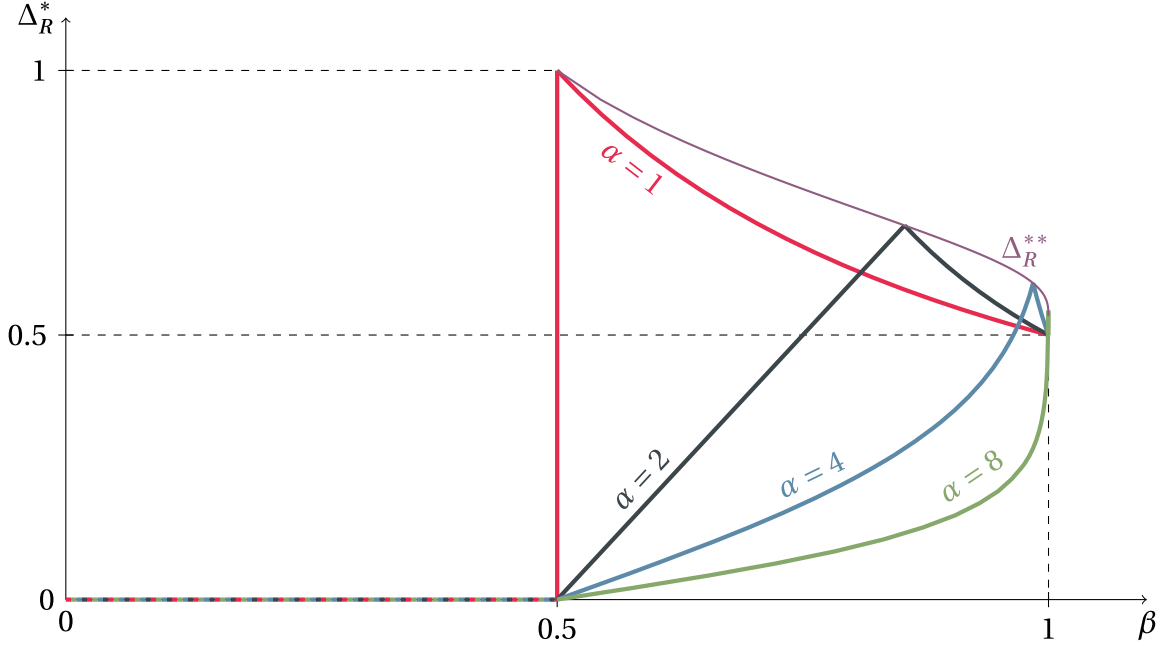


Figure 6: Optimal diversity under the Rawlsian objective, $v(d) = -d^\alpha$ for $\alpha = 1, 2, 4, 8$. Function Δ_R^{**} marks the maximum diversity under each $\alpha > 1$, see also Figure 9.

of a defibrillator does not change the quality of treatment for the patients near the center because to them, both nurses come at the same time. In contrast, the patients near the ends of the corridor receive better treatment because the nurse who comes first is now better equipped to deliver high-quality CPR for a long time. To ensure that the patients near the center and near the ends get treatment of similar quality, the hospital should increase the quality of treatment for the patients near the center by locating the nurses closer to these patients.

Comparative statics with respect to α

Figure 6 illustrates optimal diversity for $v(d) = -d^\alpha$ with different levels of α . In accordance with Theorem 2, for $\alpha > 1$ where v is strictly concave, optimal diversity is equal to 0 for $\beta \in [0, 0.5]$ and has a hump-shaped form on $\beta \in [0.5, 1]$, with a peak at some $\beta = \beta^*$.

The Rawlsian objective could invert the relationship between specialization and diversity, derived in Proposition 1 for the utilitarian objective. According to Proposition 2 (ii) and in line with Figure 6, before the peak — that is, for $\beta \in [0.5, \beta^*]$ — the effect from increasing α is the same as with the utilitarian objective: a more concave individual output function implies a lower level of optimal diversity. However, as Proposition 2 (iii) states and Figure 6 illustrates, after the peak — that is, for $\beta \in [\beta^*, 1]$ — the relationship is reversed.

Proposition 2. *With Assumption 3, under the Rawlsian objective, for individual output func-*

tion $v(d) = -d^\alpha$, for any $\alpha_2 > \alpha_1 > 1$,

(i) $\beta^*(\alpha_2) > \beta^*(\alpha_1)$;

(ii) if $0.5 < \beta < \beta^*(\alpha_1)$, then $\Delta_R^*(\beta, \alpha_2) < \Delta_R^*(\beta, \alpha_1)$;

(iii) if $\beta^*(\alpha_2) < \beta < 1$, then $\Delta_R^*(\beta, \alpha_2) > \Delta_R^*(\beta, \alpha_1)$.

Proof. See Appendix A.5. □

The intuition behind the reversed relationship in Proposition 2 (iii) is based on the corners/center trade-off, which shapes the optimal level of diversity for $\beta \in [\beta^*, 1]$. As concavity increases (that is, the agents become more narrowly specialized), the agents' individual output for both the corners and the center decreases. However, the team output for the corner tasks suffers more because it is based almost entirely on the output of one agent (since β is high), while a drop in individual performance for the center task is partially compensated by the fact that both agents equally contribute to the team output. Thus, for a sufficiently submodular production function, under the Rawlsian objective, there will be a higher level of optimal diversity in teams where specialists are hired than in teams where generalists are hired.

Proposition 2 (i) states that the value of β corresponding to the peak of the optimal diversity function, β^* , increases with α . The intuition for that is also based on the above observation that as concavity increases, the team output for the corner tasks suffers more than for the center task. This means that there is a broader range of parameters for which the middle is not the worst-case outcome.

Remark 3. Optimal diversity is unique in Theorem 2 because of the strict concavity of v (Assumption 2). This does not hold in Figure 6, as $v(d) = -d^\alpha$ is linear for $\alpha = 1$. In this case, for $\beta = 0.5$, any level of diversity is optimal.

5 Comparison of optimal diversity under different objectives

In practice, some policies can cause the principal to change her objective. For example, a minimal satisfactory service requirement could impose a Rawlsian-like objective on a firm. Likewise, privatization of a government-owned corporation could shift its objective towards the utilitarian. This section shows that under some mild condition, the Rawlsian objective unambiguously increases diversity, and more so for less specialized teams.

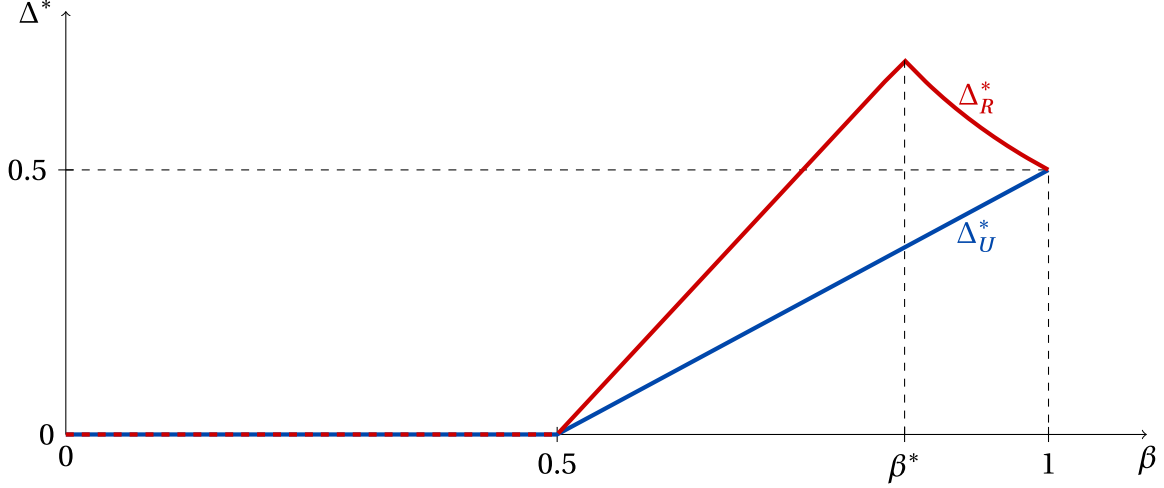


Figure 7: Optimal diversity, $v(d) = -d^2$. Blue graph corresponds to the utilitarian objective; red graph corresponds to the Rawlsian objective.

Comparison for a fixed β

By Theorems 1 and 2, optimal diversity does not change if the team output function is super-modular ($\beta \leq 0.5$) or the task is disjunctive ($\beta = 1$). Thus, for the rest of the section, we focus on $\beta \in (0.5, 1)$.

To present our result, we impose the following assumption.

Assumption 4. $A(d) = \frac{v''(d)}{v'(d)}$ is strictly decreasing in $d \in [0, 1]$.

Assumption 4 captures the idea that as task t moves from the center $1/2$ to a corner 0 , the benefit from increasing diversity grows. This benefit is measured by the ratio of a marginal gain from the increased contribution of the stronger agent, $\frac{\partial}{\partial \Delta} v\left(\frac{1-\Delta}{2} - t\right)$, over a marginal loss from the decreased contribution of the weaker agent, $-\frac{\partial}{\partial \Delta} v\left(\frac{1+\Delta}{2} - t\right)$. Assumption 4 implies that this ratio is decreasing in $t \in [0, 1/2]$:

$$\frac{\partial}{\partial t} \left(\frac{\frac{\partial}{\partial \Delta} v\left(\frac{1-\Delta}{2} - t\right)}{-\frac{\partial}{\partial \Delta} v\left(\frac{1+\Delta}{2} - t\right)} \right) = \frac{v'\left(\frac{1-\Delta}{2} - t\right)}{v'\left(\frac{1+\Delta}{2} - t\right)} \left(A\left(\frac{1+\Delta}{2} - t\right) - A\left(\frac{1-\Delta}{2} - t\right) \right) < 0. \quad (6)$$

Assumption 4 is a relatively mild condition. To see this, notice that since

$$A'(d) = \frac{v'(d)v'''(d) - v''(d)^2}{v'(d)^2}, \quad (7)$$

for a decreasing strictly concave function v , a sufficient but not necessary condition for Assumption 4 to hold is $v'''(d) \geq 0$.¹⁵

¹⁵In general, for three times differentiable decreasing function v , Assumption 4 requires v''' being not too neg-

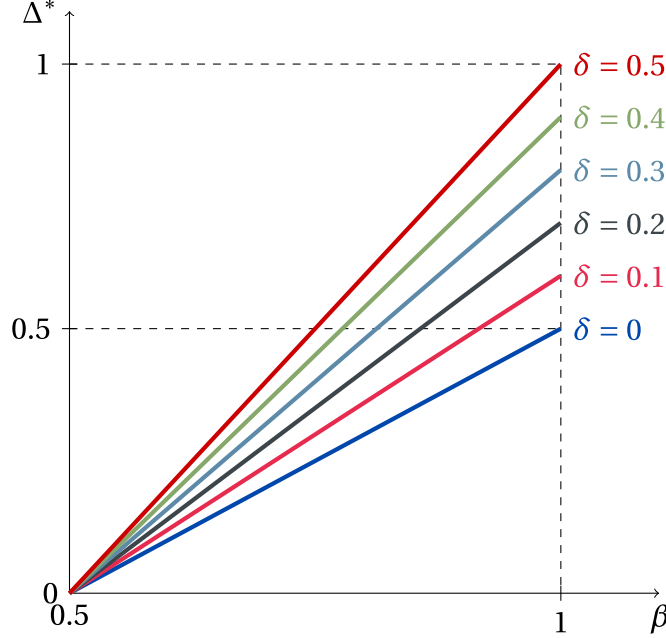


Figure 8: Optimal diversity under the δ -objective, $v(d) = -d^2$.

Theorem 3 shows that Assumption 4 is sufficient to ensure that a change in the objective from utilitarian to Rawlsian unambiguously increases diversity.

Theorem 3. *With Assumptions 1 to 4, for $\beta \in (0.5, 1)$, optimal diversity under the Rawlsian objective is always higher than under the utilitarian one: $\Delta_R^*(\beta) > \Delta_U^*(\beta)$.*

Proof. See Appendix A.6. □

Figure 7 illustrates Theorem 3. For $\beta \in (\beta^*, 1)$, Assumption 4 is not necessary, and the statement of Theorem 3 trivially follows from Theorems 1 and 2: the Rawlsian objective gives diversity greater than 0.5, while the utilitarian objective gives diversity lower than 0.5.¹⁶

To see the intuition behind Theorem 3 when $\beta \in (0.5, \beta^*)$, let us construct a parametrized family of objectives, with the utilitarian and the Rawlsian objectives at the extremes.¹⁷ The utilitarian objective averages the team outcomes across all tasks, while the Rawlsian objective focuses on the corner tasks (since $\beta \in (0.5, \beta^*)$). To connect the objectives, we remove the middle tasks. For any $\delta \in (0, 0.5)$, a δ -objective is defined as the expected team output for the tasks distributed uniformly on the support $[0, 0.5 - \delta] \cup [0.5 + \delta, 1]$. As δ increases, the focus of

ative. Most decreasing concave functions satisfy this requirement. For example, $v(d) = -d^\alpha$ satisfies Assumptions 1, 2 and 4 for all $\alpha > 1$.

¹⁶Theorem 3 can be generalized to any distribution for the required task, provided that for the utilitarian objective optimal diversity is lower than 0.5. This condition holds when, for example, the distribution density of t is symmetric around the center and weakly decreases with distance from the center.

¹⁷The proof of Theorem 3 in Appendix A.6 follows the same logic.

the principal becomes more concentrated on the tasks closer to the corners; as δ approaches 0 (0.5), the δ -objective becomes the utilitarian (Rawlsian) objective.

Figure 8 illustrates that for $\beta \in (0.5, 1)$, optimal diversity increases in δ . This result relies on Assumption 4. Intuitively, the corner tasks feature a higher difference in performance between the stronger and the weaker agents than the middle tasks. Hence, as the principal's focus becomes more concentrated on the tasks closer to the corners, the performance of the stronger agent becomes more important for her. To increase the performance of the stronger agent for the corner tasks, the principal moves the agents closer to the corners, thus increasing diversity.

Remark 4. Assumption 4 is sufficient, but not necessary, for Theorem 3. See Appendix A.7 for details.

Comparative statics with respect to α

The sensitivity of diversity to the choice of objective depends on team specialization. The change in the objective from utilitarian to Rawlsian increases diversity more if the team is comprised of less specialized agents.

To incorporate team specialization into the model, take $v(d) = -d^\alpha$, where parameter α captures the scope of specialization of each team member. Lower α describes a less specialized team.

Suppose that under each objective, the principal selects the weight β that maximizes optimal diversity. Under the utilitarian objective, by Theorem 1, the principal selects $\beta = 1$ and achieves diversity $\Delta_U^{**} = 0.5$. Under the Rawlsian objective, by Theorem 2, the principal selects $\beta = \beta^*(\alpha)$ and achieves diversity $\Delta_R^{**}(\alpha) = \Delta_R^*(\beta^*(\alpha), \alpha) > 0.5$. Theorem 4 shows that an increase in diversity, caused by a transition from the utilitarian to Rawlsian objective, falls with α .

Theorem 4. *With Assumption 3, for individual output function $v(d) = -d^\alpha$, the difference $\Delta_R^{**}(\alpha) - \Delta_U^{**}$ decreases with $\alpha > 1$.*

Proof. See Appendix A.8. □

Figures 6 and 9 illustrate Theorem 4 and show that under the Rawlsian objective, the peak of the diversity function falls with α .

Intuitively, for a less specialized team, the performance of the weaker agent (say, agent 2) on the worst-case task (task $t = 0$) is higher, while the performance of the stronger agent (agent 1) on this task is lower. Hence, the principal increases diversity trading off a lower performance of the weaker agent for a higher performance of the stronger agent.

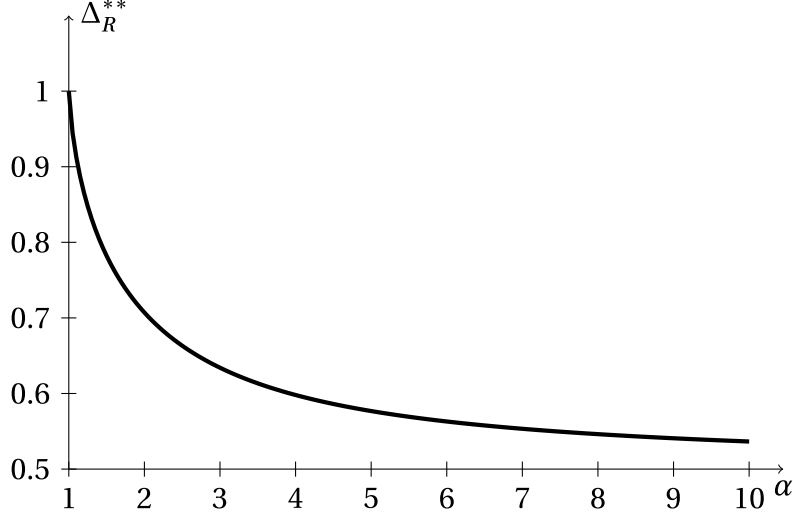


Figure 9: The highest optimal diversity under the Rawlsian objective, $\Delta_R^{**}(\alpha)$.

Another way to see that an increase in diversity, caused by a transition from the utilitarian to Rawlsian objective, falls with α is to look at the *expected* diversity. Assuming that β is distributed uniformly on $[0, 1]$, denote by $\Delta_U^*(\alpha)$ and $\Delta_R^*(\alpha)$ the expected optimal diversity under the utilitarian and Rawlsian objectives, respectively.

By Theorem 3, a transition from the utilitarian to Rawlsian objective increases the optimal diversity for $\beta \in (0.5, 1)$ and does not change it for the remaining β . Hence, the expected optimal diversity always increases, that is, $\Delta_R^*(\alpha) > \Delta_U^*(\alpha)$ for all α . In fact, as Figure 10 illustrates, the increase in diversity can be more than two times when moving to the Rawlsian objective from the utilitarian one. Moreover, this difference falls with α .¹⁸

Observation. With Assumption 3, for individual output function $v(d) = -d^\alpha$, an increase in the expected optimal diversity, caused by a transition from the utilitarian to Rawlsian objective, decreases with $\alpha > 1$; that is, $\Delta_R^*(\alpha) - \Delta_U^*(\alpha)$ and $\Delta_R^*(\alpha)/\Delta_U^*(\alpha)$ are both decreasing on $\alpha > 1$.

