

Cover Page



Universiteit Leiden



The handle <http://hdl.handle.net/1887/33204> holds various files of this Leiden University dissertation

Author: Ommen, Thijs van

Title: Better predictions when models are wrong or underspecified

Issue Date: 2015-06-10

Chapter 6

Worst-Case Optimal Probability Updating

The final three chapters of this dissertation discuss worst-case optimal probability updating, an alternative to conditioning that may be used when the distribution is not fully specified. In Chapter 6, we introduce the problem, and find how optimal solutions may be recognized for different loss functions; our main tool is convex analysis. We find that for logarithmic loss, optimality is characterized by the elegant *RCAR* (*reverse coarsening at random*) condition.

In Chapter 7, we analyse the combinatorial aspect of the probability updating problem, and present some theoretical tools that may help us find worst-case optimal solutions to a probability updating problem, as opposed to merely recognizing such solutions. Further, we see that the applicability of the RCAR condition is not restricted to the cases discovered in Chapter 6, and explore the consequences.

In Chapter 8, we give algorithms that automate the task of finding worst-case optimal solutions, for restricted classes of probability updating problems.

A more detailed overview of this first chapter will be provided in Section 6.1.2.

6.1 Introduction

There are many situations in which a decision maker receives incomplete data and still has to reach conclusions about these data. One type of incomplete data is *coarse* data: instead of the real outcome of a random event, the decision maker observes a subset of the possible outcomes, and knows only that the actual outcome is an element of this subset. An example frequently occurs in questionnaires, where people may be asked if their date of birth lies between 1950 and 1960 or between 1960 and 1970 et cetera. Their exact year of birth is unknown to us, but at least we now know for sure in which decade they

are born. We introduce another instance of coarse data with the following example.

Example 6.A (Fair die). Suppose I throw a fair die. I get to see the result of the throw, but you do not. Now I tell you that the result lies in the set $\{1, 2, 3, 4\}$. This is an example of coarse data. You know that I used a fair die and that what I tell you is true. Now you are asked to give the probability that I rolled a 3. Likely, you would say that the probability of each of the remaining possible results is $1/4$. This is the knee-jerk reaction of someone who studied probability theory, since this is standard *conditioning*. But is this always correct?

Suppose that there is only one alternative set of results I could give you after rolling the die, namely the set $\{3, 4, 5, 6\}$. I can now follow a *coarsening mechanism*: a procedure that tells me which subset to reveal given a particular result of the die roll. If the outcome is 1, 2, 5, or 6, there is nothing for me to choose. Suppose that if the outcome is 3 or 4, the coarsening mechanism I use selects set $\{1, 2, 3, 4\}$ or set $\{3, 4, 5, 6\}$ at random, each with probability $1/2$. If I throw the die 6000 times, I expect to see the outcome 3 a thousand times. Therefore I expect to report the set $\{1, 2, 3, 4\}$ five hundred times after I see the outcome 3. It is clear that I expect to report the set $\{1, 2, 3, 4\}$ 3000 times in total. So for die rolls that I told you resulted in an outcome in $\{1, 2, 3, 4\}$, the probability of the true outcome being 3 is actually $1/6$ with this coarsening mechanism. We see that the prediction in the first paragraph was not correct, in the sense that the probabilities computed there do not correspond to the long-run relative frequencies. We conclude that the knee-jerk reaction is not always correct.

In Example 6.A we have seen that standard conditioning does not always give the correct answers. Heitjan and Rubin (1991) answer the question under what circumstances standard conditioning of coarse data is correct. They discovered a necessary and sufficient condition of the coarsening mechanism, called *coarsening at random* (CAR). A coarsening mechanism satisfies the CAR condition if, for each subset y of the outcomes, the probability of choosing to report y is the same no matter which outcome $x \in y$ is the true outcome. In many situations, however, this condition cannot hold, as proved by Grünwald and Halpern (2003). It depends on the arrangement of possible revealed subsets whether a coarsening mechanism exists that satisfies CAR.

