

Dynamic project selection

ARINA NIKANDROVA

Department of Economics, Mathematics, and Statistics, Birkbeck, University of London

ROMANS PANCS

Centro de Investigación Económica, Instituto Tecnológico Autónomo de México

KEYWORDS. Internal capital market, irreversible project selection.

JEL CLASSIFICATION. D82, D83, G310, G320.

1. INTRODUCTION

Corporate finance textbooks (e.g., Webster 2003, Chapter 12) recommend that a company invest in a project if and only if the project's internal rate of return exceeds the cost of capital. If companies operated this way, their investment decisions would be independent across projects within the company, conditional on the projects' cash flows being independent. In practice, such conditional independence is the exception rather than the rule (Ozbas and Scharfstein 2010).

Investment decisions on projects with independent cash flows can be dependent for two reasons: projects may be mutually exclusive or the internal capital used to finance these projects may be scarce. We use the term “internal capital market” to describe a project-selection mechanism that deals with either situation. We are interested in the design of an optimal internal capital market for environments with independent cash flows.

We focus on the problem in which project values are initially unknown but can be learned over time. Before deciding which project to finance, a company performs due diligence on each project. The Universal Music Group faced such a situation in 2011. Universal was considering two alternative projects: the purchase of EMI Music or the purchase of Warner Music Group.¹ Purchasing both was infeasible, if only because of

Arina Nikandrova: a.nikandrova@gmail.com

Romans Páncs: rpancs@gmail.com

We thank Jeanine Miklos-Thal, Dmitry Orlov, Vladimir Parail, Marek Pycia, Juuso Toikka, and Michael Waldman for valuable input at various stages of this project. We also thank the editor for precise editorial guidance and anonymous referees for detailed comments and suggestions. This paper would not have been written without George Georgiadis's incisive modeling suggestion, made while Páncs was on sabbatical at UCLA. For their feedback, we are also grateful to seminar audiences in Rochester, Birkbeck, ITAM, Bielefeld, the University of Surrey, the Australian National University, the University of Sydney, and the University of Queensland.

¹See <http://www.completemusicupdate.com/article/bmg-and-universal-may-co-bid-for-warner-and-emi/> and <http://www.completemusicupdate.com/timeline-emisale/>.

antitrust concerns. Assessing the profitability of each purchase required costly due diligence by teams of lawyers, consultants, and accountants, who evaluated music catalogs, potential synergies, and antitrust risks.

Universal has two divisions: in London and in New York. Since EMI is based in London and Warner Music in New York, Universal could have charged the London division with performing due diligence on the purchase of EMI and the New York division with performing due diligence on Warner Music. We ask how Universal should have orchestrated its divisions' due diligence to maximize its expected cash flow.

We study Universal's problem in an auction-like environment in which HQ (the headquarters) allocates an item (the requisite funds to pursue an acquisition) to one of two divisions, denoted by D_1 and D_2 . The value of each division's project (the profitability of the acquisition) is either 0 or 1, and is distributed independently across the two divisions. Initially, each division has a belief about its project's value and revises this belief as it learns (performs due diligence). At each instant, each division can learn at a cost. Learning affects the arrival intensity of "good news," which reveals the project's value to be 1. The alternative, "no news," means that the project's value can be either 0 or 1 and causes the division to revise its value estimate downward.

HQ maximizes the expected cash flow, defined as the expected value of the winning project net of both divisions' expected cumulative costs of learning. HQ can directly control each division's learning, observe learning outcomes, and select the winning project. Thus, HQ's problem is a stochastic-control optimal-stopping problem. This problem's solution—an optimal policy—is this paper's contribution.

Figure 1 summarizes the optimal policy we identify (we claim no uniqueness) when learning is cheap, which is the most interesting case. This policy is stationary and prescribes, for every pair (x_1, x_2) of the two projects' expected values, whether either division wins immediately and, if not, which division learns. Normalizing $x_2 \geq x_1$, four prescriptions are possible:

1. *Division 2 wins immediately.* D_2 wins immediately whenever x_1 and x_2 are either both close to 0 or both close to 1. In this case, since there is little uncertainty about each project's value (probably the same), learning is not worth the cost. D_2 also wins immediately whenever x_2 is substantially larger than x_1 . In this case, since there is little uncertainty about the fact that project 2's value exceeds project 1's value, learning is unlikely to affect the decision regarding which project to select.
2. *Division 2 learns.* D_2 learns alone when $x_1 \neq x_2$, when x_1 and x_2 are close to each other (so that it is highly uncertain which project is more valuable), and when both x_1 and x_2 are far away from 0 and 1 (so that each project's value is highly uncertain). Under these conditions, the need for information is so great that it is worthwhile to ask D_2 to learn first and to plan to ask D_1 to learn later if D_2 does not observe good news. (Asking D_2 to learn without ever planning to ask D_1 is suboptimal, for such learning, while costly, would not affect the optimal allocation.)

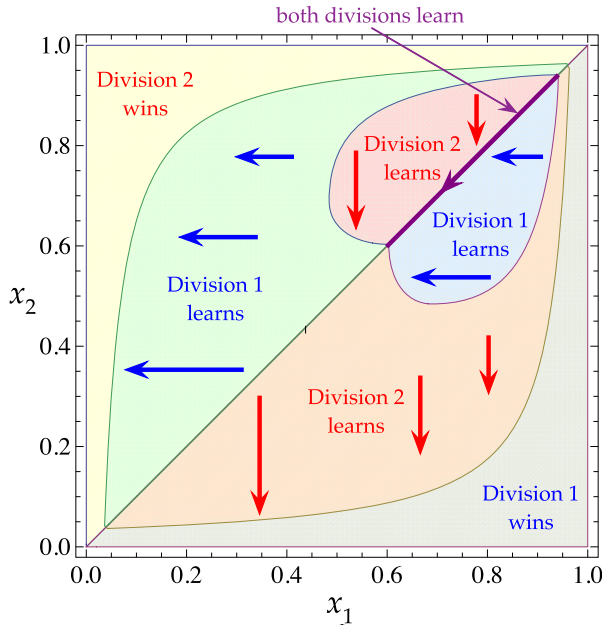


FIGURE 1. The optimal policy's prescription for each pair (x_1, x_2) of the project's expected values. The arrows indicate the direction in which the type profile is revised if the division that learns observes no good news.

3. *Both divisions learn.* Both divisions learn simultaneously when $x_1 = x_2$ (i.e., both projects appear equally valuable), when x_1 and x_2 are sufficiently large (i.e., learning by either division is informative), and when x_1 and x_2 are bounded away from 0 and 1 (i.e., each project's value is sufficiently uncertain).
4. *Division 1 learns.* D₁ learns alone when the values of x_1 and x_2 are complementary to those described in cases 1–3. Then, by asking D₁ to learn, HQ bets on having D₁ observe the good news. HQ is “insured” by D₂, which does not learn and whose project can be selected if D₁ observes no good news for a sufficiently long time.

The rest of the paper is structured as follows. This section concludes with a literature review. Section 2 sets up the problem. Section 3 solves for the optimal policy. Section 4 introduces and maps the bad-news technology case into the analyzed good-news technology case, thereby establishing that the derived results immediately apply to the former case as well. Section 5 shows numerically that the effect of exponential discounting is similar to that of costly learning, and that the introduction of discounting does not qualitatively change the optimal policy for the undiscounted case. Section 6 concludes. Auxiliary technical lemmas are provided in the Appendix, available in a supplementary file on the journal website, <http://econtheory.org/supp/2379/supplement.pdf>.

Related literature. Our paper contributes to two literatures: the corporate finance literature on internal capital markets and the economic theory literature on irreversible project selection in the presence of uncertainty. The assumptions underlying our model

of the internal capital market are motivated by the vision described by [Stein \(1997\)](#).² In particular, because internal capital is scarce (e.g., because of informational frictions associated with raising outside capital), not all profitable projects can be financed and, so, HQ must ration. At the same time, even unprofitable projects may end up being financed (e.g., because of HQ's empire-building tendencies); thus, HQ invests all available internal capital. Accordingly, we assume that HQ selects exactly one project.

The existing literature on internal capital markets is predominantly positive. In addition to [Stein \(1997\)](#), it includes [Harris and Raviv \(1996\)](#), [Rajan et al. \(2000\)](#), [Scharfstein and Stein \(2000\)](#), [de Motta \(2003\)](#), and [Inderst and Laux \(2005\)](#). The only normative dynamic model of an internal capital market that we are aware of is that of [Malenko \(2012\)](#). While we focus on learning about, and selecting between, two given projects, [Malenko \(2012\)](#) studies selection from dynamically arriving projects and does not model learning.

The economic theory literature on irreversible project selection can be interpreted as modeling internal capital markets. The real-option approach, exemplified by the work of [Dixit and Pindyck \(1994\)](#), assumes that the values of projects evolve exogenously. We extend their approach to situations in which these values evolve endogenously, as a result of learning.

Learning is the focus of multi-armed bandit problems ([Bolton and Harris 1999](#), [Keller et al. 2005](#), [Klein and Rady 2011](#), [Forand 2015](#)), which model reversible project selection.³ A solution to a bandit problem is typically an index policy that always selects the arm with the greatest value of the Gittins index. In particular, in an exponential-bandit problem with two risky arms, the optimal policy prescribes selecting myopically the project with the highest expected value.⁴ The analogous policy of always learning about the project with the highest expected value is suboptimal in our setting.

Our problem is related to sequential hypothesis testing, first formulated by [Wald \(1973\)](#). [Shiryaev \(2008, Chapter 4\)](#), [Peskir and Shiryaev \(2006, Section 23\)](#), and [Presman and Sonin \(1990\)](#) provide modern textbook accounts. Our problem shares two critical features with these testing problems. First, in our model, learning (the analogue of testing) has an explicit flow cost. Second, this cost is no longer incurred as soon as—at some optimally chosen time—one of the two projects is chosen.^{5,6}

2. MODEL

Time is continuous and indexed by $t \geq 0$. The time horizon is infinite.

²For a textbook introduction to internal capital markets, see [Tirole \(2006, Section 10.5\)](#).

³In bandit problems, the irreversibility of selecting an arm—but only a safe arm—is explored by [Murto and Välimäki \(2011\)](#) and [Rosenberg et al. \(2007\)](#).

⁴[Banks and Sundaram \(1992\)](#) show that myopic strategies are uniquely optimal in the class of bandit problems in which each of the independent arms generates rewards according to one of two reward distributions (same for both arms). By contrast, [Forand \(2015\)](#) studies a bandit-like problem with maintenance costs and finds, just as we do, that a decision maker may sometimes optimally pull a less auspicious arm.

⁵This switching off of the learning costs upon selecting a project also prevents us from mapping our problem into a multi-armed bandit problem.

⁶[Che and Mierendorff \(2017\)](#) study a sequential Wald testing problem in a Poisson environment that resembles ours but assumes perfectly negatively correlated projects.

Valuations. HQ holds an indivisible item and values it at zero. HQ allocates this item to one of two divisions, indexed by $i \in \mathcal{N} \equiv \{1, 2\}$ and denoted by D_i . D_i 's valuation, $v_i \in \{0, 1\}$, is a random variable with $\Pr\{v_i = 1\} = X_i(0)$ for some prior belief $X_i(0) \in [0, 1]$. Valuations v_1 and v_2 are independent.

Learning. Each D_i can acquire information about v_i , that is, can *learn*.⁷ At each time t , HQ allocates a unit of learning intensity between the divisions. D_i 's learning intensity is denoted by $a_i(t) \in [0, 1]$, with $a_1(t) + a_2(t) = 1$. The cumulative cost of learning incurred by D_i up to time t is $c \int_0^t e^{-rs} a_i(s) ds$, for some cost parameter $c > 0$ and discount rate $r \geq 0$.