Intuitively, an increase in the concavity of the individual output function lowers the team output for the corner tasks more than for any other task. Hence, an increase in α affects the utilitarian objective, which averages the team output over all tasks, less than the Rawlsian objective, which in the optimum is equal to the team output for the corner tasks. To increase the team output for a corner task, the principal moves the weaker agent closer to this task, thereby decreasing diversity. Since the Rawlsian objective is more sensitive to α , $\Delta_R^*(\alpha)$ decreases more than $\Delta_U^*(\alpha)$. Thus, since $\Delta_R^*(\alpha) > \Delta_U^*(\alpha)$, an increase in α lowers the gap between $\Delta_R^*(\alpha)$ and $\Delta_U^*(\alpha)$. In other words, the expected optimal diversity is less (more) sensitive to the change in objective for more (less) specialized teams.

¹⁸A formal proof of this statement is complicated and, therefore, omitted.

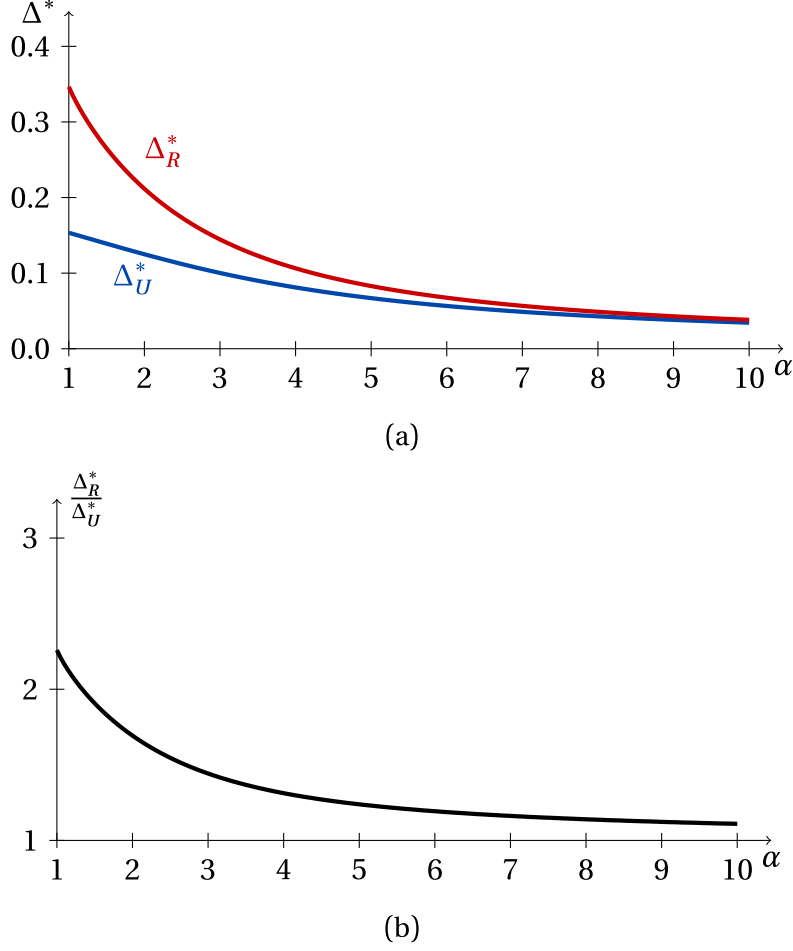


Figure 10: The expected optimal diversity, $\Delta_U^*(\alpha)$ and $\Delta_R^*(\alpha)$.

6 Concluding comments

Diversity, Equity and Inclusion (DEI) policies are now commonplace in businesses as well as in government and not-for-profit institutions around the world. The University of Sydney, for example, states in its Strategic Plan that “[w]e embrace equity, diversity and inclusion as core to our success, and our students and staff reflect the communities we serve.” Researchers such as [Hong and Page \(2001, 2004\)](#); [Page \(2008\)](#) argue that there is a diversity dividend; in complex environments, excellence requires a diversity of perspectives, meaning that teams beat individuals and that a diverse team will outperform a homogenous one.

In this paper, we try to unpick the drivers behind these arguments, analyzing the optimal cognitive diversity in a team. Firstly, the production process matters. If a team is pursuing a purely utilitarian objective, as the output of the most productive team member matters relatively more — that is, as the contributions of the team members become less complementary — the team’s optimal diversity weakly increases.

Secondly, the objective of the team will affect the optimal level of diversity and how this level of diversity responds to a change in the production process. In many situations, the team is not just concerned about maximizing the expected outcome. Rather, it could be that a bigger weight is placed on the worst outcome. We show that under the Rawlsian objective, which corresponds to maximizing the worst outcome, the optimal level of diversity may decrease as the contributions of the team members become less complementary. This result comes in contrast with the utilitarian objective, under which optimal diversity always weakly increases.

Finally, we show that optimal diversity in a team is always higher with the Rawlsian rather than the utilitarian objective, and the difference is bigger for teams with less specialized agents. This result has potentially important implications for the modern organization and the DEI policies noted above. It suggests that establishing Rawlsian-type incentives for teams or team leaders can induce greater diversity. Given that the setting of overall objectives is often less complicated and easier to monitor for senior managers than micro-managing hiring and team balance, this is a practical insight to the more effective implementation of the diversity goals of many organizations.

References

- Bel, Roland, Vladimir Smirnov, and Andrew Wait**, “Team composition, worker effort and welfare,” *International Journal of Industrial Organization*, 7 2015, 41 (C), 1–8. 7
- Cheng, Terence C., John P. Haisken-DeNew, and Jongsay Yong**, “Cream skimming and hospital transfers in a mixed public-private system,” *Social Science & Medicine*, 5 2015, 132, 156–164. 3
- Dong, Miaomiao and Tatiana Mayskaya**, “Does reducing communication barriers promote diversity?,” Working paper, 2024. Available at SSRN: <https://www.ssrn.com/abstract=4164162>. 7
- Ductor, Lorenzo and Anja Prummer**, “Gender, homophily, collaboration, and output,” Working paper, 2023. Available at <https://static1.squarespace.com/static/622744fa1588ee23cc206c1d/t/6515a4d6727dfd28ed6fb5c4/1695917271867/HomophilyRevision.pdf>. 3
- Franco, April Mitchell, Matthew Mitchell, and Galina Vereshchagina**, “Incentives and the structure of teams,” *Journal of Economic Theory*, 11 2011, 146 (6), 2307–2332. 7
- Frijns, Bart, Olga Dodd, and Helena Cimerova**, “The impact of cultural diversity in corporate boards on firm performance,” *Journal of Corporate Finance*, 12 2016, 41, 521–541. 3

Glover, Jonathan and Eunhee Kim, “Optimal team composition: diversity to foster implicit team incentives,” *Management Science*, 9 2021, 67(9), 5800–5820. 7

Glover, Jonathan C. and Eunhee Kim, “Optimal team incentives for productively heterogeneous agents,” Working paper, 2023. Available at SSRN: <https://www.ssrn.com/abstract=4433792>. 6

Guimera, Roger, Brian Uzzi, Jarrett Spiro, and Luis A. Nunes Amaral, “Team assembly mechanisms determine collaboration network structure and team performance,” *Science*, 4 2005, 308(5722), 697–702. 3

Haylen, Jo, “Tens of thousands of privatised bus services cancelled, with over a million left stranded,” Media release by Minister for Transport, NSW Government, April 2023. Retrieved from <https://www.nsw.gov.au/media-releases/tens-of-thousands-of-privatised-bus-services-cancelled-over-a-million-left-stranded>. 4

Homer-Dixon, Thomas, *The ingenuity gap*, Vintage Canada, 2000. <https://homerdixon.com/writing/books/the-ingenuity-gap>. 3

Hong, Lu and Scott E. Page, “Problem solving by heterogeneous agents,” *Journal of Economic Theory*, 3 2001, 97(1), 123–163. 5, 23

— **and** —, “Groups of diverse problem solvers can outperform groups of high-ability problem solvers,” *Proceedings of the National Academy of Sciences*, 11 2004, 101(46), 16385–16389. 5, 23

Jensen, Martin Kaae, “On the strategic benefits of diversity,” Working paper, 2020. Available at <https://drive.google.com/file/d/10kxf-vDo64ZPVNUT2Zu1AMM9TqctaJMh/view>. 6

Jones, Benjamin F., “The burden of knowledge and the “death of the Renaissance Man”: is innovation getting harder?,” *Review of Economic Studies*, 1 2009, 76(1), 283–317. 3

Kaya, Ayça and Galina Vereshchagina, “Partnerships versus corporations: moral hazard, sorting, and ownership structure,” *American Economic Review*, 1 2014, 104(1), 291–307. 7

Kim, Daehyun and Laura T. Starks, “Gender diversity on corporate boards: do women contribute unique skills?,” *American Economic Review*, 5 2016, 106(5), 267–271. 3

Knyazeva, Anzhela, Diana Knyazeva, and Lalitha Naveen, “Diversity on corporate boards,” *Annual Review of Financial Economics*, 11 2021, 13(1), 301–320. 3

- Kremer, Michael**, “The O-ring theory of economic development,” *Quarterly Journal of Economics*, 8 1993, 108 (3), 551–575. 6, 9
- LiCalzi, Marco and Oktay Surucu**, “The power of diversity over large solution spaces,” *Management Science*, 7 2012, 58 (7), 1408–1421. 5
- Lyons, Elizabeth**, “Team production in international labor markets: experimental evidence from the field,” *American Economic Journal: Applied Economics*, 7 2017, 9 (3), 70–104. 3
- Marschak, Jacob and Roy Radner**, *Economic theory of teams*, Yale University Press, 1972. <https://trid.trb.org/view/545360>. 6
- Milgrom, Paul and John Roberts**, “The economics of modern manufacturing: technology, strategy, and organization,” *American Economic Review*, 1990, 80 (3), 511–528. 6
- Nathan, Max and Neil Lee**, “Cultural diversity, innovation, and entrepreneurship: firm-level evidence from London,” *Economic Geography*, 10 2013, 89 (4), 367–394. 3
- Østergaard, Christian R., Bram Timmermans, and Kari Kristinsson**, “Does a different view create something new? The effect of employee diversity on innovation,” *Research Policy*, 4 2011, 40 (3), 500–509. 3
- Ozgen, Ceren**, “The economics of diversity: innovation, productivity and the labour market,” *Journal of Economic Surveys*, 9 2021, 35 (4), 1168–1216. 3
- Page, Scott**, *The difference: how the power of diversity creates better groups, firms, schools, and societies*, Princeton University Press, 2008. <https://doi.org/10.1515/9781400830282>. 4, 5, 6, 23
- Parrotta, Pierpaolo, Dario Pozzoli, and Mariola Pytlikova**, “The nexus between labor diversity and firm’s innovation,” *Journal of Population Economics*, 4 2014, 27 (2), 303–364. 3
- Prat, Andrea**, “Should a team be homogeneous?,” *European Economic Review*, 7 2002, 46 (7), 1187–1207. 4, 5, 6, 11, 12, 39, 45
- Rosen, Sherwin**, “The economics of superstars,” *American Economic Review*, 1981, 71 (5), 845–858. 6
- Smith, Adam**, *An inquiry into the nature and causes of the wealth of nations*, Strahan, 1776. <https://books.google.ru/books?id=C5dNAAAAcAAJ>. 3
- Steiner, Ivan Dale**, *Group process and productivity*, Academic Press, 1972. <https://archive.org/details/groupprocessprod0000stei/mode/2up>. 6, 9

Wuchty, Stefan, Benjamin F. Jones, and Brian Uzzi, “The increasing dominance of teams in production of knowledge,” *Science*, 5 2007, 316 (5827), 1036–1039. [3](#)

A Proofs

A.1 Proof of Lemma 1

We prove Lemma 1 for the generalized model in Appendix B.

We first split each objective function into four parts:

$$U = \underbrace{\int_0^{t_1} \omega(v(t_1 - t), v(t_2 - t)) dt}_{U_{01}} + \underbrace{\int_{t_1}^{\frac{t_1+t_2}{2}} \omega(v(t - t_1), v(t_2 - t)) dt}_{U_{1m}} + \underbrace{\int_{\frac{t_1+t_2}{2}}^{t_2} \omega(v(t - t_1), v(t_2 - t)) dt}_{U_{m2}} + \underbrace{\int_{t_2}^1 \omega(v(t - t_1), v(t - t_2)) dt}_{U_{21}}, \quad (\text{A.1})$$

$$R = \min \left\{ \underbrace{\min_{t \in [0, t_1]} \omega(v(t_1 - t), v(t_2 - t))}_{R_{01}}, \underbrace{\min_{t \in [t_1, \frac{t_1+t_2}{2}]} \omega(v(t - t_1), v(t_2 - t))}_{R_{1m}}, \underbrace{\min_{t \in [\frac{t_1+t_2}{2}, t_2]} \omega(v(t - t_1), v(t_2 - t))}_{R_{m2}}, \underbrace{\min_{t \in [t_2, 1]} \omega(v(t - t_1), v(t - t_2))}_{R_{21}} \right\}. \quad (\text{A.2})$$

Applying the change of variable $\tau = t_1 + t_2 - t$ and using $\omega(x_1, x_2) = \omega(x_2, x_1)$, which is followed from Assumption 7, we get that the third part of each objective is equal to the second part:

$$U_{m2} = \int_{t_1}^{\frac{t_1+t_2}{2}} \omega(v(t_2 - \tau), v(\tau - t_1)) d\tau = U_{1m}, \quad (\text{A.3})$$

$$R_{m2} = \min_{\tau \in [t_1, \frac{t_1+t_2}{2}]} \omega(v(t_2 - \tau), v(\tau - t_1)) = R_{1m}. \quad (\text{A.4})$$

We then rewrite each of the remaining three parts so that each depends on t_1 and t_2 through the sum $t_1 + t_2 \equiv s$ and the difference $t_2 - t_1 \equiv \Delta$. Since $0 \leq t_1 \leq t_2 \leq 1$, the difference $t_2 - t_1 = \Delta$ belongs to $[0, 1]$ and, for a fixed Δ , the sum $t_1 + t_2 = s$ belongs to $[\Delta, 2 - \Delta]$.

Applying the change of variable $\tau = \frac{t_1+t_2}{2} - t$, we get that the first part becomes

$$U_{01} = \int_{\Delta/2}^{s/2} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right) d\tau, \quad R_{01} = \min_{\tau \in [\Delta/2, s/2]} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right). \quad (\text{A.5})$$

Applying the change of variable $\tau = \frac{t_2-t_1}{2} - t_1 + t$, we get that the second part of each objective depends only on t_1 and t_2 only through $t_2 - t_1 \equiv \Delta$:

$$U_{1m} = \int_{\Delta/2}^{\Delta} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\frac{3\Delta}{2} - \tau\right)\right) d\tau, \quad R_{1m} = \min_{\tau \in [\Delta/2, \Delta]} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\frac{3\Delta}{2} - \tau\right)\right). \quad (\text{A.6})$$

Applying the change of variable $\tau = t - \frac{t_1+t_2}{2}$ and using $\omega(x_1, x_2) = \omega(x_2, x_1)$, we rewrite the forth part of each objective as

$$U_{21} = \int_{\Delta/2}^{1-s/2} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right) d\tau, \quad R_{21} = \min_{\tau \in [\Delta/2, 1-s/2]} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right). \quad (\text{A.7})$$

Utilitarian objective

To prove that $s = 1$ is optimal for any fixed $\Delta \in [0, 1]$, we differentiate (A.1) with respect to s and use (A.3), (A.5), (A.6) and (A.7) to get

$$\frac{\partial U(\Delta, s)}{\partial s} = \frac{1}{2} \left\{ \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right) \Big|_{\tau=s/2} - \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right) \Big|_{\tau=1-s/2} \right\}. \quad (\text{A.8})$$

Since $\omega(x_1, x_2)$ is weakly increasing in each argument by Assumption 6 and v is strictly decreasing by Assumption 1, function $\omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right)$ is weakly decreasing in τ . Thus,

$$\frac{\partial U(\Delta, s)}{\partial s} \geq 0 \text{ for } s < 1, \quad \frac{\partial U(\Delta, s)}{\partial s} \leq 0 \text{ for } s > 1, \quad (\text{A.9})$$

which implies that $s = 1$ is optimal.

Rawlsian objective

Since

$$R(\Delta, s) = \min\{R_{01}, R_{1m}, R_{21}\} \quad (\text{A.10})$$

by (A.2) and (A.4) and since R_{1m} in (A.6) does not depend on s , to prove that $R(\Delta, s)$ is maximized at $s = 1$, it is sufficient to prove that $\min\{R_{01}, R_{21}\}$ is maximized at $s = 1$. By (A.5) and (A.7),

$$\min\{R_{01}, R_{21}\} = \min_{\tau \in [\frac{\Delta}{2}, \max\{\frac{s}{2}, 1-\frac{s}{2}\}]} \omega\left(v\left(\tau - \frac{\Delta}{2}\right), v\left(\tau + \frac{\Delta}{2}\right)\right), \quad (\text{A.11})$$

which is weakly decreasing in $\max\{\frac{s}{2}, 1-\frac{s}{2}\}$. The minimum of $\max\{\frac{s}{2}, 1-\frac{s}{2}\}$ is achieved at $s = 1$.