D_i 's learning process $\{a_i(t) \mid t \geq 0\}$, denoted by a_i , controls the arrival-intensity process $\{a_i(t)v_i \mid t \geq 0\}$ of a Poisson process $\{N_i^{a_i}(t) \mid t \geq 0\}$, $N_i^{a_i}(0) = 0$.⁸ Processes $N_1^{a_1}$ and $N_2^{a_2}$ are independent. The public event when $N_i^{a_i}(t)$ is incremented is called *good news* (about v_i). The event when $N_i^{a_i}(t)$ is not incremented is called *no news*. Because event $N_i^{a_i}(t) > 0$ can occur only if $v_i = 1$, the good news reveals $v_i = 1$.

Define $X_i^{a_i}(t)$, D_i 's time- t *type*, or *belief*, to be the expectation of v_i conditional on the information revealed up to time t under some learning process a_i :

$$X_i^{a_i}(t) \equiv \mathbb{E}[v_i \mid \{N_i^{a_i}(s) \mid 0 \leq s \leq t\}].$$

For any learning-process profile $a \equiv (a_1, a_2)$, the tuple $X^a(t) \equiv (X_1^{a_1}(t), X_2^{a_2}(t))$ is a time- t *type profile*. By the law of iterated expectations, X^a is a martingale.

For any dates t and $t' > t$, D_i 's type $X_i^{a_i}(t')$ is derived from $X_i^{a_i}(t)$ by application of Bayes rule. By Bayes rule, $N_i^{a_i}(t') > 0$ implies $X_i^{a_i}(t') = 1$, whereas $N_i^{a_i}(t') = 0$ implies

$$\frac{X_i^{a_i}(t')}{1 - X_i^{a_i}(t')} = \frac{X_i^{a_i}(t)}{1 - X_i^{a_i}(t)} e^{-\int_t^{t'} a_i(s) ds}. \quad (1)$$

An optimal policy. The environment is stationary, so no generality is lost by focusing on stationary policies. A (stationary) *policy* is a tuple (α, τ) , where a learning policy $\alpha \equiv (\alpha_1, \alpha_2)$ maps a type profile $x \equiv (x_1, x_2)$ into learning decisions $(\alpha_1(x), \alpha_2(x)) \in [0, 1]^2$ with $\alpha_1(x) + \alpha_2(x) = 1$, and where τ is a stopping time that designates when the item is allocated, always to the higher-type division. A policy (α, τ) induces the type process denoted by $\{X^{\alpha, \tau}(t) \mid t \geq 0\}$.

A policy (α, τ) is *admissible* if the learning process $\{\alpha(X(t)) \mid t \geq 0\}$, induced by the learning policy α , is predictable and integrable, and if, for every $i \in \mathcal{N}$ and every $X_i(0) \in [0, 1]$, the appropriately defined stochastic differential equation for the evolution of the type process has a unique strong solution.

A policy (α, τ) and an initial type profile x induce HQ's expected discounted *cash flow*,

$$J^r(x, \alpha, \tau) \equiv \mathbb{E} \left[-c \int_0^\tau e^{-rs} ds + e^{-r\tau} \max_{i \in \mathcal{N}} \{X_i^{\alpha, \tau}(\tau)\} \mid X^{\alpha, \tau}(0) = x \right],$$

⁷Here and throughout, the *Italic* typeface highlights a definition.

⁸Henceforth, superscript a_i indicates that the superscripted process is conditional on learning process a_i .

where the expectation is with respect to the induced type process $\{X^{\alpha,\tau}(t) \mid t \geq 0\}$. For any initial type profile x , the *value function* is defined by

$$\phi^r(x) \equiv \sup_{\alpha, \tau} J^r(x, \alpha, \tau), \quad (2)$$

where the maximization is over all admissible policies. An *optimal policy* (α^{r*}, τ^{r*}) is defined to satisfy $\phi^r(x) = J(x, \alpha^{r*}, \tau^{r*})$ for all x . In the undiscounted case, we drop the subscripts, so that $\phi \equiv \phi^0$ and $(\alpha^*, \tau^*) \equiv (\alpha^{0*}, \tau^{0*})$.

3. AN OPTIMAL POLICY

The optimal policy is characterized in Propositions 1, 2, and 3, depending on the cost of learning. Proposition 3 describes the case with the richest learning dynamics—the case that prevails when the cost of learning is small. The optimal policy is inferred from the value function, which solves an HJBQVI (Hamilton–Jacobi–Bellman (HJB) quasi-variational-inequality) equation. Because, in our case, the value function is nondifferentiable, the appropriate solution concept imposes restrictions both where the function is differentiable and where it is nondifferentiable, or has *kinks*.

3.1 The HJBQVI equation

HJBQVI is the continuous-time counterpart of the Bellman equation for discrete-time settings and, just like the Bellman equation, it relies on the dynamic programming principle.

LEMMA 1. *For any type profile $x \in [0, 1]^2$ and any finite stopping time τ' , the value function ϕ^r , defined in (2), satisfies the recursive relationship that encompasses the dynamic programming principle (DPP):⁹*

$$\begin{aligned} \phi^r(x) = \sup_{\alpha, \tau} \mathbb{E} \bigg[& \mathbf{1}_{\{\tau \geq \tau'\}} e^{-r\tau'} \phi^r(X^{\alpha, \tau}(\tau')) + \mathbf{1}_{\{\tau < \tau'\}} e^{-r\tau} \max_{i \in \mathcal{N}} \{X_i^{\alpha, \tau}(\tau)\} \\ & - c \int_0^{\tau' \wedge \tau} e^{-rs} ds \mid X^{\alpha, \tau}(0) = x \bigg]. \end{aligned}$$

PROOF. The lemma's conclusion follows from the DPP in Proposition 3.1 of Pham (1998). A handful of inconsequential differences between Pham's setup and ours are worth highlighting.

Pham's focus on the finite-horizon problem is not restrictive for us because HQ's maximal feasible surplus is 1; hence, the value function of the infinite-horizon problem can be shown to be the limit of a sequence of the value functions of finite-horizon problems.

⁹Operator $\wedge$ is the binary min operator.

Pham assumes that the intensity of the Poisson jump process is independent of the state. By contrast, in our problem, the jump intensity, aX , depends on the state process, X . Nevertheless, stochastic integration with respect to this more general jump process is well defined—which is all that matters for Pham's argument.¹⁰

The argument does not require positive discounting; $r = 0$ is admissible. $\square$

The DPP says that HQ's value today equals HQ's expected discounted continuation value at an arbitrary future stopping time τ' plus the expected discounted payoffs enjoyed until that time. These intervening flow payoffs and the eventual continuation value depend on the intervening controls, chosen to maximize HQ's value today.

Relying on the DD, one can characterize the value function ϕ^r as a solution to HJBQVI. Here, we informally derive HJBQVI, which disciplines ϕ^r at the points of differentiability, and a sufficient condition for ϕ^r not to contradict optimality at kinks (i.e., whenever the function is nondifferentiable). This sufficient condition is that all kinks be convex. A *convex kink* of ϕ^r at x admits a smooth function that passes through x and lies weakly below ϕ^r .¹¹

In discrete time, with a period length $\Delta > 0$, the value function ϕ^r is characterized by the Bellman equation¹²

$$\phi^r(x) = x_1 \vee x_2 \vee \max_q \left\{ -\Delta c + e^{-r\Delta} \left[\phi^r(x^{q\Delta}) + (1 - \phi^r(x^{q\Delta})) \left(1 - \prod_{i \in \mathcal{N}} (1 - x_i(1 - e^{-q_i\Delta})) \right) \right] \right\},$$

where $q \equiv (q_1, q_2) \in [0, 1]^2$, subject to $q_1 + q_2 = 1$, is the allocation of the learning effort for the duration of a period, and where $x^{q\Delta} \equiv (x_1^{q_1\Delta}, x_2^{q_2\Delta})$, with $x_i^{q_i\Delta}$ being the type revised down from x_i according to (1). When Δ is small, the display above requires

$$r\phi^r(x) \geq -c + \max_q \left\{ \frac{\phi^r(x^{q\Delta}) - \phi^r(x)}{\Delta} + \sum_{i \in \mathcal{N}} q_i x_i (1 - \phi^r(x)) \right\}, \quad (3)$$

where the inequality is understood to be approximate, in the sense that the terms of order Δ and smaller are omitted.

If ϕ^r is differentiable at x , taking the limit $\Delta \rightarrow 0$ in (3) while using $dx^{q_i\Delta}/d\Delta|_{\Delta=0} = -q_i x_i (1 - x_i)$ (implied by Bayes rule in (1)) yields the HJB equation¹³

$$0 \geq \max_q \left\{ -r\phi^r(x) - c + \sum_{i \in \mathcal{N}} q_i x_i (1 - (1 - x_i)\phi_i^r(x) - \phi^r(x)) \right\}. \quad (4)$$

¹⁰Process X is a finite-variation process. Hence, the stochastic integral with respect to X is well defined, as a path-by-path Riemann–Stieltjes integral (see, e.g., Protter 1990, Chapter I.6).

¹¹Formally, the solution concept for HJBQVI that characterizes the value function is the viscosity solution, which disciplines the kinks. (The theory of viscosity solutions that encompasses our setting is covered by Bardi and Capuzzo-Dolcetta 1997, and Oksendal and Sulem 2005.) The argument presented here amounts to showing that, for convex-kinked functions, the viscosity solution reduces to the satisfaction of HJBQVI only at the points of differentiability. In this paper, we work only with convex-kinked functions, so we do not need to invoke the full-fledged theory of viscosity solutions.

¹²Operator $\vee$ is the binary max operator. For notational parsimony, we abuse the notation and do not index ϕ^r by Δ .

¹³Subscripts denote partial derivatives: $\phi_i^r(x) \equiv \partial \phi^r(x) / \partial x_i$. Simple derivatives are denoted by primes.

Furthermore, because the maximand in (4) is linear in q , HJB in (4) is equivalent to

$$0 \geq \max_{i \in \mathcal{N}} \{-r\phi^r(x) - c - x_i(1 - x_i)\phi_i^r(x) + x_i(1 - \phi^r(x))\}. \quad (5)$$

If ϕ^r is not differentiable at x , then, before setting $\Delta \rightarrow 0$, we use a workaround so as not to lose any implication of optimality inherent in the Bellman equation, due to the nonexistence of the limit. Assume that, at x , ϕ^r has a convex kink. In this case, for the workaround, take an arbitrary smooth function ψ —called a test function—that satisfies $\psi \leq \phi^r$ and, at the kink x , satisfies $\psi(x) = \phi^r(x)$. The inequality in (3) is reinforced if one replaces the first two appearances of ϕ^r with ψ :

$$r\phi^r(x) \geq -c + \max_q \left\{ \frac{\psi(x^{q\Delta}) - \psi(x)}{\Delta} + \sum_{i \in \mathcal{N}} q_i x_i (1 - \phi^r(x)) \right\}.$$

Taking the limit $\Delta \rightarrow 0$ and noting that the resulting maximand is linear in q gives¹⁴

$$r\phi^r(x) \geq -c + \max_{i \in \mathcal{N}} \{-x_i(1 - x_i)\psi_i(x) + x_i(1 - \phi^r(x))\}. \quad (6)$$

The convexity of the kink at x implies that $\phi_{i-}^r(x) \leq \psi_i(x) \leq \phi_{i+}^r(x)$.¹⁵ As a result, inequality (6) at x is implied by (5) near x , where ϕ^r is differentiable.

To summarize the requirements for optimality, if every kink of a candidate value function is known to be convex (or if there are no kinks), then it suffices to verify that the function solves HJBQVI at the points of differentiability. No implication of optimality is lost. That is how the analysis proceeds in this paper. We guess a value function whose kinks are all convex and verify that, whenever differentiable, the guess solves the HJBQVI equation

$$0 = (x_1 - \phi^r(x)) \vee (x_2 - \phi^r(x)) \vee \max_{i \in \mathcal{N}} \{-r\phi^r(x) - c - x_i(1 - x_i)\phi_i^r(x) + x_i(1 - \phi^r(x))\} \quad (7)$$

on $\Omega \equiv (0, 1)^2$, subject to the boundary condition

$$\phi^r(x) = x_1 \vee x_2, \quad x \in \partial\Omega,$$

where $\partial\Omega$ is the boundary of Ω .