A.2 Proof of Theorem 1 (and Remark 1)

Rewriting the expression (A.1) using (A.3), (A.5), (A.6) and (A.7), and then substituting $s = 1$ and $\omega(x_1, x_2)$ from (1) in Assumption 3, we get the following expression for the utilitarian objective

$$U(\Delta, \beta) = 2 \int_{\Delta/2}^{1/2} \left(\beta v\left(\tau - \frac{\Delta}{2}\right) + (1-\beta) v\left(\tau + \frac{\Delta}{2}\right) \right) d\tau + 2 \int_{\Delta/2}^{\Delta} \left(\beta v\left(\tau - \frac{\Delta}{2}\right) + (1-\beta) v\left(\frac{3\Delta}{2} - \tau\right) \right) d\tau. \quad (\text{A.12})$$

Differentiating (A.12) yields

$$\frac{\partial U(\Delta, \beta)}{\partial \Delta} = \beta \left(v\left(\frac{1+\Delta}{2}\right) - v\left(\frac{1-\Delta}{2}\right) \right) + (1-2\beta) \left(v\left(\frac{1+\Delta}{2}\right) - v\left(\frac{\Delta}{2}\right) \right). \quad (\text{A.13})$$

Suppose $\beta \in [0, 0.5]$. Since function v is strictly decreasing by Assumption 1, the derivative (A.13) is negative for all $\Delta > 0$. Thus, optimal diversity is unique and equal to 0.

Suppose $\beta \in [0.5, 1]$. Rewriting the derivative (A.13) as

$$\frac{\partial U(\Delta, \beta)}{\partial \Delta} = (1 - \beta)v\left(\frac{1 + \Delta}{2}\right) - \beta v\left(\frac{1 - \Delta}{2}\right) + (2\beta - 1)v\left(\frac{\Delta}{2}\right), \quad (\text{A.14})$$

we can see that it is strictly decreasing in Δ because function v is strictly decreasing by Assumption 1. Thus, $U(\Delta, \beta)$ is strictly concave in Δ and, therefore, optimal diversity is unique. Moreover, optimal diversity is equal to 0.5 for $\beta = 1$ and belongs to the interval $(0, 0.5)$ for $\beta \in (0.5, 1)$ because

$$\frac{\partial U(0, \beta)}{\partial \Delta} = (2\beta - 1) \underbrace{\left(v(0) - v\left(\frac{1}{2}\right)\right)}_{>0}, \quad \frac{\partial U(0.5, \beta)}{\partial \Delta} = (1 - \beta) \underbrace{\left(v\left(\frac{3}{4}\right) - v\left(\frac{1}{4}\right)\right)}_{<0}, \quad (\text{A.15})$$

where the signs follow from Assumption 1. Thus, optimal diversity solves $\frac{\partial U(\Delta^*, \beta)}{\partial \Delta} = 0$, which is equivalent to (4). Optimal diversity increases in β because from (A.13),

$$\frac{\partial^2 U(\Delta, \beta)}{\partial \Delta \partial \beta} = \frac{1}{\beta} \left(\frac{\partial U(\Delta, \beta)}{\partial \Delta} + \underbrace{v\left(\frac{\Delta}{2}\right) - v\left(\frac{1 + \Delta}{2}\right)}_{>0} \right), \quad (\text{A.16})$$

which implies that the mixed derivative is positive whenever $\frac{\partial U(\Delta, \beta)}{\partial \Delta} = 0$.

A.3 Proof of Proposition 1

Substituting $v(d) = -d^\alpha$ into (A.14) yields

$$\frac{\partial U(\Delta, \alpha)}{\partial \Delta} = -(1 - \beta) \left(\frac{1 + \Delta}{2}\right)^\alpha + \beta \left(\frac{1 - \Delta}{2}\right)^\alpha - (2\beta - 1) \left(\frac{\Delta}{2}\right)^\alpha, \quad (\text{A.17})$$

where, instead of β , we now emphasize the dependence of U on α . To prove that optimal diversity Δ_U^* is decreasing in α for any $\beta \in (0.5, 1)$, since $\Delta_U^* \in (0, 0.5)$ for any $\beta \in (0.5, 1)$ by Theorem 1, it is sufficient to show that

$$\frac{\partial^2 U(\Delta, \alpha)}{\partial \Delta \partial \alpha} < 0 \quad \text{whenever} \quad \frac{\partial U(\Delta, \alpha)}{\partial \Delta} = 0 \quad \text{and} \quad 0 < \Delta < \frac{1}{2}. \quad (\text{A.18})$$

Differentiating (A.17) with respect to α gives

$$\frac{\partial^2 U(\Delta, \alpha)}{\partial \Delta \partial \alpha} = -(1 - \beta) \left(\frac{1 + \Delta}{2}\right)^\alpha \ln\left(\frac{1 + \Delta}{2}\right) + \beta \left(\frac{1 - \Delta}{2}\right)^\alpha \ln\left(\frac{1 - \Delta}{2}\right) - (2\beta - 1) \left(\frac{\Delta}{2}\right)^\alpha \ln\left(\frac{\Delta}{2}\right). \quad (\text{A.19})$$

Combining (A.17) and (A.19) so to omit β , we get

$$\frac{\partial^2 U(\Delta, \alpha)}{\partial \Delta \partial \alpha} = \frac{f_1(\Delta, \alpha)}{f_2(\Delta, \alpha)} \frac{\partial U(\Delta, \alpha)}{\partial \Delta} + \frac{f_3(\Delta, \alpha)}{f_2(\Delta, \alpha)}, \quad (\text{A.20})$$

where

$$f_1(\Delta, \alpha) = \left(\frac{1 - \Delta}{2}\right)^\alpha \ln\left(\frac{1 - \Delta}{2}\right) + \left(\frac{1 + \Delta}{2}\right)^\alpha \ln\left(\frac{1 + \Delta}{2}\right) - 2 \left(\frac{\Delta}{2}\right)^\alpha \ln\left(\frac{\Delta}{2}\right), \quad (\text{A.21})$$

$$f_2(\Delta, \alpha) = \left(\frac{1 - \Delta}{2}\right)^\alpha + \left(\frac{1 + \Delta}{2}\right)^\alpha - 2 \left(\frac{\Delta}{2}\right)^\alpha, \quad (\text{A.22})$$

$$f_3(\Delta, \alpha) = \frac{1}{\alpha} \left(\left(\frac{1+\Delta}{2} \right)^\alpha - \left(\frac{\Delta}{2} \right)^\alpha \right) \left(\left(\frac{1-\Delta}{2} \right)^\alpha - \left(\frac{\Delta}{2} \right)^\alpha \right) \left(f_4 \left(\left(\frac{1-\Delta}{\Delta} \right)^\alpha \right) - f_4 \left(\left(\frac{1+\Delta}{\Delta} \right)^\alpha \right) \right), \quad (\text{A.23})$$

$$f_4(y) = \frac{y}{y-1} \ln(y). \quad (\text{A.24})$$

Function $f_2(\Delta, \alpha)$ is positive for all $\Delta \in (0, 0.5)$ and $\alpha > 0$ because

$$\frac{\partial f_2(\Delta, \alpha)}{\partial \Delta} - \frac{\alpha}{\Delta} f_2(\Delta, \alpha) = -\frac{\alpha}{2\Delta} \left(\left(\frac{1-\Delta}{2} \right)^{\alpha-1} + \left(\frac{1+\Delta}{2} \right)^{\alpha-1} \right) < 0, \quad f_2\left(\frac{1}{2}, \alpha\right) = \left(\frac{3}{4}\right)^\alpha - \left(\frac{1}{4}\right)^\alpha > 0. \quad (\text{A.25})$$

Function $f_3(\Delta, \alpha)$ is negative for all $\Delta \in (0, 0.5)$ and $\alpha > 0$ because $f_4(y)$ is increasing:

$$f_4'(y) = \frac{y-1-\ln(y)}{(1-y)^2} > 0, \quad (\text{A.26})$$

and because

$$\left(\frac{1+\Delta}{\Delta} \right)^\alpha > \left(\frac{1-\Delta}{\Delta} \right)^\alpha, \quad \left(\frac{1+\Delta}{2} \right)^\alpha > \left(\frac{1-\Delta}{2} \right)^\alpha > \left(\frac{\Delta}{2} \right)^\alpha. \quad (\text{A.27})$$

Thus, (A.18) holds.

A.4 Proof of Theorem 2 (and Remark 2)

Rewriting the expression (A.2) using (A.4), (A.5), (A.6) and (A.7), and then substituting $s = 1$ and $\omega(x_1, x_2) = \beta \max\{x_1, x_2\} + (1-\beta) \min\{x_1, x_2\}$, we get the following expression for the Rawlsian objective

$$R(\Delta, \beta) = \min \{R_{01}(\Delta, \beta), R_{1m}(\Delta, \beta)\}, \quad (\text{A.28})$$

$$R_{01}(\Delta, \beta) = \min_{\tau \in [\frac{\Delta}{2}, \frac{1}{2}]} \left\{ \beta \max \left\{ v \left(\tau - \frac{\Delta}{2} \right), v \left(\tau + \frac{\Delta}{2} \right) \right\} + (1-\beta) \min \left\{ v \left(\tau - \frac{\Delta}{2} \right), v \left(\tau + \frac{\Delta}{2} \right) \right\} \right\}, \quad (\text{A.29})$$

$$R_{1m}(\Delta, \beta) = \min_{\tau \in [\frac{\Delta}{2}, \Delta]} \left\{ \beta \max \left\{ v \left(\tau - \frac{\Delta}{2} \right), v \left(\frac{3\Delta}{2} - \tau \right) \right\} + (1-\beta) \min \left\{ v \left(\tau - \frac{\Delta}{2} \right), v \left(\frac{3\Delta}{2} - \tau \right) \right\} \right\}. \quad (\text{A.30})$$

By Assumption 1, function v is strictly decreasing. Hence,

$$v \left(\tau - \frac{\Delta}{2} \right) \geq v \left(\tau + \frac{\Delta}{2} \right) \quad \text{for all } \tau, \quad v \left(\tau - \frac{\Delta}{2} \right) \geq v \left(\frac{3\Delta}{2} - \tau \right) \quad \text{for } \tau \leq \Delta, \quad (\text{A.31})$$

which allows rewriting (A.29) and (A.30) as

$$R_{01}(\Delta, \beta) = \min_{\tau \in [\frac{\Delta}{2}, \frac{1}{2}]} \left\{ \beta v \left(\tau - \frac{\Delta}{2} \right) + (1-\beta) v \left(\tau + \frac{\Delta}{2} \right) \right\}, \quad (\text{A.32})$$

$$R_{1m}(\Delta, \beta) = \min_{\tau \in [\frac{\Delta}{2}, \Delta]} \left\{ \beta v \left(\tau - \frac{\Delta}{2} \right) + (1-\beta) v \left(\frac{3\Delta}{2} - \tau \right) \right\}. \quad (\text{A.33})$$

Since v is decreasing by Assumption 1, the objective function in (A.32) is decreasing in τ and the minimum in (A.32) is achieved at $\tau = 1/2$:

$$R_{01}(\Delta, \beta) = \beta v \left(\frac{1-\Delta}{2} \right) + (1-\beta) v \left(\frac{1+\Delta}{2} \right). \quad (\text{A.34})$$

By Assumption 2, the second derivative w.r.t. τ of the objective function in (A.33) is negative. Hence, the minimum in (A.33) is achieved either at $\tau = \Delta/2$ or at $\tau = \Delta$:

$$R_{1m}(\Delta, \beta) = \min \left\{ \beta v(0) + (1 - \beta)v(\Delta), v\left(\frac{\Delta}{2}\right) \right\}. \quad (\text{A.35})$$

Substituting (A.34) and (A.35) into (A.28) and using Assumption 1 to claim that $v(0) \geq v\left(\frac{1-\Delta}{2}\right)$ and $v(\Delta) \geq v\left(\frac{1+\Delta}{2}\right)$, we get

$$R(\Delta, \beta) = \min \left\{ \beta v\left(\frac{1-\Delta}{2}\right) + (1 - \beta)v\left(\frac{1+\Delta}{2}\right), v\left(\frac{\Delta}{2}\right) \right\}, \quad (\text{A.36})$$

which is equivalent to (5).

Since v is strictly concave by Assumption 2, the first term in (A.36), $\beta v\left(\frac{1-\Delta}{2}\right) + (1 - \beta)v\left(\frac{1+\Delta}{2}\right)$, is concave in Δ . In other words, there exists $\hat{\Delta}(\beta) \in [0, 1]$ such that the first term in (A.36) is strictly increasing in $\Delta \in [0, \hat{\Delta}(\beta))$ and strictly decreasing in $\Delta \in (\hat{\Delta}(\beta), 1]$.

Since v is strictly decreasing by Assumption 1, the second term in (A.36), $v\left(\frac{\Delta}{2}\right)$, is strictly decreasing in $\Delta \in [0, 1]$. Thus, both terms in (A.36) are strictly decreasing in $\Delta \in (\hat{\Delta}(\beta), 1]$, which implies that

$$R(\hat{\Delta}(\beta), \beta) > R(\Delta, \beta) \quad \text{for all } \Delta \in (\hat{\Delta}(\beta), 1]. \quad (\text{A.37})$$

Suppose $\beta \leq 0.5$. Then, at $\Delta = 0$, since v is strictly decreasing by Assumption 1, the first term in (A.36) is weakly decreasing in Δ because its derivative $(0.5 - \beta)v'(0.5)$ is non-positive. Hence, $\hat{\Delta}(\beta) = 0$ and both terms in (A.36) are strictly decreasing in $\Delta \in [0, 1]$, which implies that $\Delta_R^* = 0$ uniquely maximizes $R(\Delta, \beta)$.

Suppose $\beta > 0.5$. Then, at $\Delta = 0$, the first derivative of the first term in (A.36), $(0.5 - \beta)v'(0.5)$, is positive by Assumption 1. Hence, $\hat{\Delta}(\beta) > 0$. For $\Delta \in [0, \hat{\Delta}(\beta))$, the first term in (A.36) is strictly increasing, while the second term in (A.36) is strictly decreasing. At $\Delta = 0$, since v is strictly decreasing by Assumption 1, the first term in (A.36), $v(0.5)$, is strictly less than the second term in (A.36), $v(0)$. Thus, there exists $\Delta_R^*(\beta) \in (0, \hat{\Delta}(\beta)]$ such that $R(\Delta, \beta)$ is equal to the first term and strictly increasing for $\Delta \in [0, \Delta_R^*(\beta))$, and equal to the second term and strictly decreasing for $\Delta \in (\Delta_R^*(\beta), \hat{\Delta}(\beta)]$. This $\Delta_R^*(\beta)$ uniquely maximizes $R(\Delta, \beta)$ on $\Delta \in [0, \hat{\Delta}(\beta)]$ and, therefore, by (A.37), on $\Delta \in [0, 1]$.

Suppose $\beta = 1$. Then, $\hat{\Delta}(\beta) = 1$ because the first term in (A.36) is strictly increasing in $\Delta \in [0, 1]$ since v is strictly decreasing. At $\Delta = 0.5$, the first and the second terms in (A.36) are equal. Hence, $\Delta_R^*(\beta) = 0.5$.

Consider the optimal diversity function $\Delta_R^*(\beta)$ on $\beta \in (0.5, 1)$.

At $\Delta = 1$, the first derivative of the first term in (A.36) is equal to $-0.5\beta v'(0) + 0.5(1 - \beta)v'(1)$. Since v is strictly decreasing by Assumption 1, this derivative is positive for $\beta > \bar{\beta}$ and negative for $\beta < \bar{\beta}$, where

$$\bar{\beta} = \frac{v'(1)}{v'(0) + v'(1)}. \quad (\text{A.38})$$

Note that $\bar{\beta} \in (0.5, 1)$ because $v'(1) < v'(0) < 0$ since v is strictly decreasing by Assumption 1 and strictly concave by Assumption 2.

Suppose $\beta \in [\bar{\beta}, 1)$. Then, the first term in (A.36) is strictly increasing for all $\Delta \in [0, 1]$, which means that $\hat{\Delta}(\beta) = 1$. At $\Delta = 1$, the first term in (A.36), $\beta v(0) + (1 - \beta)v(1)$, is strictly greater than the second term in (A.36), $v(0.5)$, because by Assumptions 1 and 2,

$$\begin{aligned} \beta v(0) + (1 - \beta)v(1) - v(0.5) &\stackrel{v(0) > v(1), \beta \geq \bar{\beta}}{\geq} \bar{\beta}(v(0) - v(0.5)) + (1 - \bar{\beta})(v(1) - v(0.5)) \\ &\stackrel{\bar{d} \in (0.5, 1), d \in (0, 0.5)}{=} -\bar{\beta}v'(\bar{d})0.5 + (1 - \bar{\beta})v'(\bar{d})0.5 \stackrel{v'' < 0}{>} -\bar{\beta}v'(0)0.5 + (1 - \bar{\beta})v'(1)0.5 \stackrel{(\text{A.38})}{=} 0. \end{aligned} \quad (\text{A.39})$$

Suppose $\beta \in (0.5, \bar{\beta})$. Then, the maximum of the first term in (A.36) is interior, which means that $\hat{\Delta}(\beta) \in (0, 1)$ satisfies

$$(1 - \beta)v' \left(\frac{1 + \Delta}{2} \right) = \beta v' \left(\frac{1 - \Delta}{2} \right) \quad \text{at } \Delta = \hat{\Delta}(\beta). \quad (\text{A.40})$$

By the implicit function theorem applied to (A.40), by Assumptions 1 and 2,

$$\frac{d\hat{\Delta}(\beta)}{d\beta} = \frac{v' \left(\frac{1 - \Delta}{2} \right) + v' \left(\frac{1 + \Delta}{2} \right)}{\frac{\beta}{2} v'' \left(\frac{1 - \Delta}{2} \right) + \frac{1 - \beta}{2} v'' \left(\frac{1 + \Delta}{2} \right)} \stackrel{v' < 0, v'' < 0}{>} 0 \quad \text{at } \Delta = \hat{\Delta}(\beta). \quad (\text{A.41})$$

Consider the difference between the first and the second terms in (A.36) at $\Delta = \hat{\Delta}(\beta)$. This difference is increasing in β because

$$\frac{d}{d\beta} \left\{ \beta v \left(\frac{1 - \hat{\Delta}(\beta)}{2} \right) + (1 - \beta) v \left(\frac{1 + \hat{\Delta}(\beta)}{2} \right) - v \left(\frac{\hat{\Delta}(\beta)}{2} \right) \right\} \stackrel{(\text{A.40})}{=} \underbrace{v \left(\frac{1 - \hat{\Delta}(\beta)}{2} \right) - v \left(\frac{1 + \hat{\Delta}(\beta)}{2} \right)}_{> 0 \text{ since } v' < 0} - \frac{1}{2} v' \left(\frac{\hat{\Delta}(\beta)}{2} \right) \underbrace{\frac{d\hat{\Delta}(\beta)}{d\beta}}_{> 0} > 0, \quad (\text{A.42})$$

positive at $\beta = \bar{\beta}$ by (A.39) and negative at $\beta = 0.5$ because $\hat{\Delta}(0.5) = 0$ and $v(0.5) < v(0)$.

Thus, there exists a unique $\beta^* \in (0.5, \bar{\beta})$ such that at $\Delta = \hat{\Delta}(\beta)$, the first term in (A.36) is strictly greater than the second term in (A.36) for all $\beta \in (\beta^*, 1)$, and the first term is strictly lower than the second term for all $\beta \in (0.5, \beta^*)$.

Suppose $\beta \in (0.5, \beta^*)$. Then, optimal diversity $\Delta_R^*(\beta)$ is equal to $\hat{\Delta}(\beta)$ and, thus, strictly increasing by (A.41).

Suppose $\beta \in (\beta^*, 1)$. Then, optimal diversity $\Delta_R^*(\beta)$ is less than $\hat{\Delta}(\beta)$ and equalizes the first and the second terms in (A.36):

$$\beta v \left(\frac{1 - \Delta_R^*(\beta)}{2} \right) + (1 - \beta) v \left(\frac{1 + \Delta_R^*(\beta)}{2} \right) = v \left(\frac{\Delta_R^*(\beta)}{2} \right). \quad (\text{A.43})$$

By the implicit function theorem applied to (A.43),

$$\frac{d\Delta_R^*(\beta)}{d\beta} = \frac{v \left(\frac{1 - \Delta_R^*(\beta)}{2} \right) - v \left(\frac{1 + \Delta_R^*(\beta)}{2} \right)}{\frac{\beta}{2} v' \left(\frac{1 - \Delta_R^*(\beta)}{2} \right) - \frac{1 - \beta}{2} v' \left(\frac{1 + \Delta_R^*(\beta)}{2} \right) + \frac{1}{2} v' \left(\frac{\Delta_R^*(\beta)}{2} \right)}. \quad (\text{A.44})$$

The numerator in (A.44) is positive because v is strictly decreasing by Assumption 1 and $\Delta_R^*(\beta) > 0$. The denominator in (A.44) is negative because $v' \left(\frac{\Delta_R^*(\beta)}{2} \right) < 0$ by Assumption 1 and the first term in (A.36) is increasing at $\Delta = \Delta_R^*(\beta) < \hat{\Delta}(\beta)$. Thus, (A.44) is negative.

A.5 Proof of Proposition 2

It is convenient to prove (i) last.

(ii) Consider $0.5 < \beta < \beta^*$. Then, optimal diversity Δ_R^* is equal to $\hat{\Delta}$ defined as a solution to (A.40). Substituting $v(d) = -d^\alpha$ into (A.40) and rearranging yield

$$\beta = \frac{(1 + \Delta)^{\alpha-1}}{(1 + \Delta)^{\alpha-1} + (1 - \Delta)^{\alpha-1}} \equiv f_1(\Delta, \alpha) \quad \text{at } \Delta = \Delta_R^*(\beta, \alpha). \quad (\text{A.45})$$

By the implicit function theorem applied to (A.45),

$$\frac{\partial \Delta_R^*(\beta, \alpha)}{\partial \alpha} = -\frac{\frac{\partial f_1(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial f_1(\Delta, \alpha)}{\partial \Delta}} = \frac{1 - \Delta^2}{2(\alpha - 1)} \ln\left(\frac{1 - \Delta}{1 + \Delta}\right) < 0 \quad \text{at } \Delta = \Delta_R^*(\beta, \alpha), \quad (\text{A.46})$$

where the sign follows from $\alpha > 1$ and $0 < \Delta_R^* < 1$.

- (iii) Consider $\beta^* < \beta < 1$. Then, optimal diversity Δ_R^* solves (A.43). Substituting $\nu(d) = -d^\alpha$ into (A.43) and rearranging yield

$$\beta = \frac{(1 + \Delta)^\alpha - \Delta^\alpha}{(1 + \Delta)^\alpha - (1 - \Delta)^\alpha} \equiv f_2(\Delta, \alpha) \quad \text{at } \Delta = \Delta_R^*(\beta, \alpha). \quad (\text{A.47})$$

By the implicit function theorem applied to (A.47),

$$\begin{aligned} \frac{\partial \Delta_R^*(\beta, \alpha)}{\partial \alpha} = -\frac{\frac{\partial f_2(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial f_2(\Delta, \alpha)}{\partial \Delta}} &= \frac{\left(\left(\frac{\Delta}{1+\Delta}\right)^\alpha - 1 - \ln\left(\left(\frac{\Delta}{1+\Delta}\right)^\alpha\right)\right) \ln\left(\frac{\Delta}{1-\Delta}\right) + \left(\left(\frac{\Delta}{1-\Delta}\right)^\alpha - 1 - \ln\left(\left(\frac{\Delta}{1-\Delta}\right)^\alpha\right)\right) \ln\left(\frac{1+\Delta}{\Delta}\right)}{\left(\frac{\Delta}{1-\Delta}\right)^{\alpha-1} - \left(\frac{\Delta}{1+\Delta}\right)^{\alpha-1} + 2} \\ &\quad \times \frac{1 - \Delta^2}{\alpha} > 0 \quad \text{at } \Delta = \Delta_R^*(\beta, \alpha), \quad (\text{A.48}) \end{aligned}$$

where the sign follows from $\alpha > 1$, $0.5 < \Delta_R^* < 1$ and the fact that $y - 1 - \ln y > 0$ for all $y \in (0, 1)$ and for all $y > 1$.

- (i) The value of β^* is defined as β which satisfies both (A.45) and (A.47):

$$\beta^*(\alpha) \stackrel{(I)}{=} f_1(\Delta, \alpha) \stackrel{(II)}{=} f_2(\Delta, \alpha) \quad \text{at } \Delta = \Delta_R^*(\beta^*(\alpha), \alpha). \quad (\text{A.49})$$

Applying the implicit function theorem to equality (II) in (A.49) gives

$$\frac{d\Delta_R^*(\beta^*(\alpha), \alpha)}{d\alpha} = \frac{\frac{\partial f_2(\Delta, \alpha)}{\partial \alpha} - \frac{\partial f_1(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial f_1(\Delta, \alpha)}{\partial \Delta} - \frac{\partial f_2(\Delta, \alpha)}{\partial \Delta}} \quad \text{at } \Delta = \Delta_R^*(\beta^*(\alpha), \alpha). \quad (\text{A.50})$$

We use equality (I) in (A.49) and the derivative of Δ_R^* in (A.50) to calculate the derivative of β^* :

$$\begin{aligned} \frac{d\beta^*(\alpha)}{d\alpha} &= \frac{\partial f_1(\Delta, \alpha)}{\partial \alpha} + \frac{\partial f_1(\Delta, \alpha)}{\partial \Delta} \frac{\frac{\partial f_2(\Delta, \alpha)}{\partial \alpha} - \frac{\partial f_1(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial f_1(\Delta, \alpha)}{\partial \Delta} - \frac{\partial f_2(\Delta, \alpha)}{\partial \Delta}} = \frac{\overbrace{\left(-\frac{\frac{\partial f_2(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial f_2(\Delta, \alpha)}{\partial \Delta}}\right)}^{>0 \text{ by (A.48)}} - \overbrace{\left(-\frac{\frac{\partial f_1(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial f_1(\Delta, \alpha)}{\partial \Delta}}\right)}^{<0 \text{ by (A.46)}}}{\underbrace{\frac{1}{\frac{\partial f_1(\Delta, \alpha)}{\partial \Delta}}}_{>0 \text{ by (A.52)}} - \underbrace{\frac{1}{\frac{\partial f_2(\Delta, \alpha)}{\partial \Delta}}}_{<0 \text{ by (A.53)}}} > 0 \\ &\quad \text{at } \Delta = \Delta_R^*(\beta^*(\alpha), \alpha), \quad (\text{A.51}) \end{aligned}$$

$$\frac{\partial f_1(\Delta, \alpha)}{\partial \Delta} = \frac{2(\alpha - 1)(1 - \Delta^2)^{\alpha-2}}{((1 + \Delta)^{\alpha-1} + (1 - \Delta)^{\alpha-1})^2} > 0 \quad \text{for all } \alpha > 1 \text{ and } 0 < \Delta < 1, \quad (\text{A.52})$$

$$\frac{\partial f_2(\Delta, \alpha)}{\partial \Delta} = -\frac{\alpha(1 - \Delta^2)^{\alpha-1} \left(\left(\frac{\Delta}{1-\Delta}\right)^{\alpha-1} - \left(\frac{\Delta}{1+\Delta}\right)^{\alpha-1} + 2\right)}{((1 + \Delta)^\alpha - (1 - \Delta)^\alpha)^2} < 0 \quad \text{for all } \alpha > 1 \text{ and } 0 < \Delta < 1. \quad (\text{A.53})$$

A.6 Proof of Theorem 3

Step 1. Defining $\Delta^*(\delta, \beta)$.

For $\delta \in (0, 0.5)$, define a δ -objective as the expected team output, given that $t_1 = \frac{1-\Delta}{2}$, $t_2 = \frac{1+\Delta}{2}$, and t is distributed uniformly on $[0, 0.5 - \delta] \cup [0.5 + \delta, 1]$:

$$U(\Delta, \delta, \beta) = \frac{2}{1-2\delta} \int_0^{0.5-\delta} \left(\beta v\left(\left|\frac{1-\Delta}{2} - t\right|\right) + (1-\beta)v\left(\frac{1+\Delta}{2} - t\right) \right) dt. \quad (\text{A.54})$$

Differentiating (A.54) yields

$$\frac{\partial U(\Delta, \delta, \beta)}{\partial \Delta} = \frac{1}{1-2\delta} \left(\beta \left(v\left(\left|\frac{\Delta}{2} - \delta\right|\right) - v\left(\frac{1-\Delta}{2}\right) \right) - (1-\beta) \left(v\left(\frac{\Delta}{2} + \delta\right) - v\left(\frac{1+\Delta}{2}\right) \right) \right), \quad (\text{A.55})$$

which gives

$$\beta \left(v\left(\left|\frac{\Delta^*}{2} - \delta\right|\right) - v\left(\frac{1-\Delta^*}{2}\right) \right) = (1-\beta) \left(v\left(\frac{\Delta^*}{2} + \delta\right) - v\left(\frac{1+\Delta^*}{2}\right) \right) \quad (\text{A.56})$$

as the first order condition for maximizing (A.54). The solution $\Delta^*(\delta, \beta) \in (0, 1)$ to (A.56) exists, is unique and gives the maximum of $U(\Delta, \delta, \beta)$. Indeed,

$$\frac{\partial U(0, \delta, \beta)}{\partial \Delta} = \frac{2\beta-1}{1-2\delta} \left(v\left(\delta\right) - v\left(\frac{1}{2}\right) \right) \quad (\text{A.57})$$

is positive because $\beta > 0.5$, $\delta < 0.5$, and v is strictly decreasing by Assumption 1;

$$\frac{\partial U(1, \delta, \beta)}{\partial \Delta} = \frac{1}{1-2\delta} \left(\beta \left(v\left(\frac{1}{2} - \delta\right) - v(0) \right) + (1-\beta) \left(v(1) - v\left(\frac{1}{2} + \delta\right) \right) \right) \quad (\text{A.58})$$

is negative because $0 < \beta < 1$, $\delta < 0.5$, and v is strictly decreasing by Assumption 1; both

$$\frac{\partial^2 U(\Delta, \delta, \beta)}{\partial \Delta^2} = \frac{1}{2(1-2\delta)} \left(\beta \left(v'\left(\frac{1-\Delta}{2}\right) + v'\left(\frac{\Delta}{2} - \delta\right) \right) + (1-\beta) \left(v'\left(\frac{1+\Delta}{2}\right) - v'\left(\frac{\Delta}{2} + \delta\right) \right) \right) \quad \text{for } \delta < \Delta/2 \quad (\text{A.59})$$

and

$$\frac{\partial^2 U(\Delta, \delta, \beta)}{\partial \Delta^2} = \frac{1}{2(1-2\delta)} \left(\beta \left(v'\left(\frac{1-\Delta}{2}\right) - v'\left(\delta - \frac{\Delta}{2}\right) \right) + (1-\beta) \left(v'\left(\frac{1+\Delta}{2}\right) - v'\left(\frac{\Delta}{2} + \delta\right) \right) \right) \quad \text{for } \delta > \Delta/2 \quad (\text{A.60})$$

are negative because $0 < \beta < 1$, $\delta < 0.5$, v is strictly decreasing by Assumption 1 and strictly concave by Assumption 2.

Step 2. Function $\Delta^*(\delta, \beta)$ is increasing in $\delta \in (0, 0.5)$.

By the implicit function theorem, the sign of $\frac{\partial \Delta^*(\delta, \beta)}{\partial \delta}$ coincides with the sign of

$$\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U(\Delta, \delta, \beta)}{\partial \Delta} \right) = \begin{cases} -\beta v'(\Delta/2 - \delta) - (1-\beta) v'(\Delta/2 + \delta), & \delta < \Delta/2, \\ \beta v'(\delta - \Delta/2) - (1-\beta) v'(\Delta/2 + \delta), & \delta > \Delta/2 \end{cases} \quad (\text{A.61})$$

at $\Delta = \Delta^*(\delta, \beta)$. For $\delta < \Delta/2$, expression (A.61) is positive because $0 < \beta < 1$ and v is strictly decreasing by Assumption 1. For $\delta > \Delta/2$, since v' is negative by Assumption 1, the sign of (A.61) is positive at $\Delta = \Delta^*(\delta, \beta)$ if

$$\frac{\beta}{1-\beta} < \frac{v'(\frac{\Delta^*}{2} + \delta)}{v'(\delta - \frac{\Delta^*}{2})}. \quad (\text{A.62})$$

For $\delta > \Delta^*/2$, equation (A.56) allows expressing the left-hand side of (A.62) as

$$\frac{\beta}{1-\beta} = \frac{v(\frac{\Delta^*}{2} + \delta) - v(\frac{1+\Delta^*}{2})}{v(\delta - \frac{\Delta^*}{2}) - v(\frac{1-\Delta^*}{2})} \quad (\text{A.63})$$

Substituting $\frac{\beta}{1-\beta}$ from (A.63) into (A.62) and multiplying by the denominator of (A.63), which is positive by $\delta < 1/2$ and Assumption 1, we conclude that (A.62) holds if

$$g(\Delta, \delta) = \frac{v'(\frac{\Delta}{2} + \delta)}{v'(\delta - \frac{\Delta}{2})} \left(v\left(\delta - \frac{\Delta}{2}\right) - v\left(\frac{1-\Delta}{2}\right) \right) - \left(v\left(\frac{\Delta}{2} + \delta\right) - v\left(\frac{1+\Delta}{2}\right) \right) \quad (\text{A.64})$$

is positive for all $1/2 > \delta > \Delta/2 > 0$. Function $g(\Delta, \delta)$ is indeed positive in this range because the derivative

$$\frac{\partial g(\Delta, \delta)}{\partial \delta} = \frac{v'(\frac{\Delta}{2} + \delta)}{v'(\delta - \frac{\Delta}{2})} \left(v\left(\delta - \frac{\Delta}{2}\right) - v\left(\frac{1-\Delta}{2}\right) \right) \left(A\left(\delta + \frac{\Delta}{2}\right) - A\left(\delta - \frac{\Delta}{2}\right) \right) \quad (\text{A.65})$$

is negative by Assumptions 1 and 4, and at $\delta = 1/2$, $g(\Delta, \delta)$ is zero.

Step 3. Main argument for $\Delta_R^*(\beta) > \Delta_U^*(\beta)$.

Let β^* be the threshold from Theorem 2. For $\beta \in (\beta^*, 1)$, $\Delta_R^*(\beta) > \Delta_U^*(\beta)$ follows from Theorems 1 and 2: $\Delta_U^*(\beta) < 0.5$ by Theorem 1 and $\Delta_R^*(\beta) > 0.5$ by Theorem 2.

Consider $\beta \in (0.5, \beta^*)$. As we show in the proof of Theorem 2 in Appendix A.4, $\Delta_R^*(\beta)$ is equal to $\hat{\Delta}(\beta)$, which maximizes $R_{01}(\Delta, \beta)$ defined in (A.34). As δ approaches 0.5, $\Delta^*(\delta, \beta)$ converges to $\Delta_R^*(\beta)$ because the δ -objective (A.54) converges to $R_{01}(\Delta, \beta)$. As δ approaches 0, $\Delta^*(\delta, \beta)$ converges to $\Delta_U^*(\beta)$ because at $\delta = 0$, the δ -objective (A.54) is the utilitarian objective. Thus, since $\Delta^*(\delta, \beta)$ is increasing in δ , $\Delta_R^*(\beta) = \Delta^*(0.5, \beta) > \Delta^*(0, \beta) = \Delta_U^*(\beta)$.

A.7 Proof of Remark 4

Without Assumption 4, it cannot be guaranteed that Theorem 3 holds. For example, in Figure A.1, on the interval $\beta \in (0.5, \hat{\beta})$, the Rawlsian objective gives a slightly lower optimal diversity than the utilitarian objective. While the individual output function $v(d) = -d - 0.01d^2(1 + d + d^2) - 4d^5$ satisfies Assumptions 1 and 2, it does not satisfy Assumption 4, as illustrated in Figure A.3.

However, Assumption 4 could be overly restrictive. For instance, in Figure A.2, the individual output function does not satisfy Assumption 4 (see Figure A.3) but for all $\beta \in (0.5, 1)$, the Rawlsian objective gives a higher optimal diversity than the utilitarian objective.