3.2 Maintained parameter restrictions and conventions

For tractability, we neglect discounting: $r = 0$. Section 5 remarks on the case of $r > 0$.

The analysis focuses on the economically nontrivial case in which learning is sufficiently cheap to be optimal for at least some type profiles:

$$c \in (0, \bar{c}), \quad \text{where } \bar{c} \equiv 0.25.$$

¹⁴When ϕ^r is differentiable at x , (6) reduces to (5).

¹⁵By convention, ϕ_{i-}^r and ϕ_{i+}^r denote, respectively, the left and right derivatives with respect to x_i .

The case with the richest learning dynamics is when the cost of learning is smaller still.¹⁶

$$c \in (0, \underline{c}), \quad \text{where } \ln \frac{(1 - \sqrt{\underline{c}})^2}{\underline{c}} = \frac{2}{1 - \sqrt{\underline{c}}} \implies \underline{c} \approx 0.047. \quad (8)$$

By the problem's symmetry, the optimal policy is also symmetric and, so, without loss of generality, the formal arguments focus on the hyperplane defined by $x_2 \geq x_1$.

3.3 Learning is prohibitively costly

Proposition 1 shows that a sufficiently high c makes learning prohibitively costly. Intuitively, learning at a sufficiently high cost must be suboptimal because the gain from allocating optimally—and, hence, from learning—is bounded.

PROPOSITION 1. *Suppose that learning is prohibitively costly, meaning that $c \geq \bar{c}$. Then the higher-type division wins immediately. The induced value function is $\phi(x) = x_1 \vee x_2 = x_2$.*

PROOF. To verify that ϕ is, indeed, the value function, first note that the kinks of ϕ , all at $x_1 = x_2$, are convex. Hence, it suffices to verify that ϕ satisfies HJBQVI at the points of differentiability. For this, substitute $\phi(x) = x_1 \vee x_2 = x_2$ into (7).

The quasi-variational inequality (QVI) $\phi(x) \geq x_1 \vee x_2$ holds by construction.

The HJB that corresponds to D_2 's learning is $-c \leq 0$, which obviously holds.

The HJB that corresponds to D_1 's learning is

$$-c + x_1(1 - x_2) \leq 0,$$

which holds for all x with $x_2 > x_1$ if and only if $c \geq \bar{c}$. Indeed, the inequality's left-hand side is maximized at $x_1 = x_2 = 1/2$ and attains the value $1/4 - c$, which is nonpositive if and only if $c \geq \bar{c}$. $\square$

3.4 Learning is moderately costly

Assume that learning is moderately costly, meaning that $c \in [\underline{c}, \bar{c})$. Then **Proposition 2** shows (and **Figure 2** illustrates) that if x is such that there is sufficient uncertainty about the efficient allocation, then the lower-type division learns; otherwise, the higher-type division wins immediately. Intuitively, asking the lower-type division to learn amounts to betting that this division will observe good news. This bet is insured by HQ's option to allocate to the higher-type division if no good news arrives.

Now, we illustrate in some detail the arguments used in the more complex case of cheap learning. Toward **Proposition 2**, guess that, for some threshold function b , $x_1 \leq b(x_2)$ implies that D_2 wins immediately, whereas $x_1 > b(x_2)$ implies that D_1 learns until

¹⁶To see that the solution of (8) is unique, note that the left-hand side of (8) is strictly decreasing in $\underline{c}$ and maps $(0, \bar{c})$ onto $\mathbb{R}_+$, whereas the right-hand side is strictly increasing in $\underline{c}$ and maps $(0, \bar{c})$ onto $(2, 4)$, a subset of $\mathbb{R}_+$.

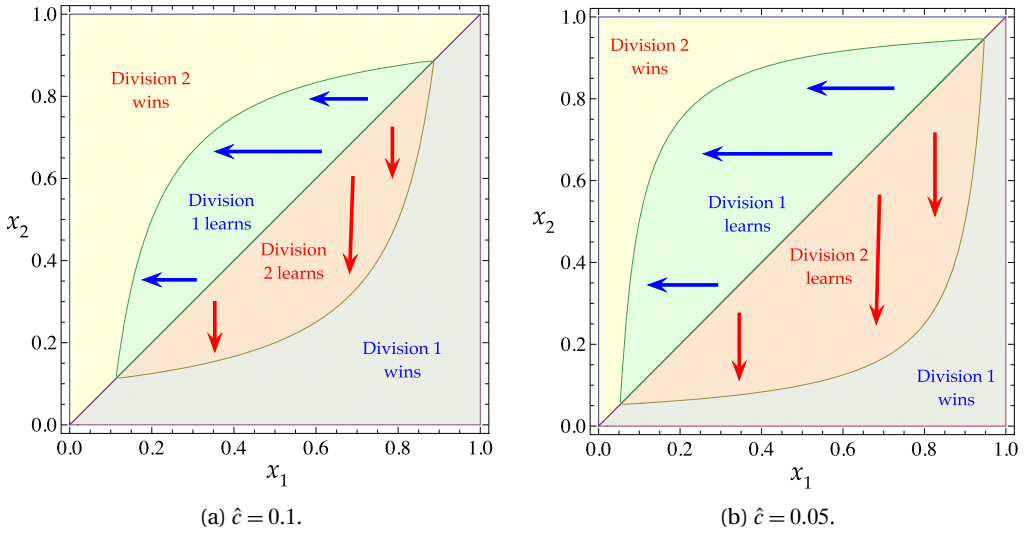


FIGURE 2. The optimal policy's prescription for each type profile when the learning cost is moderate; $c_1 \leq \hat{c} < c_2$. Within the lens-shaped region, one of the divisions learns. The arrows indicate the direction in which the type profile is revised if the division that learns observes no news. Outside the lens-shaped region, no division learns and the highest-type division wins.

either good news arrives or its revised type drops down to $b(x_2)$.¹⁷ The idea is that when $x_1 > b(x_2)$, x_1 and x_2 are close to each other, the uncertainty about the identity of the higher-value division is substantial, and, as a result, the return to learning is high.

Then $x_1 \leq b(x_2)$ implies that $\phi(x) = x_1 \vee x_2 = x_2$. When $x_1 > b(x_2)$, the value function ϕ —denoted by V on this set of type profiles—solves the HJB equation

$$0 = -c - x_1(1 - x_1)V_1(x) + x_1(1 - V(x)), \quad (9)$$

which picks out the component of HJBQVI that corresponds to D_1 's learning. The HJB equation (9) is solved subject to the boundary condition $V(b(x_2), x_2) = x_2$. The boundary condition captures the assumption that, once D_1 's type has dropped down to $b(x_2)$ (because no good news has arrived), no further learning occurs and the higher-type division wins. The solution is

$$V(x) = 1 - c + (1 - x_1) \left(c \ln \frac{b(x_2)(1 - x_1)}{(1 - b(x_2))x_1} - \frac{1 - x_2 - c}{1 - b(x_2)} \right). \quad (10)$$

To determine $b(x_2)$ in (10), we solve the *V-auxiliary problem*: choose $b(x_2)$ in $[0, x_2]$ to maximize $V(x)$ in (10). Then, if interior on the interval $[0, x_2]$, an optimal $b(x_2)$ satisfies the first-order condition

$$b(x_2) = \frac{c}{1 - x_2}. \quad (11)$$

¹⁷Recall that we assume that $x_2 \geq x_1$.

Threshold b in (11) is, indeed, interior if $x_2 \in (\underline{x}, \bar{x})$, where

$$\underline{x} \equiv \frac{1 - \sqrt{1 - 4c}}{2} \quad \text{and} \quad \bar{x} \equiv \frac{1 + \sqrt{1 - 4c}}{2}.$$

By $c < \bar{c}$, interval $(\underline{x}, \bar{x})$ is nonempty.

We now show that the kinks of the constructed value function are convex. All kinks are on the diagonal (i.e., the 45-degree line that passes through the origin). Off the diagonal, kinks could potentially occur only at type profiles $x = (b(x_2), x_2)$ (indexed by x_2), where $x_1 \vee x_2$ meets $V(x)$. Differentiation ascertains, however, that segments $x_1 \vee x_2$ and $V(x)$ paste together smoothly:

$$\begin{aligned} \phi_{1-}(x) = \phi_{1+}(x) &\iff \frac{\partial(x_1 \vee x_2)}{\partial x_1} = V_1(x), \\ \phi_{2-}(x) = \phi_{2+}(x) &\iff V_2(x) = \frac{\partial(x_1 \vee x_2)}{\partial x_2}. \end{aligned}$$

The smooth pasting is a corollary to the optimality of b —an envelope-theorem result (Milgrom and Segal 2002, Corollary 6).

By contrast, on the diagonal, each point is a kink. Among these, each kink x with $x_1 = x_2 \in [0, \underline{x}) \cup (\bar{x}, 1]$ prescribes immediate allocation. Since, in its neighborhood, the value function is $x_1 \vee x_2$, all these kinks are convex, as in the case of prohibitively costly learning.

A kink x with $x_1 = x_2 \in (\underline{x}, \bar{x})$ need not be convex in general, but is convex when

$$\phi_{1-}(x) \leq \phi_{2+}(x) \iff c \geq \underline{c},$$

as we now proceed to show. To see that $\phi_{1-}(x) \leq \phi_{2+}(x)$ captures convexity, note that, graphically, the convexity of a kink at x requires that, as one passes through x in the direction of any vector $v = (v_1, v_2)$ that traverses the diagonal from above (i.e., has a slope between $-3\pi/4$ and $\pi/4$ radians), the corresponding directional derivative of the value function experiences a jump upward (if at all):

$$v_1 \phi_{1-}(x) + v_2 \phi_{2+}(x) \leq v_1 \phi_{1+}(x) + v_2 \phi_{2-}(x).$$

Further, by the symmetry of ϕ with respect to the diagonal, $\phi_{1+}(x) = \phi_{2+}(x)$ and $\phi_{2-}(x) = \phi_{1-}(x)$. As a result, the inequality in the display above becomes

$$(v_1 - v_2)(\phi_{1-}(x) - \phi_{2+}(x)) \leq 0 \iff \phi_{1-}(x) \leq \phi_{2+}(x),$$

where the equivalence follows from $v_1 > v_2$, which is dictated by the orientation of v . The equivalence between $\phi_{1-}(x) \leq \phi_{2+}(x)$ and $c \geq \underline{c}$ follows from straightforward, if tedious, algebraic manipulations (detailed in Lemma 2).

LEMMA 2. *The following statements are equivalent:*

- (i) *For all x with $x_1 = x_2 \in [\underline{x}, \bar{x}]$, $\phi_{1-}(x) \leq \phi_{2+}(x)$.*
- (ii) *We have $c \geq \underline{c}$.*

PROOF. First, note that, at x with $x_1 = x_2 \in [\underline{x}, \bar{x}]$, $\phi_{1-}(x) = V_1(x)$ and $\phi_{2+}(x) = V_2(x)$. Define $\mathcal{D}(z) \equiv V_2(z, z) - V_1(z, z)$. We must show that $\mathcal{D}(z) \geq 0$ for all $z \in [\underline{x}, \bar{x}]$ if and only if $c \geq \underline{c}$. Substituting the functional forms gives

$$\mathcal{D}(z) = c \left(\frac{1}{z} + \frac{1-z}{1-c-z} + \ln \frac{c(1-z)}{z(1-c-z)} \right).$$

Then $\mathcal{D}(\underline{x}) = \mathcal{D}(\bar{x}) = 1$. Any critical point of $\mathcal{D}$ in $(\underline{x}, \bar{x})$ is characterized by the first-order condition $d\mathcal{D}(z)/dz = 0$, whose solutions are $z^* = 1 - \sqrt{c}$ and $z^{**} = \sqrt{1-c}$. Of these, only z^* is in $(\underline{x}, \bar{x})$. Thus, $\mathcal{D}$ is nonnegative on $[\underline{x}, \bar{x}]$ if and only if it is nonnegative at z^* .