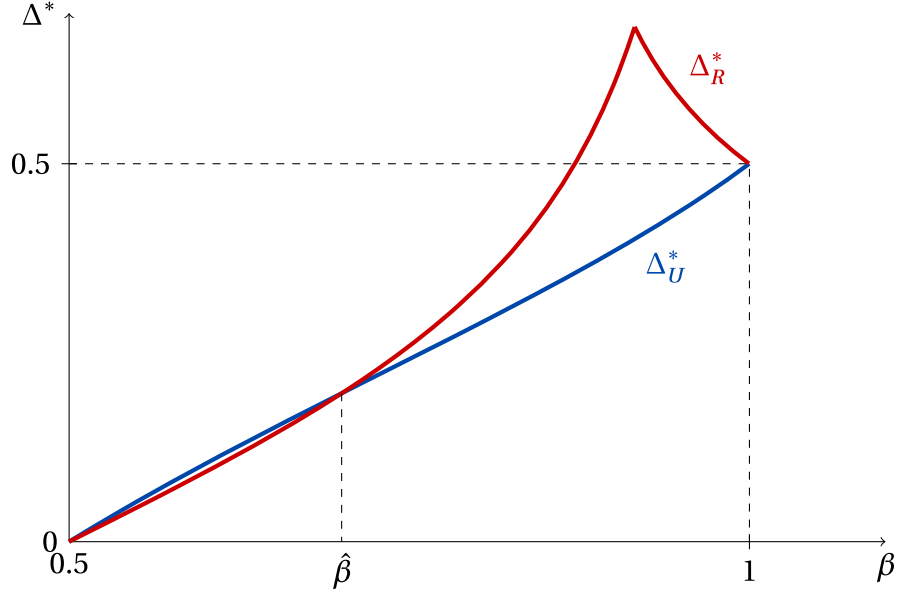


Figure A.1: Optimal diversity, $v(d) = -d - 0.01d^2(1 + d + d^2) - 4d^5$. Blue graph corresponds to the utilitarian objective; red graph corresponds to the Rawlsian objective.

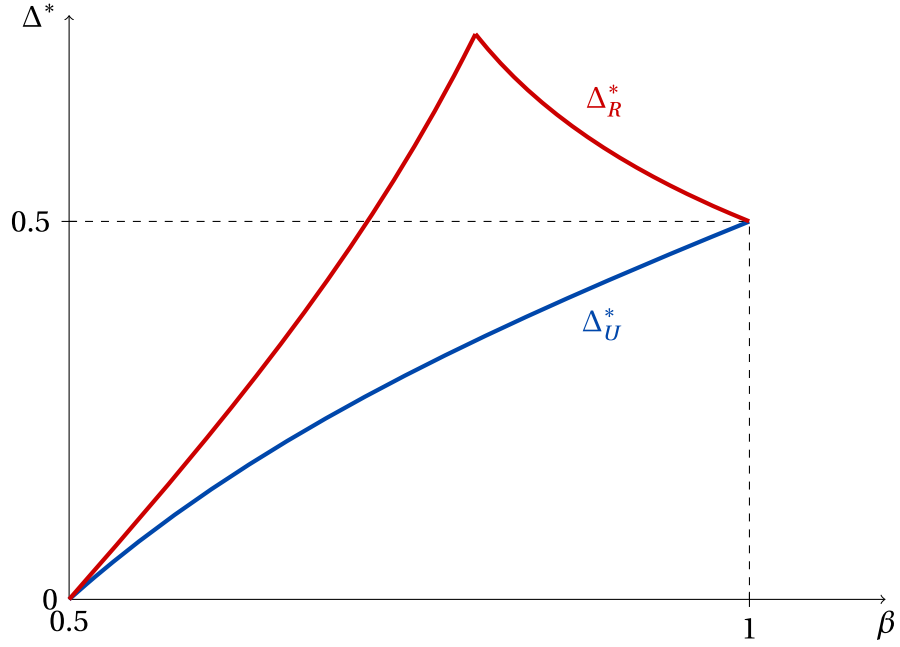


Figure A.2: Optimal diversity, $v(d) = -d - 0.01d^2(1 + d + d^2) - d^5$. Blue graph corresponds to the utilitarian objective; red graph corresponds to the Rawlsian objective.

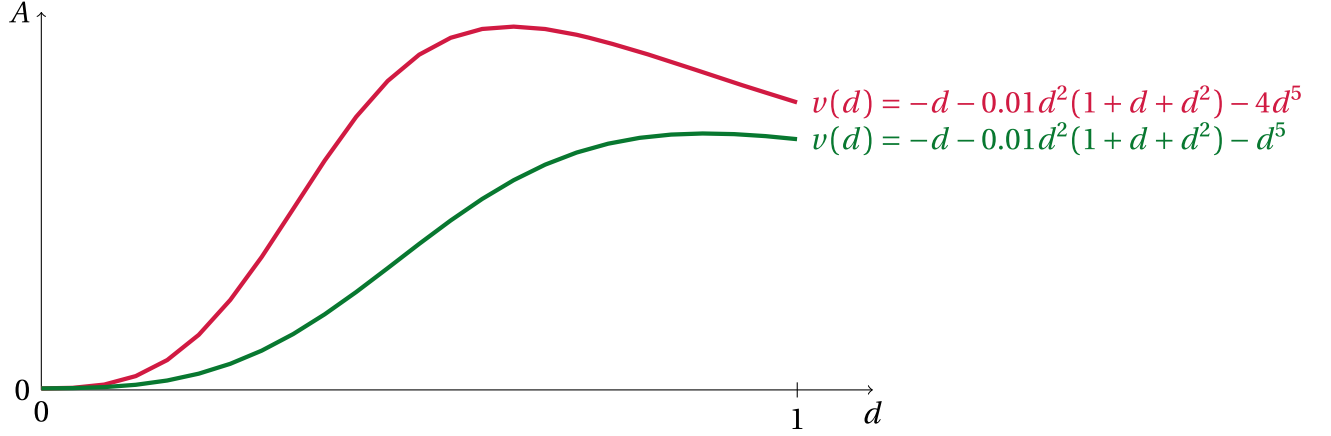


Figure A.3: $A(d)$ for functions $v(d)$ from Figures A.1 and A.2.

A.8 Proof of Theorem 4

Equation for $\Delta = \Delta_R^{**}(\alpha)$ is given by equality (II) in (A.49), which can be rewritten as follows:

$$g(\Delta, \alpha) \equiv \left(\frac{\Delta}{1-\Delta}\right)^{\alpha-1} + \left(\frac{\Delta}{1+\Delta}\right)^{\alpha-1} - \frac{2}{\Delta} = 0 \quad \text{at } \Delta = \Delta_R^{**}(\alpha). \quad (\text{A.66})$$

Applying the implicit function theorem to (A.66) yields

$$\frac{d\Delta_R^{**}(\alpha)}{d\alpha} = -\frac{\frac{\partial g(\Delta, \alpha)}{\partial \alpha}}{\frac{\partial g(\Delta, \alpha)}{\partial \Delta}} \quad \text{at } \Delta = \Delta_R^{**}(\alpha). \quad (\text{A.67})$$

Since $\Delta_R^{**}(\alpha) = \Delta_R^*(\beta^*(\alpha), \alpha) > 0.5$ by Theorem 2, and since function $g(\Delta, \alpha)$, defined in (A.66), is increasing in both arguments for all $0.5 < \Delta < 1$ and $\alpha > 1$, the derivative (A.67) is negative.

B Generalized model

B.1 Setup

In this section, we generalize Assumption 3.

Assumption 5. *The team output is given by*

$$\omega(x_1, x_2) = \beta \times \omega_s(x_1, x_2) + (1 - \beta) \times \omega_c(x_1, x_2), \quad (\text{B.1})$$

where $x_i = v(|t - t_i|)$ is agent i 's individual output.

Assumption 5 generalizes functions max and min to arbitrary functions ω_s and ω_c with the following properties; subscripts “s” and “c” in notations ω_s and ω_c stand for “substitutes” and “complements,” respectively.

Assumption 6. *Functions ω_s and ω_c are weakly increasing in each argument.*

Assumption 7. *$\omega_s(x_1, x_2) = \omega_s(x_2, x_1)$ and $\omega_c(x_1, x_2) = \omega_c(x_2, x_1)$ for any (x_1, x_2) .*

Assumption 8. Function ω_s is **submodular** in the agents' individual outputs: for any $(\hat{x}_1, \hat{x}_2)$ and $(\check{x}_1, \check{x}_2)$,

$$\omega_s(\hat{x}_1, \hat{x}_2) + \omega_s(\check{x}_1, \check{x}_2) \geq \omega_s(\min\{\hat{x}_1, \check{x}_1\}, \min\{\hat{x}_2, \check{x}_2\}) + \omega_s(\max\{\hat{x}_1, \check{x}_1\}, \max\{\hat{x}_2, \check{x}_2\}), \quad (\text{B.2})$$

while function ω_c is **supermodular** in the agents' individual outputs: for any $(\hat{x}_1, \hat{x}_2)$ and $(\check{x}_1, \check{x}_2)$,

$$\omega_c(\hat{x}_1, \hat{x}_2) + \omega_c(\check{x}_1, \check{x}_2) \leq \omega_c(\min\{\hat{x}_1, \check{x}_1\}, \min\{\hat{x}_2, \check{x}_2\}) + \omega_c(\max\{\hat{x}_1, \check{x}_1\}, \max\{\hat{x}_2, \check{x}_2\}). \quad (\text{B.3})$$

Assumption 9. $\omega_s(x_1, x_2) + \omega_c(x_1, x_2) = x_1 + x_2$ for any (x_1, x_2) .

Assumption 6 states that higher performance of an agent always weakly benefits the principal. Assumption 7 establishes that functions ω_s and ω_c are symmetric, which means that, for a given cognitive type, the identity of an agent does not affect the team output. Assumption 8 is central to our generalization. Following Prat (2002), we appeal to submodularity and supermodularity as a generalization of the traditional notions of strategic substitutability and strategic complementarity, respectively.¹⁹ Finally, Assumption 9 ensures that at $\beta = 0.5$, the team output $\omega(x_1, x_2) = 0.5(x_1 + x_2)$ satisfies (B.2) and (B.3) as equality. In other words, at $\beta = 0.5$, the submodular and supermodular components cancel each other, making the agents' contributions strategically independent. Thus, as in the benchmark model with $\omega_s(x_1, x_2) = \max\{x_1, x_2\}$ and $\omega_c(x_1, x_2) = \min\{x_1, x_2\}$, $\beta = 0.5$ corresponds to an additive task.

Given Assumption 9, we can rewrite (B.1) as

$$\omega(x_1, x_2) = (2\beta - 1) \times \omega_s(x_1, x_2) + (1 - \beta)(x_1 + x_2). \quad (\text{B.4})$$

Thus, function $\omega(x_1, x_2)$ is supermodular for $\beta \leq 0.5$ and submodular for $\beta \geq 0.5$.

Lemma 1, proved in Appendix A.1, assures the optimality of $t_1 = (1 - \Delta)/2$ and $t_2 = (1 + \Delta)/2$ in the general model; that is, Lemma 1 holds with Assumptions 1 and 5 to 7.

To preserve uniqueness of optimal diversity, we impose the following sufficient condition:

Assumption 10. For all $t \in [0, \frac{1}{2}]$, function $\omega_s\left(v\left(\left|\frac{1-\Delta}{2} - t\right|\right), v\left(\frac{1+\Delta}{2} - t\right)\right)$ is strictly concave with respect to $\Delta \in [0, 1]$.

Assumption 10 says that for any task, the team output for a submodular production function is concave with respect to diversity. Due to symmetry, it is sufficient to limit attention only to tasks from $[0, \frac{1}{2}]$. For these tasks, agent 1's individual output $v\left(\left|\frac{1-\Delta}{2} - t\right|\right) = v(|t_1 - t|)$ is higher than agent 2's individual output $v\left(\frac{1+\Delta}{2} - t\right) = v(t_2 - t)$. In the benchmark model, $\omega_s(x_1, x_2) = \max\{x_1, x_2\}$, and so, Assumption 10 holds because $\omega_s\left(v\left(\left|\frac{1-\Delta}{2} - t\right|\right), v\left(\frac{1+\Delta}{2} - t\right)\right) = v\left(\left|\frac{1-\Delta}{2} - t\right|\right)$ is strictly concave by Assumption 2.

Although Assumption 10 guarantees the uniqueness of optimal diversity, it does not exclude corner solutions. In particular, optimal diversity can be 0. To ensure that optimal diversity is positive for all $\beta > 0.5$, we impose the following sufficient condition:

Assumption 11. For all $t \in [0, \frac{1}{2}]$, function $\omega_s\left(v\left(\left|\frac{1-\Delta}{2} - t\right|\right), v\left(\frac{1+\Delta}{2} - t\right)\right)$ is strictly increasing at $\Delta = 0^+$.

Assumption 11 says that for any task, the team output for a submodular production function is increasing as diversity grows from 0 to some positive level. Assumptions 10 and 11 together ensure that optimal diversity is unique and belongs to $(0, 1)$ for $\beta > 0.5$.

¹⁹For twice differentiable functions, condition (B.2) is equivalent to $\partial^2 \omega_s(x_1, x_2) / \partial x_1 \partial x_2 \leq 0$, while condition (B.3) is equivalent to $\partial^2 \omega_c(x_1, x_2) / \partial x_1 \partial x_2 \geq 0$.

In the benchmark model, Assumption 11 holds because $\omega_s\left(v\left(\left|\frac{1-\Delta}{2}-t\right|\right), v\left(\frac{1+\Delta}{2}-t\right)\right) = v\left(\left|\frac{1-\Delta}{2}-t\right|\right)$ is strictly increasing in $\Delta \in (0, 1-2t)$ by Assumption 1.

In general, Assumption 11 requires that function $\omega_s(x_1, x_2)$ has a convex kink at $x_1 = x_2$. To see this, we rewrite Assumption 11 as

$$v'\left(\frac{1}{2}-t\right)\left(\frac{\partial\omega_s(x, x-0)}{\partial x_2} - \frac{\partial\omega_s(x+0, x)}{\partial x_1}\right) > 0, \quad (\text{B.5})$$

where $x = v\left(\frac{1}{2}-t\right)$ and notation “ ± 0 ” means that the derivative with respect to x_1 is taken from above (i.e., as x_1 approaches x from above) and the derivative with respect to x_2 is taken from below. In light of Assumptions 1 and 7, condition (B.5) is equivalent to

$$\frac{\partial\omega_s(x+0, x)}{\partial x_1} > \frac{\partial\omega_s(x-0, x)}{\partial x_1}. \quad (\text{B.6})$$

Inequality (B.6) implies that function $\omega_s(x_1, x_2)$ has a convex kink at $x_1 = x_2$.

The following example demonstrates that all the above assumptions are not too restrictive, and the benchmark functions $\omega_s(x_1, x_2) = \max\{x_1, x_2\}$ and $\omega_c(x_1, x_2) = \min\{x_1, x_2\}$ are not the only admissible ones.

Example 1. Assuming $v(d) = -d^\alpha$ with $\alpha \geq 2$, functions

$$\omega_s(x_1, x_2) = (1-\gamma)\frac{x_1+x_2}{2} + \gamma \max\{x_1, x_2\} - kx_1x_2, \quad (\text{B.7})$$

$$\omega_c(x_1, x_2) = (1-\gamma)\frac{x_1+x_2}{2} + \gamma \min\{x_1, x_2\} + kx_1x_2, \quad (\text{B.8})$$

satisfy Assumptions 6 to 11 for all $0 < \gamma \leq 1$ and $0 \leq k \leq \frac{1-\gamma}{2}$ (see Appendix B.5.1 for the proof). In particular, restriction $k \leq \frac{1-\gamma}{2}$ ensures that functions ω_s and ω_c are weakly increasing in each argument (Assumption 6), while restriction $\alpha \geq 2$ makes function v sufficiently concave to ensure that Assumption 10 holds.

B.2 Utilitarian objective

Theorem 1 also holds in the general model.

Theorem 1*. With Assumptions 1 and 5 to 11, under the utilitarian objective, diversity $\Delta_U^* = 0$ is optimal for $\beta \in [0, 0.5]$; for $\beta \in [0.5, 1]$, optimal diversity $\Delta_U^*(\beta)$ is unique and strictly increasing in $\beta \in [0.5, 1]$, from $\Delta_U^*(0.5) = 0$.

Proof. See Appendix B.5.2. □

Unlike the benchmark model, here we cannot guarantee that $\Delta_U^*(1) = 0.5$. Moreover, Assumption 10 ensures uniqueness only for $\beta > 0.5$; for $\beta < 0.5$, we can only prove that $\Delta_U^* = 0$ is one of the maximizers.

B.3 Rawlsian objective

When applying the Rawlsian objective to the general model, we focus only on submodular team output — that is, when $\beta \geq 0.5$. This is for two reasons. Firstly, we cannot guarantee the uniqueness of optimal diversity when $\beta < 0.5$. Secondly, and more importantly, in the benchmark model in Sections 2 to 5, only the region $\beta \geq 0.5$ provides non-trivial comparative statics with respect to β .

For Theorem 2 to apply to the general model, we need an additional assumption. This assumption assures that β^* is below 1.

Assumption 12. *Let $\hat{\Delta}$ be the value of $\Delta \in [0, 1]$ that maximizes $\omega_s\left(v\left(\frac{1-\Delta}{2}\right), v\left(\frac{1+\Delta}{2}\right)\right)$.*

Then, $\omega_s\left(v\left(\frac{1-\hat{\Delta}}{2}\right), v\left(\frac{1+\hat{\Delta}}{2}\right)\right) > \omega_s\left(v\left(\frac{\hat{\Delta}}{2}\right), v\left(\frac{\hat{\Delta}}{2}\right)\right)$.