From

$$\mathcal{D}(z^*) = c \left(\frac{2}{1-\sqrt{c}} - \ln \frac{(1-\sqrt{c})^2}{c} \right),$$

conclude that

$$\mathcal{D}(z^*) \geq 0 \iff \frac{2}{1-\sqrt{c}} \geq \ln \frac{(1-\sqrt{c})^2}{c} \iff c \geq \underline{c}. \quad \square$$

To recap, we have conjectured the optimal policy and the associated value function, which has been verified to be convex-kinked. To validate the conjecture, it remains to verify that ϕ satisfies HJBQVI. The verification is split into two cases: $x_1 \leq b(x_2)$ and $x_1 > b(x_2)$.

- When $x_1 \leq b(x_2)$, $\phi(x) = x_1 \vee x_2 = x_2$, which, by construction, satisfies the QVI $\phi(x) \geq x_1 \vee x_2$.

The HJB that corresponds to D_1 's learning is $-c + x_1(1 - x_2) \leq 0$, which is implied by $x_1 \leq b(x_2)$.

The HJB that corresponds to D_2 's learning is $-c \leq 0$, which obviously holds.

- When $x_1 > b(x_2)$, $\phi = V$, which, by construction, satisfies the HJB that corresponds to D_1 's learning.

QVI $V(x) \geq x_2$ holds by the optimality of b .

QVI $V(x) \geq x_1$ is implied by $V(x) \geq x_2$ and (by convention) $x_2 \geq x_1$.

The HJB that corresponds to D_2 's learning requires that

$$-c - x_2(1 - x_2)V_2(x) + x_2(1 - V(x)) \leq 0.$$

By the envelope theorem applied to (10), $V_2(x) = (1 - x_1)/(1 - b(x_2))$. Substituting $V_2(x)$ and $V(x)$ into the display above and dividing by $-c(1 - x_2)$ gives the equivalent inequality

$$\Phi^A(x, c) \equiv \frac{1 - c - x_1 x_2}{1 - c - x_2} + x_2 \frac{1 - x_1}{1 - x_2} \ln \frac{c(1 - x_1)}{x_1(1 - c - x_2)} \geq 0. \quad (12)$$

By part (iii) of Lemma A.1 in the Appendix, (12) holds if and only if $c \geq \underline{c}$.

Because ϕ is convex-kinked and satisfies HJBQVI, Proposition 2 follows.

PROPOSITION 2. *Suppose that learning is moderately costly, meaning that $c \in [\underline{c}, \bar{c})$. Then, if $x_1 > b(x_2)$, the lower-type division, D_1 , learns and, if it observes good news, wins. If $x_1 \leq b(x_2)$, D_2 wins immediately.*

The economic content of condition $c \geq \underline{c}$ in [Proposition 2](#) is the suboptimality of learning by D_2 , the higher-type division. What $\Phi^A(x, c) \geq 0$ in (12) expresses and $c \geq \underline{c}$ guarantees is that a one-off (literally, infinitesimal) deviation toward learning by D_2 is unprofitable.

3.5 Learning is cheap

Assume that $c < \underline{c}$. Then [Proposition 3](#) shows that it may also be optimal for D_2 , the higher-type division, to learn. Intuitively, it is suboptimal to ask D_2 to learn when the intended period of learning is insufficiently long to flip the ranking of the divisions' revised types; such learning would not affect the allocation decision, but would entail a wasteful learning cost. Lengthy learning is only ever justified, however, if it is sufficiently cheap—which, here, means that $c < \underline{c}$ —and if the gains from learning are sufficiently large. These gains are large when x_1 and x_2 are close to each other (so that the identity of the more valuable project is highly uncertain), when both x_1 and x_2 are far away from 0 and 1 (so that each project's value is highly uncertain), and when x_1 and x_2 are rather large (so that learning is rather informative; the good news arrives with a high probability, and if it does not arrive, then the type is revised downward fast).

[Figure 3](#) illustrates the optimal policy. For the type profiles in the heart-shaped region, the higher-type division learns. Elsewhere in the lens-shaped region, the lower-type division learns. On the diagonal that traverses the heart-shaped region, both divisions learn simultaneously.¹⁸ If neither division learns, the higher-type division wins immediately. The boundary of the lens-shaped region is demarcated by b defined in (11). The boundary of the heart-shaped region is derived in the remainder of this section.

The (rather technical) intuition for the heart-shaped region in [Figure 3](#) can be gleaned from studying the set of type profiles on which $c < \underline{c}$ causes the verification of the conjectured value function in [Proposition 2](#) to fail by causing some kinks to be nonconvex. That is, we are interested in the *failure set*

$$\mathcal{F} \equiv \{x \mid \Phi^A(x, c) < 0\},$$

on which the HJB component that corresponds to D_2 's learning fails in [Proposition 2](#) when $c < \underline{c}$. [Figure 4\(a\)](#) illustrates $\mathcal{F}$ (and, by the problem's symmetry, its reflection about the diagonal), which is heart-shaped. On $\mathcal{F}$, infinitesimal learning by D_2 followed by D_1 's learning is a profitable deviation from the policy in which only D_1 learns, as in [Proposition 2](#). [Proposition 3](#) “patches” the failure set $\mathcal{F}$ by making the higher-type division learn on a heart-shaped region that covers $\mathcal{F}$. The region on which the higher-type division learns according to [Proposition 3](#) exceeds $\mathcal{F}$ (see [Figure 4\(b\)](#)) because one can chain together infinitesimal deviations to obtain a deviation that is profitable even at a type profile at which a single infinitesimal deviation is unprofitable.

¹⁸This simultaneous learning is not nongeneric.

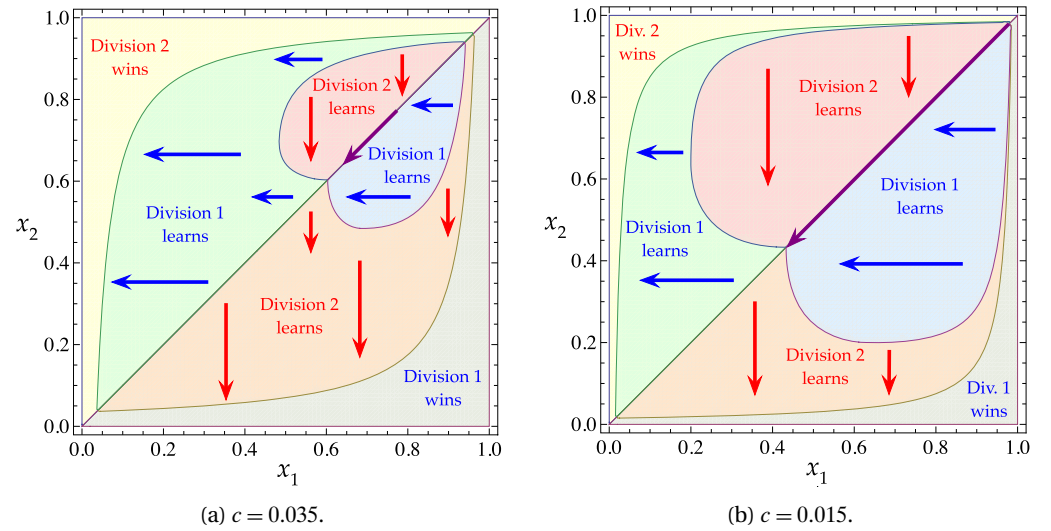
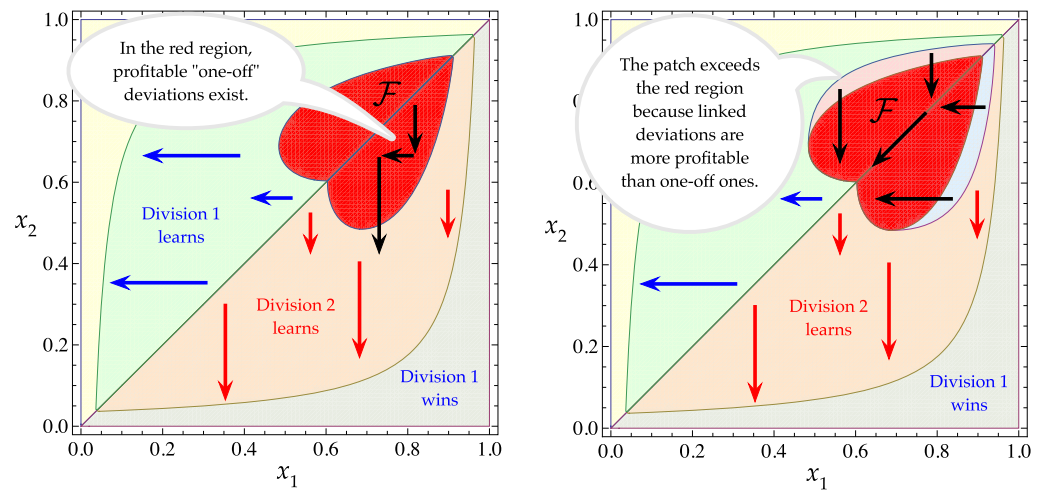


FIGURE 3. The optimal policy's prescription for each type profile when learning is cheap: $c < \underline{c}$. Within the lens-shaped region (which encompasses the heart-shaped region), at least one of the divisions learns. Each arrow indicates the direction in which the type profile is revised if the division that learns observes no news. Outside the lens-shaped region, no division learns and the higher-type division wins immediately.



(a) On $\mathcal{F}$, instead of asking the lower-type division to learn (as Proposition 2 would have it), HQ can achieve a higher payoff by momentarily asking the higher-type division to learn and then reverting to asking the lower-type division to learn.
(b) The heart-shaped region on which Proposition 3 prescribes that the higher-type division learns exceeds $\mathcal{F}$.

FIGURE 4. The policy prescribed by Proposition 2 is no longer optimal when $c < \underline{c}$.

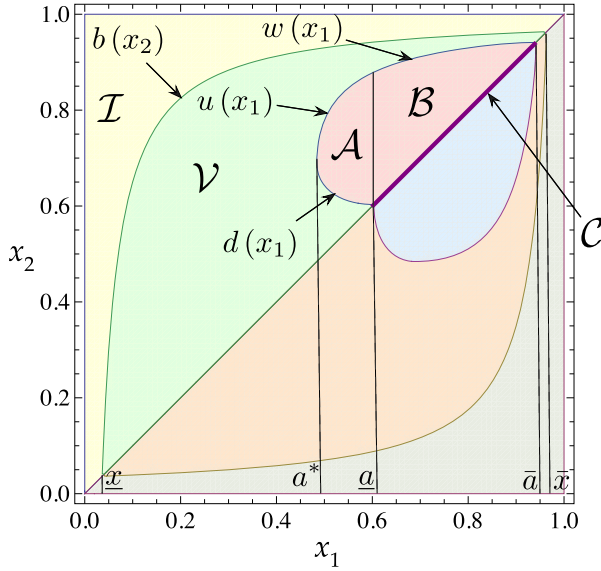


FIGURE 5. Each of the sets $\mathcal{I}$, $\mathcal{V}$, $\mathcal{A}$, $\mathcal{B}$, and $\mathcal{C}$ collects the type profiles at which the optimal policy makes identical prescriptions. Set $\mathcal{A}$ is bounded by curve u above and by curve d below. Set $\mathcal{B}$ is bounded by curve w above and by the diagonal below. Set $\mathcal{V}$ is bounded by curve b above and to the left. D_2 learns on $\mathcal{A}$ and $\mathcal{B}$. Both divisions learn on $\mathcal{C}$. D_1 learns on $\mathcal{V}$. D_2 wins immediately on $\mathcal{I}$.

The optimal policy is formally described in terms of five type-profile sets, or regions: $\mathcal{A}$, $\mathcal{B}$, $\mathcal{C}$, $\mathcal{I}$, and $\mathcal{V}$. These regions are depicted in Figure 5. Region $\mathcal{I}$ is the region on which D_2 wins immediately in Proposition 2. Region $\mathcal{V}$ is the region on which D_1 leans in Proposition 2 less $\mathcal{A}$, $\mathcal{B}$, and $\mathcal{C}$. To characterize $\mathcal{A}$, $\mathcal{B}$, and $\mathcal{C}$, we consider three auxiliary stopping problems: $\mathcal{A}$ auxiliary, $\mathcal{B}$ auxiliary, and $\mathcal{C}$ auxiliary.

The $\mathcal{C}$ -auxiliary problem is defined on the subset

$$\hat{\mathcal{C}} \equiv \{(x_1, x_2) \in [0, 1]^2 \mid \underline{x} \leq x_1 = x_2 \leq \bar{x}\}$$

of the diagonal. In the $\mathcal{C}$ -auxiliary problem, both divisions learn until either one observes the good news or until both revised types drop down to some optimally chosen threshold, denoted by $\underline{a}$ —whichever happens first. At $\underline{a}$, the strategy described in Proposition 2 is followed: D_1 learns. Let C denote the value function of the $\mathcal{C}$ -auxiliary problem.

While both divisions learn, along the diagonal (z, z) indexed by z , the value function satisfies the HJB equation

$$0 = -c - z(1 - z) \frac{C'(z)}{2} + z(1 - C(z)), \quad z \in [\underline{x}, \bar{x}],$$

subject to $C(\underline{a}) = V(\underline{a}, \underline{a})$. To find the optimal $\underline{a}$, first solve the differential equation in the display above for an arbitrary a : $C(a) = V(a, a)$. The solution is

$$C(z) = 1 + (1 - z)^2 \underbrace{\left(2c\sigma(a) - \frac{1 - V(a, a)}{(1 - a)^2} - 2c\sigma(z) \right)}_{\equiv M^C(a)}, \quad (13)$$

where V is defined in (10) and

$$\sigma(s) \equiv \frac{2s}{1-s} + \frac{1}{2} \left(\frac{s}{1-s} \right)^2 + \ln \frac{s}{1-s}, \quad s \in (0, 1).$$

The threshold $\underline{a}$ maximizes C over a and satisfies the first-order condition $dM^C(\underline{a})/da = 0$, where M^C is defined in (13). Equivalently, by

$$\frac{dM^C(a)}{da} = \frac{c\Phi^C(a, c)}{a(1-a)^2}, \quad \text{where } \Phi^C(a, c) \equiv \Phi^A(a, a, c), \quad (14)$$

the threshold $\underline{a}$ satisfies $\Phi^C(\underline{a}, c) = 0$.

Lemma 3 characterizes the solution of the C -auxiliary problem in terms of the thresholds

$$\underline{a} \equiv \min\{a \in [\underline{x}, \bar{x}] \mid \Phi^C(a, c) = 0\}, \quad (15)$$

$$\bar{a} \equiv \max\{a \in [\underline{a}, \bar{x}] \mid M^C(a) = M^C(\underline{a})\}, \quad (16)$$

and the type subset

$$\mathcal{C} \equiv \{(z, z) \in \hat{\mathcal{C}} \mid z \in (\underline{a}, \bar{a})\}.$$

LEMMA 3. On $\mathcal{C}$, D_1 and D_2 both learn; C is given by (13) with $a = \underline{a}$. On $\hat{\mathcal{C}} \setminus \mathcal{C}$, D_1 learns as in Proposition 2; C coincides with V .

PROOF. To solve $\max_{a \in [\underline{\theta}, z]} M^C(a)$, let us examine the shape of M^C on $[\underline{x}, \bar{x}]$; Figure 6 previews M^C .

By (14), the sign of $dM^C(a)/da$ coincides with the sign of $\Phi^C(a, c)$. By Lemma A.1 in Appendix, $\Phi^C(\cdot, c)$ is positive at first, then intersects zero at a point, then is negative, then intersects zero at a point, and then is positive again. As a result, M^C is wave-shaped, with local maxima at $\underline{a}$, where $\underline{a}$ is the smallest of the two roots of $\Phi^C(\cdot, c) = 0$, and at $\bar{x}$. Furthermore, by Lemma A.2 in the Appendix, $\bar{x}$ is the unique global maximum: $M^C(\bar{x}) > M^C(\underline{a})$.

Lemma A.2 and the wave shape of M^C imply the existence of a unique $\bar{a} \in (\underline{a}, \bar{x})$ such that $M^C(\bar{a}) = M^C(\underline{a})$, as defined in (16).

The described properties of M^C have the following implications for the C -auxiliary problem. When $z \notin (\underline{a}, \bar{a})$, $\arg \max_{a \in [\underline{x}, z]} M^C(a) = \{z\}$ and, so, $C(z) = V(z, z)$; D_1 learns. When $z \in (\underline{a}, \bar{a})$, $\arg \max_{a \in [\underline{x}, z]} M^C(a) = \{\underline{a}\}$ and, so, $C(z) > V(z, z)$; D_1 and D_2 learn simultaneously until either division observes the good news or until both divisions' types fall to $\underline{a}$, whereupon (say) D_1 learns. $\square$

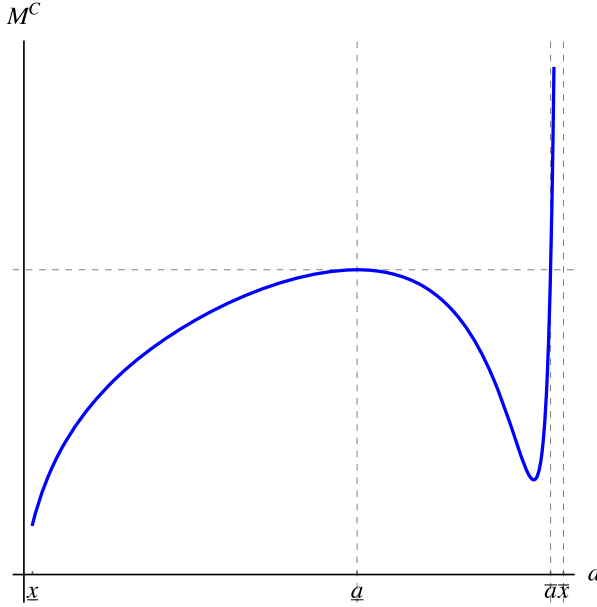


FIGURE 6. The maximand, M^C , in the C -auxiliary problem. The sign of dM^C/da coincides with the sign of $\Phi^C(\cdot, c)$.

To characterize set $\mathcal{A}$, we formulate the A -auxiliary problem on the set

$$\hat{\mathcal{A}} \equiv \{(x_1, x_2) \in (a^*, \underline{a}] \times [0, 1] \mid x_2 \geq x_1\}.$$

In this problem, either D_1 learns as in Proposition 2, or D_2 learns until either it observes good news or its revised type reaches some optimally chosen threshold, denoted by $d(x_1)$. At that threshold, the strategy described in Proposition 2 is followed: D_1 learns. Let A denote the value function of the A -auxiliary problem.

While D_2 learns, the associated value function satisfies the HJB equation

$$0 = -c - x_2(1 - x_2)A_2(x) + x_2(1 - A(x)), \quad x \in \hat{\mathcal{A}}, \quad (17)$$

subject to $A(x_1, d(x_1)) = V(x_1, d(x_1))$. To find the optimal $d(x_1)$, first let us solve the differential equation in the display above for an arbitrary a : $A(x_1, a) = V(x_1, a)$. The solution is

$$A(x_1, x_2) \equiv 1 + (1 - x_2) \underbrace{\left(c\eta(a) - \frac{1 - V(x_1, a)}{1 - a} \right)}_{\equiv M^A(x_1, a)} - c\eta(x_2), \quad (18)$$

where

$$\eta(s) \equiv \frac{s}{1 - s} + \ln \frac{s}{1 - s}, \quad s \in (0, 1).$$

For each x_1 , $d(x_1)$ maximizes A over a and satisfies the first-order condition $M_2^A(x_1, d(x_1)) = 0$, where M^A is defined in (18). Equivalently, by

$$M_2^A(x_1, x_2) = \frac{c\Phi^A(x, c)}{x_2(1 - x_2)},$$

$d(x_1)$ satisfies $\Phi^A(x_1, d(x_1), c) = 0$.

Lemma 4 characterizes the solution to the A -auxiliary problem in terms of the type subset

$$\mathcal{A} \equiv \{x \in (a^*, \underline{a}] \times [0, 1] \mid x_2 \in (d(x_1), u(x_1))\},$$

where

$$\begin{aligned} a^* &\equiv \min\{x_1 \in [0, 1] \mid \exists x_2 \in [x_1, 1] \text{ s.t. } \Phi^A(x_1, x_2, c) = 0\} \\ d(x_1) &\equiv \min\{a \in [x_1, 1] \mid \Phi^A(x_1, a, c) = 0\}, \quad x_1 \in (a^*, \bar{a}), \end{aligned} \quad (19)$$

and, letting b^{-1} denote the inverse of b in (11),

$$u(x_1) \equiv \max\{a \in [x_1, b^{-1}(x_1)] \mid M^A(x_1, a) = M^A(x_1, d(x_1))\}, \quad x_1 \in (a^*, \bar{a}). \quad (20)$$

LEMMA 4. *On $\mathcal{A}$, D_2 learns; A is given by (18) with $a = d(x_1)$. On $\hat{\mathcal{A}} \setminus \mathcal{A}$, D_1 learns as in Proposition 2; A coincides with V . Moreover, on $\mathcal{A}$, $A \geq V$, $\mathcal{F} \cap \hat{\mathcal{A}} \subset \mathcal{A}$, and, for $x_1 \in (a^*, \underline{a})$, $u(x_1) < b^{-1}(x_1)$.*

PROOF. To solve $\max_{a \in [x_1, \bar{x}]} M^A(x_1, a)$, let us examine the shape of M^A .

By (12), the sign of $M_2^A(x)$ coincides with the sign of $\Phi^A(x, c)$. By part (iv) of Lemma A.1 in the Appendix, $\Phi^A(x_1, \cdot, c)$ is quasi-convex; thus, $M^A(x_1, \cdot)$ is wave-shaped, as depicted in Figure 7.

Because $\Phi^A(x_1, \cdot, c)$ is quasi-convex, it intersects zero at most twice, in which case the first intersection is a local maximum of M^A . We denote this local maximum by $d(x_1)$, defined in (19). This local maximum is not global; Lemma A.3 in the Appendix implies that $M^A(x_1, d(x_1)) < M^A(x_1, b^{-1}(x_1))$.

The wave shape of $M^A(x_1, \cdot)$ implies the existence of a unique $u(x_1) \in (d(x_1), b^{-1}(x_1))$ such that $M^A(x_1, u(x_1)) = M^A(x_1, d(x_1))$, as in (20).

The derived properties of M^A have the following implications for the maximization problem $\max_{a \in [x_1, x_2]} M^A(x_1, a)$. When $x_2 \notin (d(x_1), u(x_1))$, $M^A(x_1, \cdot)$ is maximized on $[x_1, x_2]$ at x_2 and $A(x) = V(x)$. When $x_2 \in (d(x_1), u(x_1))$, $M^A(x_1, \cdot)$ is uniquely maximized on $[x_1, x_2]$ at $d(x_1)$ and $A(x) > V(x)$. Thus, on $\mathcal{A}$, $A \geq V$, as claimed in the “moreover” part of the lemma.

To complete the “moreover” part of the lemma, note that, because the slope of $M^A(x_1, \cdot)$ coincides with the sign of $\Phi^A(x_1, \cdot, c)$, the shape of $M^A(x_1, \cdot)$, summarized in Figure 7, implies that $\mathcal{F} \cap \hat{\mathcal{A}} \subset \mathcal{A}$.