Assumption 12 says that when $\beta = 1$ and diversity is set to maximize the team output for the corner tasks, this output is larger than for the center task.²⁰ In other words, Assumption 12 assures that at $\beta = 1$, the principal faces the corners/center trade-off: at the diversity level chosen with only the corner tasks in mind, the minimum team output is achieved at the center. Consequently, to increase the output at the center, the principal should lower diversity so as to equalize the team output at the corners and at the center.

Note that in the benchmark model, Assumption 12 follows from Assumption 1, as

$\omega_s\left(v\left(\frac{1-\Delta}{2}\right), v\left(\frac{1+\Delta}{2}\right)\right) = v\left(\frac{1-\Delta}{2}\right)$ is maximized at $\hat{\Delta} = 1$, and

$\omega_s\left(v\left(\frac{1-\hat{\Delta}}{2}\right), v\left(\frac{1+\hat{\Delta}}{2}\right)\right) - \omega_s\left(v\left(\frac{\hat{\Delta}}{2}\right), v\left(\frac{\hat{\Delta}}{2}\right)\right) = v\left(\frac{1-\hat{\Delta}}{2}\right) - v\left(\frac{\hat{\Delta}}{2}\right)$ is positive at $\hat{\Delta} = 1$.

The fact that $\beta^* < 1$ requires an additional assumption is unsurprising. Indeed, function $\omega_s(x_1, x_2) = \beta' \max\{x_1, x_2\} + (1 - \beta') \min\{x_1, x_2\}$ for any $\beta' \in (0.5, 1]$, however close to 0.5, satisfies all the assumptions in Appendix B.1. With this function, the team output is $\omega(x_1, x_2) = \check{\beta} \max\{x_1, x_2\} + (1 - \check{\beta}) \min\{x_1, x_2\}$, where $\check{\beta} = 1 - \beta' + \beta(2\beta' - 1)$. From Theorem 2 for the benchmark model, we know that there exists β^* such that optimal diversity $\Delta_R^*(\check{\beta})$ is increasing on $\check{\beta} \in (0.5, \beta^*)$. Thus, since $\check{\beta} = 1 - \beta' + \beta(2\beta' - 1)$, if $\beta' < \beta^*$, then $\Delta_R^*(\beta)$ is increasing on $\beta \in (0.5, 1)$. Thus, intuitively, for the decreasing part of the optimal diversity function to appear, we should require that at $\beta = 1$, the team output function is sufficiently submodular.

Indeed, Assumption 12 requires that the submodular component of the team output function is sufficiently submodular. Since $v\left(\frac{1+\hat{\Delta}}{2}\right) < v\left(\frac{\hat{\Delta}}{2}\right)$ by Assumption 1 and function ω_s is weakly increasing in each argument by Assumption 6, condition $\omega_s\left(v\left(\frac{1-\hat{\Delta}}{2}\right), v\left(\frac{1+\hat{\Delta}}{2}\right)\right) > \omega_s\left(v\left(\frac{\hat{\Delta}}{2}\right), v\left(\frac{\hat{\Delta}}{2}\right)\right)$ holds only if (1) $v\left(\frac{1-\hat{\Delta}}{2}\right) > v\left(\frac{\hat{\Delta}}{2}\right)$, and (2) the dependence of ω_s on the first argument is relatively higher than on the second argument. Since function v is decreasing by Assumption 1, condition (1) is equivalent to $\hat{\Delta} > 0.5$. If $\omega_s\left(v\left(\frac{1-\hat{\Delta}}{2}\right), v\left(\frac{1+\hat{\Delta}}{2}\right)\right)$ relies more on the first argument, then $\hat{\Delta}$ is indeed high because the first argument is maximized at $\Delta = 1$. Hence, roughly speaking, condition (1) follows from informally formulated condition (2). Condition (2) says that the team output relies relatively more on the individual output of the stronger agent than one of the weaker agent. Informally speaking, a heavier weight on the stronger agent's output increases the degree of submodularity of the team output function.

Formally, in Example 1, Assumption 12 holds if γ is sufficiently close to 1 because it holds for the benchmark model where $\gamma = 1$. High γ means that $\omega_s(x_1, x_2)$, defined in (B.7), depends more on $\max\{x_1, x_2\}$ than on $\min\{x_1, x_2\}$.

²⁰Note that by Assumption 10 with $t = 0$, $\hat{\Delta}$ is unique.

Theorem 2*. With Assumptions 1, 2 and 5 to 11, under the Rawlsian objective, optimal diversity $\Delta_R^*(\beta)$ is unique for all $\beta \in [0.5, 1]$ and equal to 0 at $\beta = 0.5$; there exists $\beta^* \in (0.5, 1]$ such that optimal diversity $\Delta_R^*(\beta)$

- strictly increases for $\beta \in (0.5, \beta^*)$, and
- strictly decreases and greater than 0.5 for $\beta \in (\beta^*, 1)$.

If, in addition, Assumption 12 holds, then $\beta^* < 1$.

Proof. See Appendix B.5.3. □

Remark 2* extends Remark 2 in that for any diversity, the team output is minimized either at the corners or at the center.

Remark 2*. In Appendix B.5.3, in the proof of Theorem 2*, we show that the Rawlsian objective can be written as

$$R = \min \left\{ \omega \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right), \omega \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) \right\}, \quad (\text{B.9})$$

where the first term is the team output for the corner tasks ($t = 0$ and $t = 1$) and the second term is the team output for the task at the center ($t = 1/2$).

For $\beta \in (0.5, 1)$, the behavior of (B.9) is similar to the behavior of (5) in the benchmark model. The team output for the center task, $\omega \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right)$, is decreasing in Δ and, thus, is maximized at $\Delta = 0$. The team output for the corners, $\omega \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right)$, is concave and maximized at some positive $\hat{\Delta}(\beta)$, which is increasing in β .

For $\beta \in (0.5, \beta^*)$, at $\Delta = \hat{\Delta}(\beta)$, the team output at the corners is lower than in the center. Hence, optimal diversity maximizes the team output at the corners, $\Delta_R^*(\beta) = \hat{\Delta}(\beta)$, and, at optimal diversity, the minimum in (B.9) is achieved only at the first term.

For $\beta \in (\beta^*, 1)$, at $\Delta = \hat{\Delta}(\beta)$, the team output at the corners is higher than in the center. Hence, the principal lowers optimal diversity $\Delta_R^*(\beta)$ below $\hat{\Delta}(\beta)$ to equalize the team output at the corners and in the center — so that, at optimal diversity, the minimum in (B.9) is achieved at both terms.

As in the benchmark model, optimal diversity decreases whenever the minimum team output is achieved at both the corners and the center, and increases whenever the minimum team output is achieved only at the corners.

B.4 Comparison of optimal diversity under different objectives

To generalize Theorem 3, we impose the following two assumptions.

Assumption 13. For all $\Delta \in [0, 1]$, function $\frac{\frac{\partial}{\partial \Delta} \omega_s \left(v \left(\frac{1-\Delta}{2} - t \right), v \left(\frac{1+\Delta}{2} - t \right) \right)}{-\frac{\partial}{\partial \Delta} \left(v \left(\frac{1-\Delta}{2} - t \right) + v \left(\frac{1+\Delta}{2} - t \right) \right)}$ is strictly decreasing in $t \in [0, \frac{1-\Delta}{2}]$.

Assumption 13 says that as diversity increases, a marginal change in the team output for a sub-modular production function, normalized to a marginal change in the team output for an additive production function, is increasing as tasks become more extreme. In other words, as tasks become more extreme, the benefit from increasing diversity grows.

Note that for an additive production function, the benefit from increasing diversity is negative (so that optimal diversity is 0): $\frac{\partial}{\partial \Delta} \left(v \left(\frac{1-\Delta}{2} - t \right) + v \left(\frac{1+\Delta}{2} - t \right) \right) < 0$. Thus, in Assumption 13, we take this benefit with minus sign to keep the normalization term (i.e., denominator) positive.

In the benchmark model, Assumption 13 follows from Assumptions 2 and 4 because

$$\begin{aligned} \frac{\partial}{\partial t} \left(\frac{\frac{\partial}{\partial \Delta} \omega_s \left(v \left(\frac{1-\Delta}{2} - t \right), v \left(\frac{1+\Delta}{2} - t \right) \right)}{-\frac{\partial}{\partial \Delta} \left(v \left(\frac{1-\Delta}{2} - t \right) + v \left(\frac{1+\Delta}{2} - t \right) \right)} \right) &= -\frac{\partial}{\partial t} \left(\frac{\frac{\partial}{\partial \Delta} v \left(\frac{1-\Delta}{2} - t \right)}{\frac{\partial}{\partial \Delta} \left(v \left(\frac{1-\Delta}{2} - t \right) + v \left(\frac{1+\Delta}{2} - t \right) \right)} \right) \\ &= \frac{v' \left(\frac{1-\Delta}{2} - t \right) v' \left(\frac{1+\Delta}{2} - t \right)}{\underbrace{\left(v' \left(\frac{1-\Delta}{2} - t \right) - v' \left(\frac{1+\Delta}{2} - t \right) \right)^2}_{>0 \text{ by Assumption 2}}} \underbrace{\left(A \left(\frac{1+\Delta}{2} - t \right) - A \left(\frac{1-\Delta}{2} - t \right) \right)}_{<0 \text{ by Assumption 4}}. \quad (\text{B.10}) \end{aligned}$$

Assumption 14. Function $\int_0^{0.5} \omega_s \left(v \left(\left| \frac{1-\Delta}{2} - t \right| \right), v \left(\frac{1+\Delta}{2} - t \right) \right) dt$ is non-increasing at $\Delta = 0.5$.

Assumption 14 says that for $\beta = 1$, the utilitarian objective is non-increasing at $\Delta = 0.5$. Together with Assumption 10, Assumption 14 implies that at $\beta = 1$, optimal diversity under the utilitarian objective is less or equal to 0.5: $\Delta_U^*(1) \leq 0.5$.

In the benchmark model, Assumption 14 trivially holds because

$$\frac{d}{d\Delta} \int_0^{0.5} \omega_s \left(v \left(\left| \frac{1-\Delta}{2} - t \right| \right), v \left(\frac{1+\Delta}{2} - t \right) \right) dt = \frac{d}{d\Delta} \int_0^{0.5} v \left(\left| \frac{1-\Delta}{2} - t \right| \right) dt = \frac{1}{2} \left(v \left(\frac{\Delta}{2} \right) - v \left(\frac{1-\Delta}{2} \right) \right) \quad (\text{B.11})$$

is equal to 0 at $\Delta = 0.5$.

Both Assumptions 13 and 14 hold for Example 1, without additional restrictions on parameters (see Appendix B.5.1 for the proof).

Theorem 3*. Let Assumptions 1, 2, 5 to 11, 13 and 14 hold. Then, for $\beta \in (0.5, 1)$, optimal diversity under the Rawlsian objective is always higher than under the utilitarian one: $\Delta_R^*(\beta) > \Delta_U^*(\beta)$.

Proof. See Appendix B.5.4. □

In Theorem 3*, Assumption 13 is crucial, playing the same role as Assumption 4 in Theorem 3. Intuitively, as δ increases, the δ -objective becomes concentrated on more extreme tasks. By Assumption 13, the benefit from increasing diversity grows. Thus, optimal diversity increases in δ , and so, $\Delta_R^*(\beta) > \Delta_U^*(\beta)$ for $\beta \in (0.5, \beta^*)$.

For $\beta \in (\beta^*, 1)$, $\Delta_R^*(\beta) > \Delta_U^*(\beta)$ follows from Assumption 14. Indeed, by Theorem 2*, for $\beta \in (\beta^*, 1)$, optimal diversity under the Rawlsian objective belongs to the interval $(0.5, 1]$. By Theorem 1*, optimal diversity under the utilitarian objective is increasing in β and, thus, it is less than 0.5 for all $\beta \in (\beta^*, 1)$ if and only if it is less or equal to 0.5 at $\beta = 1$. The latter is guaranteed by Assumption 14.

B.5 Proofs

B.5.1 Example 1

Given that function $v(d) = -d^\alpha$ is decreasing in $d \in [0, 1]$ from 0 to -1 , the range on which functions $\omega_s(x_1, x_2)$ and $\omega_c(x_1, x_2)$ must satisfy the desired properties is

$$-1 \leq x_1 \leq 0, \quad -1 \leq x_2 \leq 0. \quad (\text{B.12})$$

Assumption 7 holds because functions $\omega_s(x_1, x_2)$ and $\omega_c(x_1, x_2)$ defined in (B.7) and (B.8) are symmetric in their arguments.

By symmetry, to verify Assumption 6, it is sufficient to show the monotonicity of functions ω_s and ω_c for $x_1 \geq x_2$:

$$\frac{\partial \omega_s(x_1, x_2)}{\partial x_1} \stackrel{(B.7)}{=} \frac{1+\gamma}{2} - kx_2 \stackrel{(B.12), \gamma \geq 0, k \geq 0}{\geq} 0, \quad (B.13)$$

$$\frac{\partial \omega_s(x_1, x_2)}{\partial x_2} \stackrel{(B.7)}{=} \frac{1-\gamma}{2} - kx_1 \stackrel{(B.12), \gamma \leq 1, k \geq 0}{\geq} 0, \quad (B.14)$$

$$\frac{\partial \omega_c(x_1, x_2)}{\partial x_1} \stackrel{(B.8)}{=} \frac{1-\gamma}{2} + kx_2 \stackrel{(B.12), k \geq 0}{\geq} \frac{1-\gamma}{2} - k \geq 0, \quad (B.15)$$

$$\frac{\partial \omega_c(x_1, x_2)}{\partial x_2} \stackrel{(B.8)}{=} \frac{1+\gamma}{2} + kx_1 \stackrel{(B.12), k \geq 0}{\geq} \frac{1+\gamma}{2} - k \geq 0. \quad (B.16)$$

Assumption 8 holds because a sum of submodular (supermodular) functions is a submodular (supermodular) function. We know that $\gamma \max\{x_1, x_2\}$ and $\gamma \min\{x_1, x_2\}$ are submodular and supermodular, respectively. Function $(1-\gamma)\frac{x_1+x_2}{2}$ is both submodular and supermodular. Finally, functions $\pm kx_1x_2$ are supermodular and submodular, respectively, by the remark in footnote 19.

Assumption 9 trivially holds.

Assumption 10 requires that the derivative

$$\begin{aligned} \frac{\partial^2}{\partial \Delta^2} \omega_s \left(v \left(\left| \frac{1-\Delta}{2} - t \right| \right), v \left(\frac{1+\Delta}{2} - t \right) \right) &= -\frac{\alpha(\alpha-1)}{8} \left((1+\gamma) \left(\left| \frac{1-\Delta}{2} - t \right| \right)^{\alpha-2} + (1-\gamma) \left(\frac{1+\Delta}{2} - t \right)^{\alpha-2} \right) \\ &\quad - \underbrace{\frac{\alpha k}{2} \left(\frac{(\alpha-1)\Delta^2}{2} + \left(t - \frac{1-\Delta}{2} \right) \left(\frac{1+\Delta}{2} - t \right) \right)}_{(\star)} \left(\left| \frac{1-\Delta}{2} - t \right| \right)^{\alpha-2} \left(\frac{1+\Delta}{2} - t \right)^{\alpha-2} \end{aligned} \quad (B.17)$$

is negative for all $t \in [0, \frac{1}{2}]$ and $\Delta \in [0, 1]$. If $(\star) \geq 0$, then expression (B.17) is negative because $\alpha > 1$, $k \geq 0$ and $0 \leq \gamma \leq 1$. If $(\star) < 0$ — which is possible only if $t < \frac{1-\Delta}{2}$ — then the sign of expression (B.17) coincides with the sign of

$$\frac{(B.17)}{\frac{\alpha(\alpha-1)}{8} \left(\frac{1-\Delta}{2} - t \right)^{\alpha-2} \left(\frac{1+\Delta}{2} - t \right)^{\alpha-2}} = -\frac{1+\gamma}{\left(\frac{1+\Delta}{2} - t \right)^{\alpha-2}} - \frac{1-\gamma}{\left(\frac{1-\Delta}{2} - t \right)^{\alpha-2}} - 2k\Delta^2 + \frac{4k}{\alpha-1} \left(\frac{1-\Delta}{2} - t \right) \left(\frac{1+\Delta}{2} - t \right). \quad (B.18)$$

Expression (B.18) is decreasing in $\alpha > 1$. Moreover, under restriction $\alpha \geq 2$, expression (B.18) is decreasing in $t \in (0, \frac{1-\Delta}{2})$. Hence, expression (B.18) achieves its maximum at $t = 0$ and $\alpha = 2$:

$$(B.18) \stackrel{\alpha \geq 2}{\leq} k - 2 - 3k\Delta^2 \stackrel{0 \leq k < 2}{<} 0. \quad (B.19)$$

Assumption 11 holds because the derivative

$$\frac{\partial}{\partial \Delta} \omega_s \left(v \left(\left| \frac{1-\Delta}{2} - t \right| \right), v \left(\frac{1+\Delta}{2} - t \right) \right) \Big|_{\Delta=0} = -\frac{\gamma}{2} v' \left(\frac{1}{2} - t \right) \quad (B.20)$$

is positive since $\gamma > 0$ and function v is decreasing.

With $\omega_s(x_1, x_2)$ defined in (B.7), Assumption 13 follows from Assumption 4, which holds for function $v(d) = -d^\alpha$ for all $\alpha > 1$. Indeed, the derivative

$$\begin{aligned} & \frac{\partial}{\partial t} \left(\frac{\frac{\partial}{\partial \Delta} \omega_s \left(v \left(\frac{1-\Delta}{2} - t \right), v \left(\frac{1+\Delta}{2} - t \right) \right)}{-\frac{\partial}{\partial \Delta} \left(v \left(\frac{1-\Delta}{2} - t \right) + v \left(\frac{1+\Delta}{2} - t \right) \right)} \right) \\ &= \frac{v' \left(\frac{1-\Delta}{2} - t \right) v' \left(\frac{1+\Delta}{2} - t \right)}{\underbrace{\left(v' \left(\frac{1-\Delta}{2} - t \right) - v' \left(\frac{1+\Delta}{2} - t \right) \right)^2}_{>0 \text{ by Assumptions 1 and 2}}} \underbrace{\left(\gamma + k \left(v \left(\frac{1-\Delta}{2} - t \right) - v \left(\frac{1+\Delta}{2} - t \right) \right) \right)}_{>0 \text{ by Assumption 1 and since } \gamma > 0 \text{ and } k \geq 0} \left(\frac{v'' \left(\frac{1+\Delta}{2} - t \right)}{v' \left(\frac{1+\Delta}{2} - t \right)} - \frac{v'' \left(\frac{1-\Delta}{2} - t \right)}{v' \left(\frac{1-\Delta}{2} - t \right)} \right) \end{aligned} \quad (\text{B.21})$$

is negative if $A(d) = \frac{v''(d)}{v'(d)}$ is strictly decreasing.