Finally, by Lemma A.3, $\arg \max_{a \in [x_1, b^{-1}(x_1)]} M^A(x_1, a) = \{b^{-1}(x_1)\}$, which, together with the wave shape of $M^A(x_1, \cdot)$, implies that $u(x_1) < b^{-1}(x_1)$. $\square$

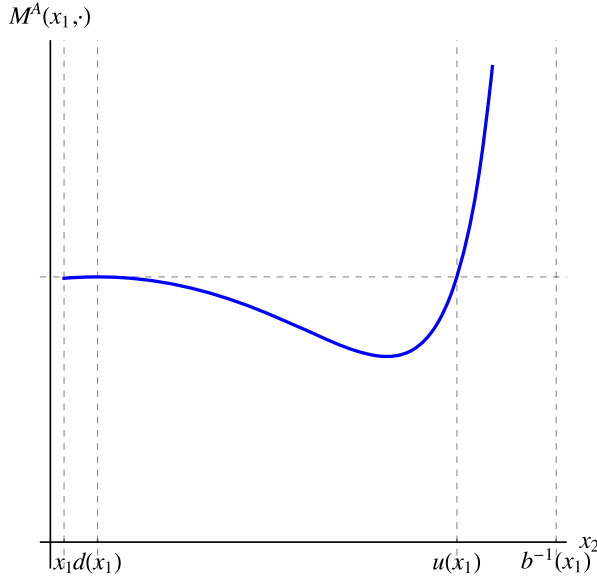


FIGURE 7. The maximand, M^A , in the A -auxiliary problem. The sign of $M_2^A(x_1, \cdot)$ coincides with the sign of $\Phi^A(x_1, \cdot, c)$.

To characterize set $\mathcal{B}$, we formulate the B -auxiliary problem, defined on the set

$$\hat{\mathcal{B}} \equiv \{(x_1, x_2) \in (\underline{a}, \bar{a}) \times [0, b^{-1}(x_1)] \mid x_2 \geq x_1\}.$$

In the B -auxiliary problem, either D_1 immediately learns, as in Proposition 2, or D_2 learns until it observes good news or until its revised type reaches x_1 , whereupon both divisions learn as prescribed by the C -auxiliary problem. Let B denote the value function of the B -auxiliary problem.

When only D_2 learns, the associated value function satisfies the HJB equation

$$0 = -c - x_2(1 - x_2)B_2(x) + x_2(1 - B(x)) \quad (21)$$

subject to $B(z, z) = C(z)$ for each $z \in [\underline{a}, \bar{a}]$. The solution is

$$B(x) \equiv 1 - (1 - x_2) \left(\frac{1 - C(x_1)}{1 - x_1} + c\eta(x_2) - c\eta(x_1) \right). \quad (22)$$

Because the threshold at which D_2 stops learning is assumed to be x_1 , it remains only to identify the threshold, denoted by $w(x_1)$, such that D_2 does not learn if $x_2 \geq w(x_1)$. This threshold is

$$w(x_1) \equiv \min\{a \in [x_1, b^{-1}(x_1)] \mid V(x_1, a) = B(x_1, a)\}, \quad x_1 \in (\underline{a}, \bar{a}). \quad (23)$$

Lemma 5 summarizes the solution to the B -auxiliary problem in terms of the type subset

$$\mathcal{B} \equiv \{(x_1, x_2) \in (\underline{a}, \bar{a}) \times [0, 1] \mid x_2 \in (x_1, w(x_1))\},$$

where $\underline{a}$, $\bar{a}$, and w are defined in (15), (16), and (23), respectively.

LEMMA 5. On $\mathcal{B}$, D_2 learns, B is given by (22) and satisfies $B > V$, and $w(x_1) < b^{-1}(x_1)$. On $\hat{\mathcal{B}} \setminus \mathcal{B}$, D_1 learns as in Proposition 2; B coincides with V .

PROOF. The proof is in the text above, except for the claim that $w(x_1) < b^{-1}(x_1)$, which is Lemma A.4 in the Appendix. $\square$

We can now assemble the pieces to form the conjectured value function when learning is cheap:

$$\phi(x) = \mathbf{1}_{\{x \in \mathcal{A}\}} A(x) + \mathbf{1}_{\{x \in \mathcal{B}\}} B(x) + \mathbf{1}_{\{x \in \mathcal{C}\}} C(x) + \mathbf{1}_{\{x \in \mathcal{V}\}} V(x) + \mathbf{1}_{\{x \in \mathcal{I}\}} x_2. \quad (24)$$

Function ϕ is extended to the hyperplane $x_2 < x_1$ by the symmetry about the 45-degree line.¹⁹ Function ϕ has kinks where sets $\mathcal{A}$ and $\mathcal{V}$ meet at curve u , where sets $\mathcal{B}$ and $\mathcal{V}$ meet at curve w , and on the 45-degree line outside the heart-shaped region in Figure 5. The kinks at these boundaries comply with the rule of thumb of Peskir and Shiryaev (2006, Chapter IV.9): whenever the type process is certain to move away from a boundary, the value function at the boundary is liable to be nondifferentiable. Alternatively, here, all the kinks of ϕ are at the boundaries that have not been explicitly determined as solutions to optimal stopping; thus, at those boundaries, the envelope theorem does not guarantee smooth pasting.

To verify the conjecture in (24), we must perform the same two-step procedure that leads to Proposition 2: verify that ϕ 's kinks are convex and verify that ϕ solves HJBQVI. The requisite steps are contained in the proof of Proposition 3.

PROPOSITION 3. Suppose that learning is cheap, meaning that $c < \underline{c}$. Then, on $\mathcal{A}$ and $\mathcal{B}$, the higher-type division, D_2 , learns and, if it observes good news, wins. On $\mathcal{V}$, D_1 learns and, if it observes good news, wins. On $\mathcal{I}$, D_2 wins immediately.

PROOF. The the proof proceeds in steps collected into three groups. Step 1 is concerned with the convexity of kinks. Step 2 is concerned with the satisfaction of the QVIs. Step 3 is concerned with the satisfaction of the HJB equations.

Step 1.1. By the argument leading up to Proposition 2, there are no kinks at the boundary of $\mathcal{I}$ and $\mathcal{V}$.

Step 1.2. We show that there are no kinks where $\mathcal{A}$ and $\mathcal{V}$ meet along d . Indeed,

$$\begin{aligned} A_2(x_1, d(x_1)) &= 1 - x_1 - \frac{c(1 - x_1)}{1 - d(x_1)} \ln \frac{c(1 - x_1)}{x_1(1 - c - d(x_1))} - \frac{c}{d(x_1)} \\ &= \frac{(1 - d(x_1))(1 - x_1)}{1 - c - d(x_1)} = V_2(x_1, d(x_1)), \end{aligned}$$

where the first equality follows by differentiating A , the second equality uses $\Phi^A(x_1, d(x_1), c) = 0$ (from the definition of d) to substitute out the logarithmic term, and the third equality follows by differentiating V . Furthermore, direct differentiation (without

¹⁹The value function in (24) also describes the intermediate-cost case, in which $\mathcal{A}$, $\mathcal{B}$, and $\mathcal{C}$ are empty, and the prohibitive-cost case, in which $\mathcal{V}$ is empty as well.

using the condition for d) establishes that $A_1(x) = V_1(x)$, including at $x = (x_1, d(x_1))$. Thus, there are no kinks where $\mathcal{A}$ and $\mathcal{V}$ meet along d .

Step 1.3. We show that there are no kinks where $\mathcal{A}$ and $\mathcal{B}$ meet, a vertical boundary. We shall ascertain that $A_1(x) = B_1(x)$ along that boundary. Indeed,

$$\begin{aligned} A_1(x) &= \frac{1-x_2}{1-d(x_1)} V_1(x_1, d(x_1)) \\ &=_{x_1 \rightarrow \underline{a}} \frac{1-x_2}{1-\underline{a}} V_1(\underline{a}, \underline{a}) \\ &= \frac{1-x_2}{1-\underline{a}} \left(1-\underline{a} - \frac{c}{\underline{a}} - c \ln \frac{c(1-\underline{a})}{\underline{a}(1-c-\underline{a})} \right), \end{aligned}$$

where the first equality follows by the envelope theorem ($d(x_1)$ has been chosen optimally), the second equality uses $\lim_{x_1 \rightarrow \underline{a}} d(x_1) = \underline{a}$, and the last equality is by differentiation and rearranging.

Further,

$$\begin{aligned} B_1(x) &= (1-x_2) \left(\frac{C'(x_1)}{1-x_1} - \frac{1-C(x_1)}{(1-x_1)^2} + c\eta'(x_1) \right) \\ &= \frac{(1-x_2)(1-C(x_1)-c/x_1)}{(1-x_1)^2} \\ &=_{x_1 \rightarrow \underline{a}} \frac{1-x_2}{1-\underline{a}} \left(1-\underline{a} - \frac{c}{\underline{a}} - c \ln \frac{c(1-\underline{a})}{\underline{a}(1-c-\underline{a})} \right), \end{aligned}$$

where the first equality is by differentiation, the second equality is by $\eta'(x_1) = 1/(x_1(1-x_1)^2)$ and by

$$C'(z) = \frac{2(1-C(z)-c/z)}{1-z},$$

and the final equality follows by substituting C . As a result, $A_1(\underline{a}, x_2) = B_1(\underline{a}, x_2)$ for $x_2 \in (\underline{a}, b^{-1}(\underline{a}))$.

Step 1.4. We show that there are no kinks where $\mathcal{B}$ meets its reflection about the diagonal, on $\mathcal{C}$. Indeed,

$$B_1(z, z) = \frac{1-C(z)-c/z}{1-z} = B_2(z, z),$$

where the first equality follows by the computations in Step 1.2 and the second equality follows by differentiation.

Step 1.5. We show that the kinks where $\mathcal{V}$ meets its reflection about the 45-degree line, on $\hat{\mathcal{C}} \setminus \mathcal{C}$, are convex. By [Lemma 2](#), we show that $V_1(z, z) \leq V_2(z, z)$ for $(z, z) \in \hat{\mathcal{C}} \setminus \mathcal{C}$:

$$\begin{aligned} V_1(z, z) &= 1-z - \frac{c}{z} - c \ln \frac{c(1-z)}{z(1-c-z)}, \\ V_2(z, z) &= -\frac{(1-z)^2}{1-c-z}. \end{aligned}$$

Rearranging implies that $V_1(z, z) \leq V_2(z, z)$ is equivalent to $\Phi^C(z, c) \geq 0$, which is implied by Lemma A.1.

Step 1.6. We argue that the kinks where $\mathcal{V}$ meets $\mathcal{A}$ along u are convex. Indeed, the surfaces constructed in the $\mathcal{A}$ -auxiliary and the V -auxiliary problems, characterized by $\mathcal{A}$ and V , are both smooth. As x_2 increases, the $\mathcal{A}$ -induced surface cuts into the V -induced surface from above, by construction, and because $u(x_1) < b^{-1}(x_1)$ (Lemma A.3). Thus, the kinks along u are convex.

Step 1.7. We argued that the kinks where $\mathcal{V}$ meets $\mathcal{B}$ along w are convex. Indeed, the surfaces constructed in the $\mathcal{B}$ -auxiliary and the V -auxiliary problems, characterized by $\mathcal{B}$ and V , are both smooth. As x_2 increases, the $\mathcal{B}$ -induced surface cuts into the V -induced surface from above, by construction, and because $w(x_1) < b^{-1}(x_1)$ (Lemma A.4). Thus, the kinks along w are convex.

Step 2.1. On $\mathcal{I}$ and $\mathcal{V}$, the QVIs follow by the arguments leading up to the proof of Proposition 2 because the specification of ϕ on $\mathcal{I} \cup \mathcal{V}$ is the same here and in that proposition.