Assumption 14 holds:

$$\begin{aligned} & \left| \frac{d}{d\Delta} \left(\int_0^{0.5} \omega_s \left(v \left(\left| \frac{1-\Delta}{2} - t \right| \right), v \left(\frac{1+\Delta}{2} - t \right) \right) dt \right) \right|_{\Delta=0.5} \\ &= \frac{1+\gamma}{4} \left(\int_{1/4}^{1/2} v' \left(t - \frac{1}{4} \right) dt - \int_0^{1/4} v' \left(\frac{1}{4} - t \right) dt \right) + \frac{1-\gamma}{4} \int_0^{1/2} v' \left(\frac{3}{4} - t \right) dt \\ &+ \frac{k}{2} \left(\int_0^{1/4} \left\{ v' \left(\frac{1}{4} - t \right) v \left(\frac{3}{4} - t \right) - v \left(\frac{1}{4} - t \right) v' \left(\frac{3}{4} - t \right) \right\} dt - \int_{1/4}^{1/2} \left\{ v \left(t - \frac{1}{4} \right) v' \left(\frac{3}{4} - t \right) + v' \left(t - \frac{1}{4} \right) v \left(\frac{3}{4} - t \right) \right\} dt \right) \\ &= -\frac{1-\gamma}{4} \left(v \left(\frac{1}{4} \right) - v \left(\frac{3}{4} \right) \right) + \frac{k}{2} \int_0^{1/4} \left\{ v' \left(\frac{1}{4} - t \right) \left(v \left(\frac{3}{4} - t \right) - v \left(\frac{1}{4} + t \right) \right) - v \left(\frac{1}{4} - t \right) \left(v' \left(\frac{3}{4} - t \right) + v' \left(\frac{1}{4} + t \right) \right) \right\} dt \\ &= -\frac{1-\gamma}{4} \left(v \left(\frac{1}{4} \right) - v \left(\frac{3}{4} \right) \right) + \frac{k}{2} \left(v \left(\frac{1}{4} \right) \left(v \left(\frac{1}{4} \right) - v \left(\frac{3}{4} \right) \right) + 2 \int_0^{1/4} v' \left(\frac{1}{4} - t \right) \underbrace{\left(v \left(\frac{3}{4} - t \right) - v \left(\frac{1}{4} + t \right) \right)}_{\geq v \left(\frac{3}{4} \right) - v \left(\frac{1}{4} \right)} dt \right) \\ &\stackrel{v' < 0, k \geq 0}{\leq} -\frac{1-\gamma}{4} \left(v \left(\frac{1}{4} \right) - v \left(\frac{3}{4} \right) \right) + \frac{k}{2} \left(v \left(\frac{1}{4} \right) \left(v \left(\frac{1}{4} \right) - v \left(\frac{3}{4} \right) \right) + 2 \left(v \left(\frac{3}{4} \right) - v \left(\frac{1}{4} \right) \right) \int_0^{1/4} v' \left(\frac{1}{4} - t \right) dt \right) \\ &= -\frac{1}{2} \left(\frac{1-\gamma}{2} - 2k v \left(0 \right) + k v \left(\frac{1}{4} \right) \right) \left(v \left(\frac{1}{4} \right) - v \left(\frac{3}{4} \right) \right) \quad (\text{B.22}) \end{aligned}$$

is non-positive because $v(0) = 0$, $v(\frac{1}{4}) > v(\frac{3}{4}) > -1$ and $k \leq \frac{1-\gamma}{2}$.

B.5.2 Proof of Theorem 1*

Proof for $\beta \in [0, 0.5)$.

The proof for $\beta \in [0, 0.5)$ is an adapted version of the proof of Proposition 1 in Prat (2002).

For $\beta \leq 0.5$, function $\omega(x_1, x_2)$ is supermodular. Hence, for any t_1, t_2 and t ,

$$\omega(v(|t-t_1|), v(|t-t_2|)) + \omega(v(|t-t_2|), v(|t-t_1|)) \leq \omega(v(|t-t_1|), v(|t-t_1|)) + \omega(v(|t-t_2|), v(|t-t_2|)). \quad (\text{B.23})$$

By Assumption 7, function $\omega(x_1, x_2)$ is symmetric, which implies

$$\omega(v(|t-t_2|), v(|t-t_1|)) = \omega(v(|t-t_1|), v(|t-t_2|)). \quad (\text{B.24})$$

Substituting $\omega(v(|t - t_2|), v(|t - t_1|))$ from (B.24) into (B.23), taking the expectation of (B.23) over t and using the definition (2) of the utilitarian objective, we get

$$2U(t_1, t_2) \leq U(t_1, t_1) + U(t_2, t_2), \quad (\text{B.25})$$

which implies that

$$U(t_1, t_2) \leq \max\{U(t_1, t_1), U(t_2, t_2)\} \leq \max_{t \in [0,1]} U(t, t). \quad (\text{B.26})$$

Since (B.26) holds for any $0 \leq t_1 \leq t_2 \leq 1$ and the right-hand side of (B.26) does not depend on t_1 and t_2 , we can take the maximum over t_1 and t_2 of the left-hand side of (B.26) and get

$$\max_{0 \leq t_1 \leq t_2 \leq 1} U(t_1, t_2) \leq \max_{t \in [0,1]} U(t, t). \quad (\text{B.27})$$

The right-hand side of (B.27) is obviously less or equal the left-hand side of (B.27). Hence, inequality (B.27) must always hold as equality. Result (B.27) implies that there always exist t_1^* and t_2^* such that $t_1^* = t_2^*$ and they maximize $U(t_1, t_2)$ over $0 \leq t_1 \leq t_2 \leq 1$. By Lemma 1, $t = 0.5$ maximizes $U(t, t)$. Hence, $t_1^* = t_2^* = 0.5$ maximize $U(t_1, t_2)$ over $0 \leq t_1 \leq t_2 \leq 1$.

Proof for $\beta \in [0.5, 1]$.

Rewriting the expression (A.1) with $t_1 = (1 - \Delta)/2$ and $t_2 = (1 + \Delta)/2$, then using $U_{1m} = U_{m2}$ (see (A.3)) and $U_{01} = U_{21}$ (see (A.5) and (A.7) with $s = 1$), we get the following expression for the utilitarian objective

$$U(\Delta, \beta) = 2 \int_0^{1/2} \omega\left(v\left(\left|\frac{1-\Delta}{2} - t\right|\right), v\left(\frac{1+\Delta}{2} - t\right)\right) dt. \quad (\text{B.28})$$

Denote by $U_s(\Delta) = U(\Delta, 1)$ the utilitarian objective when the weight on the submodular component is 1, and by $U_a(\Delta) = U(\Delta, 0.5)$ the utilitarian objective for an additive task. Then, decomposition (B.4) implies that for any $\beta \in [0, 1]$, $U(\Delta, \beta)$ from (B.28) is a weighted sum of $U_s(\Delta)$ and $U_a(\Delta)$:

$$U(\Delta, \beta) = (2\beta - 1)U_s(\Delta) + 2(1 - \beta)U_a(\Delta), \quad (\text{B.29})$$

where

$$U_s(\Delta) = 2 \int_0^{1/2} \omega_s\left(v\left(\left|\frac{1-\Delta}{2} - t\right|\right), v\left(\frac{1+\Delta}{2} - t\right)\right) dt, \quad U_a(\Delta) = \int_0^{1/2} \left(v\left(\left|\frac{1-\Delta}{2} - t\right|\right) + v\left(\frac{1+\Delta}{2} - t\right)\right) dt. \quad (\text{B.30})$$

Differentiating $U_a(\Delta)$ yields

$$U_a'(\Delta) = \frac{1}{2} \left(v\left(\frac{1+\Delta}{2}\right) - v\left(\frac{1-\Delta}{2}\right) \right) \leq 0, \quad U_a''(\Delta) = \frac{1}{4} \left(v'\left(\frac{1+\Delta}{2}\right) + v'\left(\frac{1-\Delta}{2}\right) \right) < 0 \quad (\text{B.31})$$

for all $\Delta \in [0, 1]$. Both derivatives $U_a'(\Delta)$ and $U_a''(\Delta)$ are negative for all $\Delta > 0$ because v is decreasing by Assumption 1.

By Assumption 10, function $U_s(\Delta)$ is strictly concave for all $\Delta \in [0, 1]$. By (B.31), function $U_a(\Delta)$ is strictly concave for all $\Delta \in [0, 1]$. Thus, for $\beta \in [0.5, 1]$, function $U(\Delta, \beta)$ in (B.29) is strictly concave for all $\Delta \in [0, 1]$, which implies that it has a unique maximum $\Delta_U^*(\beta)$.

Suppose $\beta = 0.5$. By (B.31), function $U(\Delta, \beta) = U_a(\Delta)$ is strictly decreasing and, thus, it is uniquely maximized at $\Delta_U^*(0.5) = 0$.

Suppose $\beta \in (0.5, 1]$. Optimal diversity $\Delta_U^*(\beta)$ is positive because

$$\frac{\partial U(0, \beta)}{\partial \Delta} \stackrel{(B.29)}{=} (2\beta - 1)U'_s(0) + 2(1 - \beta)U'_a(0) \stackrel{U'_a(0)=0 \text{ by (B.31)}}{=} (2\beta - 1)U'_s(0) \quad (B.32)$$

is positive by Assumption 11. Optimal diversity $\Delta_U^*(\beta)$ is less than 1 because

$$U'_s(1) \stackrel{(B.30)}{=} 2 \int_0^{1/2} \underbrace{\frac{\partial}{\partial \Delta} \omega_s \left(v \left(t - \frac{1 - \Delta}{2} \right), v \left(\frac{1 + \Delta}{2} - t \right) \right)}_{\leq 0 \text{ by Assumptions 1 and 6}} dt \Big|_{\Delta=1} \leq 0, \quad (B.33)$$

and so,

$$\frac{\partial U(1, \beta)}{\partial \Delta} \stackrel{(B.29)}{=} (2\beta - 1)U'_s(1) + 2(1 - \beta)U'_a(1) \stackrel{U'_a(1) < 0 \text{ by (B.31)}}{<} (2\beta - 1)U'_s(1) \stackrel{(B.33)}{\leq} 0. \quad (B.34)$$

Since optimal diversity $\Delta_U^*(\beta)$ is interior, it satisfies the first order condition:

$$0 = \frac{\partial U(\Delta_U^*(\beta), \beta)}{\partial \Delta} \stackrel{(B.29)}{=} (2\beta - 1)U'_s(\Delta_U^*(\beta)) + 2(1 - \beta)U'_a(\Delta_U^*(\beta)). \quad (B.35)$$

Hence, since function $U(\Delta, \beta)$ is strictly concave in Δ , $\Delta_U^*(\beta)$ is strictly increasing if

$$0 < \frac{\partial^2 U(\Delta_U^*(\beta), \beta)}{\partial \Delta \partial \beta} \stackrel{(B.29)}{=} 2(U'_s(\Delta_U^*(\beta)) - U'_a(\Delta_U^*(\beta))). \quad (B.36)$$

Substituting $U'_s(\Delta_U^*(\beta))$ from (B.35), we get that inequality (B.36) holds because $U'_a(\Delta_U^*(\beta)) < 0$ by (B.31):

$$U'_s(\Delta_U^*(\beta)) - U'_a(\Delta_U^*(\beta)) \stackrel{U'_s \text{ from (B.35)}}{=} -\frac{2(1 - \beta)}{2\beta - 1}U'_a(\Delta_U^*(\beta)) - U'_a(\Delta_U^*(\beta)) \stackrel{U'_a < 0 \text{ by (B.31)}}{>} 0. \quad (B.37)$$

B.5.3 Proof of Theorem 2* (and Remark 2*)

By Lemma 1 for the generalized model, since Assumptions 1 and 5 to 7 hold, in the optimum, $t_1 = (1 - \Delta)/2$ and $t_2 = (1 + \Delta)/2$. Moreover, $R_{01} = R_{21}$ when $s = 1$ by (A.5) and (A.7), $R_{m2} = R_{1m}$ by (A.4). Hence, we can focus on $t \in [0, 0.5]$:

$$R(\Delta, \beta) = \min_{t \in [0, \frac{1}{2}]} \omega \left(v \left(\left| \frac{1 - \Delta}{2} - t \right| \right), v \left(\frac{1 + \Delta}{2} - t \right) \right). \quad (B.38)$$

By Assumptions 5 and 9, decomposition (B.4) holds. By decomposition (B.4), the objective function (B.38) can be rewritten as

$$\begin{aligned} & \omega \left(v \left(\left| \frac{1 - \Delta}{2} - t \right| \right), v \left(\frac{1 + \Delta}{2} - t \right) \right) \\ &= (2\beta - 1)\omega_s \left(v \left(\left| \frac{1 - \Delta}{2} - t \right| \right), v \left(\frac{1 + \Delta}{2} - t \right) \right) + (1 - \beta) \left(v \left(\left| \frac{1 - \Delta}{2} - t \right| \right) + v \left(\frac{1 + \Delta}{2} - t \right) \right). \end{aligned} \quad (B.39)$$

In (B.39), by Assumption 2, functions $v \left(\left| \frac{1 - \Delta}{2} - t \right| \right)$ and $v \left(\frac{1 + \Delta}{2} - t \right)$ are strictly concave in $t \in [0, \frac{1}{2}]$ for any fixed $\Delta \in [0, 1]$. By Assumption 10, function

$$\omega_s \left(v \left(\left| \frac{1 - \Delta}{2} - t \right| \right), v \left(\frac{1 + \Delta}{2} - t \right) \right) = \omega_s \left(v \left(\left| \frac{1 - \tilde{\Delta}}{2} - \tau \right| \right), v \left(\frac{1 + \tilde{\Delta}}{2} - \tau \right) \right) \quad \text{for } t = \frac{1 - \tilde{\Delta}}{2}, \Delta = 1 - 2\tau, \quad (B.40)$$

is strictly concave

- in $\tilde{\Delta} \in [0, 1]$ for any fixed $\tau \in [0, \frac{1}{2}]$,

or, equivalently, given the connection $t = \frac{1-\tilde{\Delta}}{2}$, and $\Delta = 1 - 2\tau$,

- in $t \in [0, \frac{1}{2}]$ for any fixed $\Delta \in [0, 1]$.

Hence, for any $\beta \in [0.5, 1]$, function (B.39) is strictly concave in $t \in [0, \frac{1}{2}]$ for any fixed $\Delta \in [0, 1]$. Therefore, the minimum in (B.38) is achieved either at $t = 0$ or at $t = \frac{1}{2}$:

$$R(\Delta, \beta) = \min \left\{ \omega \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right), \omega \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) \right\}, \quad (\text{B.41})$$

which is equivalent to (B.9).

Applying decomposition (B.4), we conclude that the first term in (B.41),

$$\omega \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right) \stackrel{(\text{B.4})}{=} (2\beta - 1) \omega_s \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right) + (1 - \beta) \left(v \left(\frac{1-\Delta}{2} \right) + v \left(\frac{1+\Delta}{2} \right) \right), \quad (\text{B.42})$$

is strictly concave on $\Delta \in [0, 1]$ because (1) $0.5 \leq \beta \leq 1$; (2) function $\omega_s \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right)$ is strictly concave in $\Delta \in [0, 1]$ by Assumption 10 (applied with $t = 0$); and (3) functions $v \left(\frac{1-\Delta}{2} \right)$ and $v \left(\frac{1+\Delta}{2} \right)$ are strictly concave in $\Delta \in [0, 1]$ by Assumption 2. Hence, there exists $\hat{\Delta}(\beta) \in [0, 1]$ such that the first term in (B.41) is strictly increasing in $\Delta \in [0, \hat{\Delta}(\beta)]$ and strictly decreasing in $\Delta \in (\hat{\Delta}(\beta), 1]$.

Applying decomposition (B.4), we conclude that the second term in (B.41),

$$\omega \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) \stackrel{(\text{B.4})}{=} (2\beta - 1) \omega_s \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) + 2(1 - \beta) v \left(\frac{\Delta}{2} \right), \quad (\text{B.43})$$

is strictly decreasing on $\Delta \in [0, 1]$ because (1) $0.5 \leq \beta \leq 1$; (2) function $\omega_s \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right)$ is strictly decreasing in $\Delta \in [0, 1]$ because it is strictly concave in $\Delta \in [0, 1]$ by Assumption 10 (applied with $t = \frac{1}{2}$), and because it is weakly decreasing in $\Delta \in [0, 1]$ by Assumptions 1 and 6; and (3) function $v \left(\frac{\Delta}{2} \right)$ is strictly decreasing in $\Delta \in [0, 1]$ by Assumption 1.