Step 2.2. On $\mathcal{A}$, the QVI requires that $A(x) \geq x_1 \vee x_2 = x_2$. Inequality $A(x) \geq V(x)$ follows by the optimality of stopping—or, as we say, by revealed preference (of the maximizer)—in the $\mathcal{A}$ -auxiliary problem. Inequality $A(x) \geq x_2$ follows by revealed preference in the V -auxiliary problem. Combining the preceding two inequalities gives $A(x) \geq x_2$.

Step 2.3. On $\mathcal{B}$, the QVI requires that $B(x) \geq x_2$. To verify the inequality, first, extend the $\mathcal{A}$ -auxiliary problem (originally defined on $\mathcal{A}$) to $\mathcal{B}$. Using the same arguments as in Lemma 4, this problem's solution can be verified to imply that, on $\mathcal{B}$, D_2 learns until either he observes the good news or his belief drops down to x_1 , at which point the prescription of the V -auxiliary problem is followed and delivers continuation value $V(x_1, x_1)$. By revealed preference, on $\mathcal{B}$, $A(x) \geq V(x)$ (where A is the value function of the $\mathcal{A}$ -auxiliary problem extended to $\mathcal{B}$). Moreover, on $\mathcal{B}$, the $\mathcal{B}$ -auxiliary problem has the same threshold (x_1 , by assumption) as the $\mathcal{A}$ -auxiliary problem extended to $\mathcal{B}$ (x_1 , now as a result). But once this threshold has been reached, it delivers a higher continuation value, $C(x_1)$, which satisfies $C(x_1) \geq V(x_1, x_1)$ by revealed preference in the C -auxiliary problem. So, on $\mathcal{B}$, $B(x) \geq A(x)$, which, combined with $A(x) \geq V(x) \geq x_2$, gives $B(x) \geq x_2$, as desired.

Step 2.4. On $\mathcal{C}$, the QVI requires that $C(z) \geq z$. Inequality $C(z) \geq V(z, z)$ follows by revealed preference in the C -auxiliary problem. Inequality $V(z, z) \geq z$ follows by revealed preference in the V -auxiliary problem. Combining the preceding two inequalities gives $C(z) \geq z$.

Step 3.1. On $\mathcal{I}$, the HJBs for D_1 and for D_2 hold by the argument in the proof of Proposition 2.

Step 3.2. On $\mathcal{V}$, the HJB for D_1 , (9), holds by construction.

HJB for D_2 holds if and only if $\Phi^A(x, c) \geq 0$, as was established in the discussion preceding (12). Recall that $\Phi^A(x, c) < 0 \iff x \in \mathcal{F}$. Because $\mathcal{F} \subset \mathcal{A} \cup \mathcal{B}$ (on $\hat{\mathcal{A}}$, $\mathcal{F} \subset \mathcal{A}$ by Lemma A.3; on $\hat{\mathcal{B}}$, $\mathcal{F} \subset \mathcal{B}$ by Lemma A.4) and $\mathcal{V} \cap (\mathcal{A} \cup \mathcal{B}) = \emptyset$ (by the definition of $\mathcal{V}$), $x \in \mathcal{V}$ implies that $x \notin \mathcal{F}$. As a result, $\Phi^A(x, c) \geq 0$ for all $x \in \mathcal{V}$, as desired.

Step 3.3. On $\mathcal{A}$, the HJB for D_2 , (17), holds by construction.

The HJB for D_1 is

$$-c - x_1(1 - x_1)A_1(x) + x_1(1 - A(x)) \leq 0. \quad (25)$$

By the envelope theorem applied to A in (18),

$$A_1(x) = \frac{(1 - x_2)V_1(x_1, d(x_1))}{1 - d(x_1)}.$$

Substituting the display above and the definition of A in (18) into (25) and dividing by c yields

$$\begin{aligned} & x_1(1 - x_2)(\eta(x_2) - \eta(d(x_1))) \\ & + \frac{x_1(1 - x_2)}{c(1 - d(x_1))} [1 - V(x_1, d(x_1)) - (1 - x_1)V_1(x_1, d(x_1))] \leq 1, \end{aligned}$$

which further simplifies by substituting V from (10) and V_1 from (9), and by dividing both sides by $1 - x_2$:

$$\frac{1}{1 - x_2} \geq x_1(\eta(x_2) - \eta(d(x_1))) + \frac{1}{1 - d(x_1)}. \quad (26)$$

If $x_2 = d(x_1)$, then (26) holds trivially, as equality. To show that (26) also holds for $x_2 > d(x_1)$, it suffices to show that its left-hand side increases in x_2 faster than its right-hand side does. Indeed, by $x_1 < x_2$, the left-hand side's derivative, $1/(1 - x_2)^2$, exceeds the right-hand side's derivative, $x_1/(x_2(1 - x_2)^2)$. Thus, (26) is verified.

Step 3.4. On $\mathcal{B}$, the HJB for D_2 , (21), holds by construction.

The HJB for D_2 is

$$-c - x_1(1 - x_1)B_1(x) + x_1(1 - B(x)) \leq 0. \quad (27)$$

Differentiating B in (22) and using the expression for $C'(x_1)$ gives

$$B_1(x) = \frac{1 - x_2}{(1 - x_1)^2} \left(1 - C(x_1) - \frac{c}{x_1} \right).$$

Substituting the display above and the definition of B in (22) into (27) leads to inequality (26) with $d(x_1)$ replaced by x_1 ; that inequality was verified in the preceding step (for any $d(x_1)$, including $d(x_1) = x_1$).

Because, by Steps 1, 2, and 3, ϕ is convex-kinked and satisfies HJBQVI, the conclusion of the proposition follows. $\square$

4. ALTERNATIVE LEARNING TECHNOLOGIES

So far, we have assumed that both divisions operate the *good-news technology* (GNT), which, when $v_i = 1$, sometimes reveals the good news that $v_i = 1$ but never reveals $v_i = 0$. Both from the conceptual standpoint and motivated by applications, one may wonder about HQ's optimal policy when both divisions operate the *bad-news technology* (BNT),

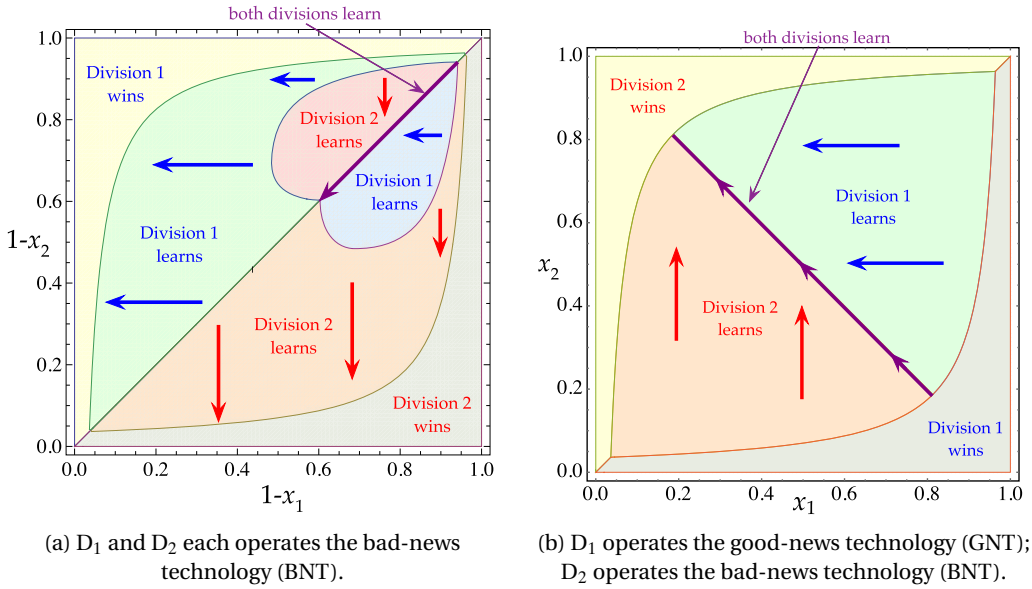


FIGURE 8. Optimal policies for bad-news and hybrid (good-news–bad-news) technologies; $c = 0.035$. Each arrow indicates the direction in which the type profile is revised if the division that learns observes no news. The region in which some agent learns is the same for good-news (not shown), bad-news, and hybrid learning technologies.

which, when $v_i = 0$, sometimes reveals the bad news that $v_i = 0$ but never reveals $v_i = 1$. Examples of the BNT are a clinical drug trial whose goal is to determine whether a drug has serious side effects, a press investigation whose goal is to discover a political candidate's disqualifying trait, and a company's due diligence about whether the Department of Justice will block a merger.

It turns out that the BNT case is a corollary (Corollary 1) to the GNT case (of Propositions 1, 2, and 3). The key to the result is to observe that, by the symmetry of the two technologies, the benefit from learning optimally relative to allocating immediately with the GNT when $\{\Pr\{v_i = 1\} = y_i\}_{i=1,2}$ is the same as with the BNT when $\{\Pr\{v_i = 0\} = y_i\}_{i=1,2}$ for any $(y_1, y_2) \in [0, 1]^2$. Formally, denoting by ω the value function for the BNT case, the symmetry between the GNT and BNT cases is captured by²⁰

$$\phi(x) - \max\{x_1, x_2\} \equiv \omega(1 - x) - \max\{1 - x_1, 1 - x_2\}, \quad (28)$$

where, as before, $x_i = \Pr\{v_i = 1\}$, $i = 1, 2$. Roughly speaking, what matters for the relative benefit of learning is how likely the states are that generate the news, not whether these states correspond to high or low project values.

The observation in (28) implies that, graphically, the optimal-policy map for the BNT is obtained from the optimal-policy map for the GNT (Figure 3(a)) by reversing the direction of each axis and by swapping the two divisions' areas for immediate allocation. Figure 8(a) illustrates.

²⁰See the proof of Corollary 1 for more on this symmetry.

COROLLARY 1. *Optimal policies for the BNT and the GNT are such that the following statements hold:*

- (i) *D_i learns with the BNT at a type profile x if and only if it learns with the GNT at the type profile $1 - x$.*
- (ii) *At any type profile x , D_i wins with the BNT if and only if it wins with the GNT.*

REMARK 1. The corollary's conclusion does not survive discounting. For a rough intuition, suppose that $r > 0$ and that x is "large," so that $\phi(x) > \omega(1 - x)$. Then the opportunity cost of learning—and, thus, delaying allocation—is higher at x with the GNT than at $1 - x$ with the BNT.

PROOF OF COROLLARY 1. Let HJBQVI-GNT and HJBQVI-BNT stand for HJBQVI equations for good-news and bad-news technologies, respectively.

Let ω , given in (28), be a conjectured value function for the BNT case. Equivalently, when $x_2 \geq x_1$,

$$\omega(x) \equiv \phi(1 - x) + x_1 + x_2 - 1. \quad (29)$$

By inspection of (29), the kinks of ω inherit the properties of the corresponding kinks of ϕ and, so, are convex.

It remains to verify that ω solves HJBQVI-BNT at the points of differentiability or that

$$0 = \max_{i \in \mathcal{N}} \{x_i - \omega(x), -c + x_i(1 - x_i)\omega_i(x) + (1 - x_i)(x_{-i} - \omega(x))\}$$

holds. Substituting (29) into the display above and setting $y = 1 - x$ gives

$$0 = \max_{i \in \mathcal{N}} \{y_{-i} - \phi(y), -c - y_i(1 - y_i)\phi_i(y) + y_i(1 - \phi(y))\},$$

which is HJBQVI-GNT, satisfied by ϕ .

The substitutions leading to the equivalence of HJBQVI-BNT and HJBQVI-GNT imply that D_i wins with the BNT at x whenever it would win with the GNT at x , and that D_i learns with the BNT at x whenever it would learn with the GNT at $1 - x$. $\square$

Finally, one can conceive of a hybrid technology, with D_1 operating the GNT and D_2 operating the BNT. An example is a drug trial in which D_1 tests an existing drug, known to be safe, for off-label efficacy, whereas D_2 tests a new drug, known to be efficacious, for side effects. Figure 8(b) illustrates the optimal policy. In this case, simultaneous learning occurs along the backward-bending diagonal segment. The region on which some division learns is the same in all three learning technologies considered in the paper.

5. DISCOUNTING

So far, our analysis of the optimal policy has focused on the undiscounted problem. This section suggests that the described results are robust to the introduction of discounting, and that discounting affects the optimal policy in intuitive ways. The comprehensive

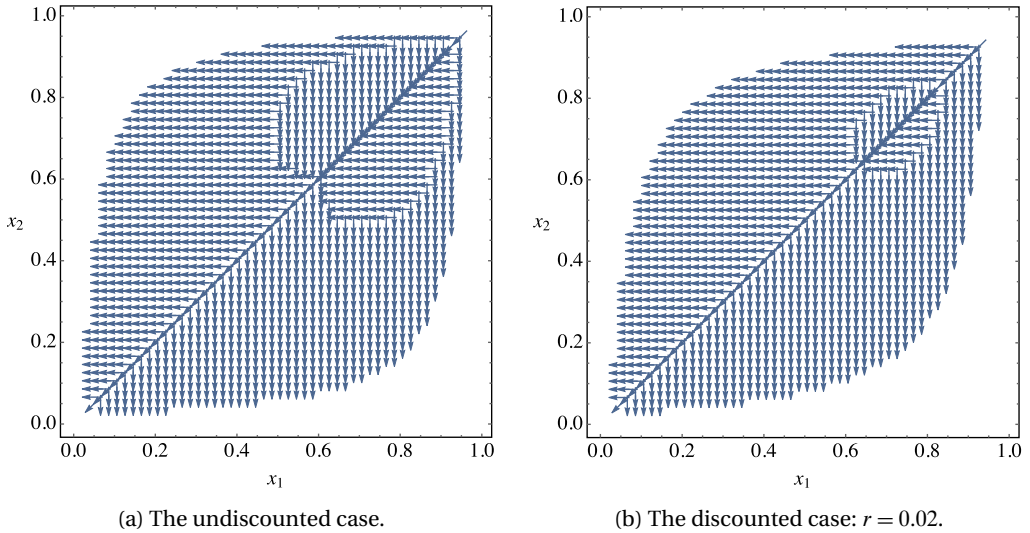


FIGURE 9. The optimal policy's prescription for each type profile when learning is cheap; $c = 0.035$. Within the lens-shaped region, one of the divisions learns. Each arrow indicates the direction in which the type profile is revised if the division that learns observes no news. Outside the lens-shaped region, no division learns and the higher-type division wins.

analysis of the discounted case is conceptually no different from the undiscounted one, but its execution is beyond both our ability to perform algebraic manipulations and this paper's scope. To illustrate the optimal policy with discounting, we resort to numerical analysis. Figure 9 illustrates the broad lessons.

According to Figure 9, an increase in the discount factor is qualitatively similar to an increase in the cost of learning. In particular, as r rises, the region on which some division learns (i.e., the lens-shaped region) shrinks.²¹ Furthermore, as r rises, the region on which the higher-type division learns (i.e., the heart-shaped region) also shrinks.

Discounting does not make simultaneous learning by both divisions more prevalent than it is in the undiscounted problem. Formally, whenever it is differentiable at a type profile x , the value function ϕ^r solves the HJB equation (4). Because the maximand in the HJB is linear in q , the allocation of learning effort, an interior solution is never strictly optimal.

Discounting would favor simultaneous learning if the model were changed so that learning capacity were division-specific instead of being fixed in the aggregate and allocated between the two divisions. In that case, HQ might tolerate the redundancy of simultaneous learning so as to avoid the delay associated with sequential learning.

6. CONCLUSIONS

The paper solves the cash-flow maximization problem of a company that faces an irreversible project-selection decision with information acquisition about each project. In

²¹This lens-shaped region is nonempty if and only if $c < (1 - r)^2/4$, which is the analogue of $c < \bar{c}$ in Proposition 1 and is derived analogously.

practice, strategic decisions pertaining to project selection (e.g., a merger or an acquisition) constitute sensitive information, which companies guard against outsiders. This lack of observability limits the scope for testing the model's predictions and makes the paper's focus largely normative.

Nevertheless, one can tentatively ask the positive questions of whether the actions assumed to be available to HQ are observed in practice and whether HQ's derived optimal strategy can rationalize observed outcomes. As noted in the Introduction, in 2011, Universal Music Group was choosing between two projects: buying EMI Music and buying Warner Music Group. Industry rumors suggest that, initially, Universal was learning about both projects simultaneously but quickly focused its efforts on learning about EMI. Late in 2011, Universal announced that it would buy EMI; Universal's consultants must have gotten good news. Universal's behavior is consistent with the model's optimal policy for the case when learning is cheap, so that simultaneous learning can be optimal. Thus, the actions of taking time to learn about projects and of learning either sequentially or simultaneously were available to Universal.

This paper derives the optimal policy by assuming that HQ both directly controls divisions' learning and observes divisions' news, if any. What if HQ can do neither? The optimal policy can still be implemented in a dynamic auction. This auction is a special case of the Vickrey–Clarke–Groves (VCG) mechanism's dynamic extensions (see [Athey and Segal 2013](#), and [Bergemann and Välimäki 2010](#)). Because of the good-news nature of the learning technology, our special case requires much less communication than a direct dynamic mechanism would suggest for a general learning technology. In particular, the auction begins with indicative bidding, followed by self-enforcing optimal learning by the divisions, and then firm bidding. Firm bidding occurs at a deadline that HQ sets, given the indicative bids, or as soon as either division calls for early firm bidding—whichever occurs first.²²

A dynamic auction is a plausible implementation instrument in practice. Internal auctions have been successfully deployed to predict sales (Hewlett–Packard), manage manufacturing capacity (Intel), generate business ideas (General Electric), select marketing campaigns (Starwood), and predict project completion (Microsoft) or external events (Google).²³ The efficacy of an internal auction relies on the company's ability to commit to refraining from subsequently undoing any payments received from its divisions in the course of the auction. One would expect the requisite commitment to be available to successful companies with a developed reputation for executing cash-flow-maximizing decisions.

REFERENCES

Athey, Susan and Ilya Segal (2013), “An efficient dynamic mechanism.” *Econometrica*, 81, 2463–2485. [141]

²²The details of the implementation are described in the working-paper version of this paper.

²³The Wikipedia entry “Prediction Market” and Cowgill and Zitzewitz (2015) contain further examples and evidence on the efficacy of internal auctions.

- Banks, Jeffrey S. and Rangarajan K. Sundaram (1992), "A class of bandit problems yielding myopic optimal strategies." *Journal of Applied Probability*, 29, 625–632. [118]
- Bardi, Martino and Italo Capuzzo-Dolcetta (1997), *Optimal Control and Viscosity Solutions of Hamilton–Jacobi–Bellman Equations*. Systems & Control: Foundations & Applications. Birkhäuser, Boston. [121]
- Bergemann, Dirk and Juuso Välimäki (2010), "The dynamic pivot mechanism." *Econometrica*, 78, 771–789. [141]
- Bolton, Patrick and Christopher Harris (1999), "Strategic experimentation." *Econometrica*, 67, 349–374. [118]
- Che, Yeon-Koo and Konrad Mierendorff (2017), "Optimal sequential decision with limited attention." Unpublished paper, University College London. [118]
- Cowgill, Bo and Eric Zitzewitz (2015), "Corporate prediction markets: Evidence from Google, Ford, and firm X." *Review of Economic Studies*, 82, 1309–1341. [141]
- de Motta, Adolfo (2003), "Managerial incentives and internal capital markets." *The Journal of Finance*, 58, 1193–1220. [118]
- Dixit, Avinash K. and Robert S. Pindyck (1994), *Investment Under Uncertainty*. Princeton University Press. [118]
- Forand, Jean Guillaume (2015), "Keeping your options open." *Journal of Economic Dynamics and Control*, 53, 47–68. [118]
- Harris, Milton and Artur Raviv (1996), "The capital budgeting process: Incentives and information." *The Journal of Finance*, 51, 1139–1174. [118]
- Inderst, Roman and Christian Laux (2005), "Incentives in internal capital markets: Capital constraints, competition, and investment opportunities." *The RAND Journal of Economics*, 36, 215–228. [118]
- Keller, Godfrey, Sven Rady, and Martin Cripps (2005), "Strategic experimentation with exponential bandits." *Econometrica*, 73, 39–68. [118]
- Klein, Nicolas A. and Sven Rady (2011), "Negatively correlated bandits." *Review of Economic Studies*, 78, 693–732. [118]
- Malenko, Andrey (2012), "Optimal design of internal capital markets." Unpublished paper, Massachusetts Institute of Technology. [118]
- Milgrom, Paul R. and Ilya Segal (2002), "Envelope theorems for arbitrary choice sets." *Econometrica*, 70, 583–601. [125]
- Murto, Pauli and Juuso Välimäki (2011), "Learning and information aggregation in an exit game." *The Review of Economic Studies*, 78, 1426–1461. [118]
- Oksendal, Bernt and Agn  s Sulem (2005), *Applied Stochastic Control of Jump Diffusions*. Springer, Berlin. [121]

- Ozbas, Oguzhan and David S. Scharfstein (2010), “Evidence on the dark side of internal capital markets.” *The Review of Financial Studies*, 23, 581–599. [115]
- Peskir, Goran and Albert Shiryaev (2006), *Optimal Stopping and Free-Boundary Problems*. Lectures in Mathematics. ETH Zürich. Birkhäuser, Basel. [118, 134]
- Pham, Huyen (1998), “Optimal stopping of controlled jump diffusion processes: A viscosity solution approach.” *Journal of Mathematical Systems, Estimation, and Control*, 8, 1–27. [120]
- Presman, Ernst L. and Isaac M. Sonin (1990), *Sequential Control With Incomplete Information: The Bayesian Approach to Multi-Armed Bandit Problems* (Elena Medova-Dempster and Michael Alan Howarth Dempster, eds.). Economic Theory, Econometrics, and Mathematical Economics. Academic Press, Cambridge, Massachusetts. [118]
- Protter, Philip E. (1990), *Stochastic Integration and Differential Equations: A New Approach*, first edition, Vol. 21 of Stochastic Modelling and Applied Probability. Springer, Berlin. [121]
- Rajan, Raghuram, Henri Servaes, and Luigi Zingales (2000), “The cost of diversity: The diversification discount and inefficient investment.” *The Journal of Finance*, 55, 35–80. [118]
- Rosenberg, Dinah, Eilon Solan, and Nicolas Vieille (2007), “Social learning in one-armed bandit problems.” *Econometrica*, 75, 1591–1611. [118]
- Scharfstein, David S. and Jeremy C. Stein (2000), “The dark side of internal capital markets: Divisional rent-seeking and inefficient investment.” *The Journal of Finance*, 55, 2537–2563. [118]
- Shiryaev, Albert N. (2008), *Optimal Stopping Rules*, volume 8 of Stochastic Modelling and Applied Probability. Springer, Berlin. [118]
- Stein, Jeremy C. (1997), “Internal capital markets and the competition for corporate resources.” *The Journal of Finance*, 52, 111–133. [118]
- Tirole, Jean (2006), *The Theory of Corporate Finance*. Princeton University Press, Princeton. [118]
- Wald, Abraham (1973), *Sequential Analysis*. Dover, Mineola, New York. [118]
- Webster, Thomas J. (2003), *Managerial Economics: Theory and Practice*. Academic Press, San Diego, California. [115]

Co-editor Johannes Hörner handled this manuscript.

Manuscript received 9 January, 2016; final version accepted 23 April, 2017; available online 1 May, 2017.