Thus, both terms in (B.41) are strictly decreasing in $\Delta \in (\hat{\Delta}(\beta), 1]$, which implies that

$$R(\hat{\Delta}(\beta), \beta) > R(\Delta, \beta) \quad \text{for all } \Delta \in (\hat{\Delta}(\beta), 1]. \quad (\text{B.44})$$

Define function $F(\Delta, \beta)$ as the difference between the first and the second terms in (B.41):

$$\begin{aligned} F(\Delta, \beta) &= \omega \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right) - \omega \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) \\ &\stackrel{(\text{B.42}), (\text{B.43})}{=} (2\beta - 1) \left\{ \omega_s \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right) - \omega_s \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) \right\} + (1 - \beta) \left\{ v \left(\frac{1-\Delta}{2} \right) + v \left(\frac{1+\Delta}{2} \right) - 2v \left(\frac{\Delta}{2} \right) \right\}. \end{aligned} \quad (\text{B.45})$$

Since on $\Delta \in [0, \hat{\Delta}(\beta)]$, the first term in (B.41) is strictly increasing and the second term in (B.41) is strictly decreasing,

- (i) $F(\Delta, \beta)$ is strictly increasing in $\Delta \in [0, \hat{\Delta}(\beta)]$,

and there exists $\Delta_R^*(\beta) \in [0, \hat{\Delta}(\beta)]$ such that

- (ii) for $\Delta \in [0, \Delta_R^*(\beta)]$, $F(\Delta, \beta) < 0$ and $R(\Delta, \beta)$ is equal to the first term and strictly increasing in Δ ;

(iii) for $\Delta \in (\Delta_R^*(\beta), \hat{\Delta}(\beta)]$, $F(\Delta, \beta) > 0$ and $R(\Delta, \beta)$ is equal to the second term and strictly decreasing in Δ .

This $\Delta_R^*(\beta)$ uniquely maximizes $R(\Delta, \beta)$ on $\Delta \in [0, \hat{\Delta}(\beta)]$ and, therefore, by (B.44), on $\Delta \in [0, 1]$.

Suppose $\beta = 0.5$. Then, expression (B.42) achieves its maximum at $\hat{\Delta}(0.5) = 0$ because

$$\left. \frac{d}{d\Delta} \left(v \left(\frac{1-\Delta}{2} \right) + v \left(\frac{1+\Delta}{2} \right) \right) \right|_{\Delta=0} = 0. \quad (\text{B.46})$$

Hence, $\Delta_R^*(0.5) = \hat{\Delta}(0.5) = 0$.

Suppose $\beta \in (0.5, 1]$.

Then, $\hat{\Delta}(\beta) > 0$. Indeed, the first term in (B.41), as written in (B.42), is strictly increasing at $\Delta = 0^+$ by (B.46) and because $\omega_s \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right)$ is strictly increasing at $\Delta = 0^+$ by Assumption 11, applied with $t = 0$. Hence, $\hat{\Delta}(\beta) > 0$.

As long as $0 < \hat{\Delta}(\beta) < 1$, — that is, $\hat{\Delta}(\beta)$ is an interior maximum of (B.42), — $\hat{\Delta}(\beta)$ is strictly increasing in β because the derivative

$$\frac{\partial^2 (\text{B.42})}{\partial \beta \partial \Delta} = \frac{2}{2\beta - 1} \underbrace{\frac{\partial (\text{B.42})}{\partial \Delta}}_{=0 \text{ at } \Delta = \hat{\Delta}(\beta)} + \frac{1}{2(2\beta - 1)} \underbrace{\left(v' \left(\frac{1-\Delta}{2} \right) - v' \left(\frac{1+\Delta}{2} \right) \right)}_{>0 \text{ for } \Delta > 0 \text{ by Assumption 2}} \quad (\text{B.47})$$

is positive at $\Delta = \hat{\Delta}(\beta)$.

For $\beta \in (0.5, 1]$, $\Delta_R^*(\beta) > 0$. Indeed, the first term in (B.41) at $\Delta = 0$ is equal to (B.43) at $\Delta = 1$; the second term in (B.41) at $\Delta = 0$ is equal to (B.43) at $\Delta = 0$; expression (B.43) is strictly decreasing on $\Delta \in [0, 1]$ and, thus, (B.43) at $\Delta = 1$ is strictly less than at $\Delta = 0$. Hence, $F(0, \beta) < 0$, which means that $\Delta_R^*(\beta) > 0$ by (iii).

As long as $0 < \Delta_R^*(\beta) < \hat{\Delta}(\beta)$, $\Delta_R^*(\beta)$ is strictly decreasing in β . Indeed, $0 < \Delta_R^*(\beta) < \hat{\Delta}(\beta)$ implies that $F(\Delta, \beta) = 0$ at $\Delta = \Delta_R^*(\beta)$. Hence, by the implicit function theorem,

$$\frac{d\Delta_R^*(\beta)}{d\beta} = - \frac{\partial F(\Delta, \beta)}{\partial \beta} \bigg/ \frac{\partial F(\Delta, \beta)}{\partial \Delta} \bigg|_{\Delta = \Delta_R^*(\beta)}. \quad (\text{B.48})$$

The derivative $\frac{\partial F(\Delta, \beta)}{\partial \Delta}$ is positive by (i). The derivative

$$\frac{\partial F(\Delta, \beta)}{\partial \beta} \stackrel{(\text{B.45})}{=} \frac{2}{2\beta - 1} \underbrace{F(\Delta, \beta)}_{=0 \text{ at } \Delta = \Delta_R^*(\beta)} + \frac{1}{2\beta - 1} \left\{ \underbrace{2v \left(\frac{\Delta}{2} \right)}_{\geq 2v(\frac{1}{2}) \text{ by Assumption 1}} \underbrace{- v \left(\frac{1-\Delta}{2} \right) - v \left(\frac{1+\Delta}{2} \right)}_{\substack{\text{increasing by Assumption 2} \\ \Rightarrow > -2v(\frac{1}{2}) \text{ for } \Delta > 0}} \right\} \quad (\text{B.49})$$

is positive at $\Delta = \Delta_R^*(\beta)$. Hence, $\Delta_R^*(\beta)$ is strictly decreasing in β .

Moreover, by (iii), if $\Delta_R^*(\beta) < \hat{\Delta}(\beta)$, then $\Delta_R^*(\beta) \geq 0.5$ because $F(\Delta, \beta) \leq 0$ for all $\Delta \in [0, 0.5]$. Indeed, whenever $F(\Delta, \beta) \geq 0$, the derivative (B.49) is non-negative for all $\beta \in (0.5, 1]$ and $\Delta \in [0, 1]$. At $\beta = 1$, $F(\Delta, \beta) \leq 0$ for all $\Delta \in [0, 0.5]$:

$$F(\Delta, 1) \stackrel{(\text{B.45})}{=} \omega_s \left(v \left(\frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} \right) \right) - \omega_s \left(v \left(\frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} \right) \right) \leq 0 \quad (\text{B.50})$$

because (1) $\frac{1-\Delta}{2} \geq \frac{\Delta}{2}$ and $\frac{1+\Delta}{2} \geq \frac{\Delta}{2}$ for all $\Delta \in [0, 0.5]$; (2) function v is strictly decreasing by Assumption 1; and (3) function ω_s is weakly increasing in each argument by Assumption 6. Hence, $F(\Delta, \beta) \leq 0$ for all $\Delta \in [0, 0.5]$ and all $\beta \in (0.5, 1]$.

Consider the difference $\Delta_R^*(\beta) - \hat{\Delta}(\beta)$. By construction, it cannot be positive. Once this difference is negative, it is strictly decreasing because (1) $\Delta_R^*(\beta)$ is strictly decreasing whenever $\Delta_R^*(\beta) < \hat{\Delta}(\beta)$; and (2) $\hat{\Delta}(\beta)$ is weakly increasing, as $\hat{\Delta}(\beta)$ is strictly increasing as long as it is less than 1. Hence, there exists $\beta^* \in [0.5, 1]$ such that

- $\Delta_R^*(\beta) = \hat{\Delta}(\beta)$ for all $\beta \in (0.5, \beta^*)$,
- $\Delta_R^*(\beta) < \hat{\Delta}(\beta)$ for all $\beta \in (\beta^*, 1)$.

Moreover, $\beta^* > 0.5$ because $\Delta_R^*(\beta) - \hat{\Delta}(\beta) = 0$ at $\beta = 0.5$.

To finish the proof, it is sufficient to show that $\beta^* < 1$ under Assumption 12. Assumption 12 says that $F(\Delta, 1) > 0$ at $\Delta = \hat{\Delta}(1)$. Hence, by (ii), $\Delta_R^*(1) < \hat{\Delta}(1)$, which means that $\beta^* < 1$.

B.5.4 Proof of Theorem 3*

The proof of Theorem 3* closely follows the proof of Theorem 3 in Appendix A.6.

Step 1. Defining $\Delta^*(\delta, \beta)$.

For $\delta \in (0, 0.5)$, define a δ -objective as the expected team output, given that $t_1 = \frac{1-\Delta}{2}$, $t_2 = \frac{1+\Delta}{2}$, and t is distributed uniformly on $[0, 0.5 - \delta] \cup [0.5 + \delta, 1]$:

$$U(\Delta, \delta, \beta) = \frac{2}{1-2\delta} \int_0^{0.5-\delta} \omega \left(\left| \frac{1-\Delta}{2} - t \right| \right), v \left(\frac{1+\Delta}{2} - t \right) dt. \quad (\text{B.51})$$

Denote by $U_s(\Delta, \delta) = U(\Delta, \delta, 1)$ the δ -objective when the weight on the submodular component is 1, and by $U_a(\Delta, \delta) = U(\Delta, \delta, 0.5)$ the δ -objective for an additive task. Then, decomposition (B.4) implies that for any $\beta \in [0, 1]$, $U(\Delta, \delta, \beta)$ is a weighted sum of $U_s(\Delta, \delta)$ and $U_a(\Delta, \delta)$:

$$U(\Delta, \delta, \beta) = (2\beta - 1) U_s(\Delta, \delta) + 2(1 - \beta) U_a(\Delta, \delta). \quad (\text{B.52})$$

Differentiating (B.52) yields

$$(2\beta - 1) \frac{\partial U_s(\Delta, \delta)}{\partial \Delta} + 2(1 - \beta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} = 0 \quad (\text{B.53})$$

as the first order condition for maximizing (B.51). The solution $\Delta^*(\delta, \beta) \in (0, 1)$ to (B.53) exists, is unique and gives the maximum of $U(\Delta, \delta, \beta)$ by Assumptions 10 and 11. Indeed, function $U_s(\Delta, \delta)$ is strictly concave in Δ by Assumption 10 and its derivative is positive at $\Delta = 0$ by Assumption 11. At $\Delta = 1$, the derivative

$$\frac{\partial U_s(1, \delta)}{\partial \Delta} = \frac{2}{1-2\delta} \int_0^{0.5-\delta} \frac{\partial}{\partial \Delta} \omega_s \left(v \left(t - \frac{1-\Delta}{2} \right), v \left(\frac{1+\Delta}{2} - t \right) \right) \Big|_{\Delta=1} dt \quad (\text{B.54})$$

is non-positive because $\delta < 0.5$, ω_s is weakly increasing in each argument by Assumption 6, and v is decreasing by Assumption 1. At $\beta = 0.5$, $\omega(x_1, x_2) = 0.5(x_1 + x_2)$ is the same as in the benchmark model, which implies that function $U_a(\Delta, \delta)$ has all the properties described in Appendix A.6:

$$\frac{\partial U_a(\Delta, \delta)}{\partial \Delta} = \frac{1}{2(1-2\delta)} \left(v \left(\left| \frac{\Delta}{2} - \delta \right| \right) - v \left(\frac{1-\Delta}{2} \right) - v \left(\frac{\Delta}{2} + \delta \right) + v \left(\frac{1+\Delta}{2} \right) \right), \quad (\text{B.55})$$

$$\frac{\partial U_a(0, \delta)}{\partial \Delta} = 0, \quad \frac{\partial U_a(1, \delta)}{\partial \Delta} < 0, \quad \frac{\partial^2 U_a(\Delta, \delta)}{\partial \Delta^2} < 0 \quad (\text{B.56})$$

(see (A.55), (A.57), (A.58), (A.59) and (A.60) for $\beta = 0.5$). Hence, for any $0.5 < \beta < 1$ and $0 < \delta < 0.5$, function $U(\Delta, \delta, \beta)$ is strictly concave in $\Delta \in (0, 1)$, increasing at $\Delta = 0$ and decreasing at $\Delta = 1$ — all of which guarantee that the solution $\Delta^*(\delta, \beta) \in (0, 1)$ to (B.53) exists, is unique and gives the maximum of $U(\Delta, \delta, \beta)$ on $\Delta \in [0, 1]$.

Step 2. Function $\Delta^*(\delta, \beta)$ is increasing in $\delta \in (0, 0.5)$.

By the implicit function theorem, the sign of $\frac{\partial \Delta^*(\delta, \beta)}{\partial \delta}$ coincides with the sign of

$$\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U(\Delta, \delta, \beta)}{\partial \Delta} \right) \stackrel{(\text{B.52})}{=} (2\beta-1) \frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_s(\Delta, \delta)}{\partial \Delta} \right) + 2(1-\beta) \frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} \right) \quad (\text{B.57})$$

at $\Delta = \Delta^*(\delta, \beta)$, where

$$\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_s(\Delta, \delta)}{\partial \Delta} \right) = -2 \frac{\partial}{\partial \Delta} \omega_s \left(v \left(\left| \delta - \frac{\Delta}{2} \right| \right), v \left(\frac{\Delta}{2} + \delta \right) \right) \quad (\text{B.58})$$

and

$$2 \frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} \right) = \begin{cases} -v'(\Delta/2 - \delta) - v'(\Delta/2 + \delta), & \delta < \Delta/2, \\ v'(\delta - \Delta/2) - v'(\Delta/2 + \delta), & \delta > \Delta/2 \end{cases} \quad (\text{B.59})$$

(see (A.61) for $\beta = 0.5$). For $\delta < \Delta/2$, expression (B.58) is non-negative and expression (B.59) is positive — and, thus, expression (B.57) is positive — because $0.5 < \beta < 1$, ω_s is weakly increasing in each argument by Assumption 6, and v is strictly decreasing by Assumption 1.

To find the sign of (B.57) at $\Delta = \Delta^*$ for $\delta > \Delta^*/2$, we express β from (B.53):

$$\frac{1-\beta}{2\beta-1} = -\frac{1}{2} \frac{\frac{\partial U_s(\Delta^*, \delta)}{\partial \Delta}}{\frac{\partial U_a(\Delta^*, \delta)}{\partial \Delta}}. \quad (\text{B.60})$$

Substituting (B.60) into (B.57), we get

$$\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U(\Delta^*, \delta, \beta)}{\partial \Delta} \right) = (2\beta-1) \left\{ \frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_s(\Delta^*, \delta)}{\partial \Delta} \right) - \frac{\frac{\partial U_s(\Delta^*, \delta)}{\partial \Delta}}{\frac{\partial U_a(\Delta^*, \delta)}{\partial \Delta}} \frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_a(\Delta^*, \delta)}{\partial \Delta} \right) \right\}. \quad (\text{B.61})$$

Since $\frac{\partial U_a(\Delta^*, \delta)}{\partial \Delta}$ is negative by (B.56), and since $\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_a(\Delta^*, \delta)}{\partial \Delta} \right) \stackrel{(\text{B.59})}{=} \frac{1}{2} v' \left(\delta - \frac{\Delta^*}{2} \right) - \frac{1}{2} v' \left(\frac{\Delta^*}{2} + \delta \right)$ is positive because v is strictly concave by Assumption 2, the sign of (B.61) coincides with the sign of

$$g(\Delta, \delta) = \left((1-2\delta) \frac{\partial U_s(\Delta, \delta)}{\partial \Delta} \right) - \left((1-2\delta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} \right) \frac{\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_s(\Delta, \delta)}{\partial \Delta} \right)}{\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} \right)} \quad (\text{B.62})$$

at $\Delta = \Delta^*$. In sum, to prove that expression (B.57) is positive at $\Delta = \Delta^*$ for $\delta > \Delta^*/2$, it is sufficient to show that $g(\Delta, \delta)$ is positive for all $1/2 > \delta > \Delta/2 > 0$. Function $g(\Delta, \delta)$ is indeed positive in this range because the derivative

$$\begin{aligned} \frac{\partial g(\Delta, \delta)}{\partial \delta} &= -(1-2\delta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} \frac{\partial}{\partial \delta} \left(\frac{\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_s(\Delta, \delta)}{\partial \Delta} \right)}{\frac{\partial}{\partial \delta} \left((1-2\delta) \frac{\partial U_a(\Delta, \delta)}{\partial \Delta} \right)} \right) \\ &\stackrel{(\text{B.58}), (\text{B.59})}{=} 2(1-2\delta) \underbrace{\frac{\partial U_a(\Delta, \delta)}{\partial \Delta}}_{<0 \text{ by (B.56)}} \underbrace{\frac{\partial}{\partial \delta} \left(\frac{\frac{\partial}{\partial \Delta} \omega_s \left(v \left(\delta - \frac{\Delta}{2} \right), v \left(\frac{\Delta}{2} + \delta \right) \right)}{-\frac{\partial}{\partial \Delta} \left(v \left(\delta - \frac{\Delta}{2} \right) + v \left(\frac{\Delta}{2} + \delta \right) \right)} \right)}_{>0 \text{ by Assumption 13 with } t = 0.5 - \delta} \end{aligned} \quad (\text{B.63})$$

is negative, and at $\delta = 0.5$, $g(\Delta, \delta)$ is zero because by (B.51), $(1 - 2\delta) \frac{\partial U(\Delta, \delta, \beta)}{\partial \Delta} = 0$ as $\delta \rightarrow 0.5$ for any β , including $\beta = 0.5$ and $\beta = 1$.

Step 3. Main argument for $\Delta_R^*(\beta) > \Delta_U^*(\beta)$.

The argument here is almost the same as in Step 3 in Appendix A.6. The only difference is that in the general model, we do not necessarily have $\Delta_R^*(1) = \Delta_U^*(1) = 0.5$, and so, the argument for $\beta \in (\beta^*, 1)$ is slightly different.

Since $\Delta_R^*(\beta) > 0.5$ for all $\beta \in (\beta^*, 1)$, to ensure $\Delta_R^*(\beta) > \Delta_U^*(\beta)$ for all $\beta \in (\beta^*, 1)$, it is sufficient to require $\Delta_U^*(\beta) < 0.5$ for all $\beta \in (\beta^*, 1)$. Assumption 14 ensures $\Delta_U^*(1) \leq 0.5$ and the monotonicity of $\Delta_U^*(\beta)$ does the rest.