We continue with an example generating much debate among both the general public and professional probabilists: the Monty Hall puzzle, posed by Selvin (1975) and popularized years later when it appeared in Ask Marilyn, a weekly column in Parade Magazine by Marilyn vos Savant, in September 1990 (Gill, 2011). In this example, no coarsening mechanism satisfies the CAR condition.

Example 6.B (Monty Hall). Suppose you are on a game show and you may choose one of three doors. Behind one of the doors a car can be found, but the other two only hide a goat. Initially the car is equally likely to be behind each of the doors. After you have picked one of the doors, the host Monty Hall, who

knows the location of the prize, will open one of the other doors, revealing a goat. Now you are asked if you would like to switch from the door you chose to the other unopened door. Is it a good idea to switch?

At this moment we will not answer this question, but we show that the problem of choosing whether to switch doors is an example of the coarse data problem. The unknown random value we are interested in is the location of the car: one of the three doors. When the host opens a door different from the one you picked, revealing a goat, this is equivalent to reporting a subset. The subset he reports is the set of the two doors that are still closed. For example, if he opens door 2, this tells us that the true value, the location of the car, is in the subset $\{1, 3\}$. Note that if you have by chance picked the correct door, there are two possible doors Monty Hall can open, so also two subsets he can report. This implies that Monty has a choice in reporting a subset. How does Monty's coarsening mechanism influence your prediction of the true location of the car?

The CAR condition can only be satisfied for very particular distributions of where the prize is: the probability that the prize is hidden behind the initially chosen door must be either 0 or 1, otherwise no CAR coarsening mechanism exists (Grünwald and Halpern, 2003, Example 3.3)¹. If the prize is hidden in any other way, for example uniformly at random as we assume, then CAR cannot hold, and standard conditioning will result in an incorrect conclusion for at least one of the two subsets.

Examples 6.A and 6.B are just two instances of a more general problem: the number of outcomes may be different; the initial distribution of the true outcome may be any distribution; and the subsets of outcomes that may be reported to the decision maker may be any family of sets. Our goal in this chapter is to define general procedures that tell us how to update the probabilities on the outcomes after making a coarse observation, in situations where standard conditioning is not adequate.

These probabilities may be either frequentist or Bayesian probabilities. In Example 6.A, they were frequentist probabilities: the original distributions were known, and the updated probabilities we found were again frequencies over many repetitions of the same experiment. The original distribution of the outcomes could also express our Bayesian prior belief of how likely each outcome is. This is the more powerful interpretation in Example 6.B, where the uniform distribution of the location of the car requires an assumption on the frequentist's part, while it is a reasonable choice of prior for a Bayesian (Gill, 2011). In this case, the updated probabilities after a coarse observation take the role of the Bayesian posterior distribution. In either case, we will refer to the initial probability of an outcome, regardless of observations, as the *marginal* probability.

Without any assumptions on the quizmaster's strategy (i.e. the coarsening mechanism), the conditional distributions of outcomes given observations will

¹This uses the *weak* version of CAR in the terminology of Jaeger (2005a), in which outcomes with probability 0 are exempt from the equality constraint. A *strong* CAR coarsening mechanism does not exist regardless of the probabilities with which the prize is hidden.

be unknown, and this uncertainty cannot be fully expressed by a single probability distribution over the outcomes. So to get a single answer to the question how to update our probabilities, we need to make some assumption about how the quizmaster chooses his strategy. Assuming that the coarsening mechanism satisfies CAR is one such approach, but as we saw in the two examples, there are scenarios where this assumption cannot hold. We instead take a *worst-case approach*, treating the coarsening of the observation and the subsequent probability update as a game between two players: the *quizmaster* and the *contestant* (named for their roles in the Monty Hall scenario). The subset of outcomes communicated by the quizmaster to the contestant will be called the *message*.

In this fictional game, the quizmaster's goal is the opposite of the contestant's, namely to make predicting the true outcome as hard as possible for the contestant. Such situations are rare in practice: The sender of a message might be motivated by interests other than informing us (for example, a newspaper may be trying to optimize its sales figures, or a company may want to present its performance in the best light), but rarely by trying to be as uninformative as possible. (Though see Section 7.4.3 in the next chapter, where we consider the case that the players' goals are not diametrically opposed.) In other situations, the 'sender' might not be a rational being at all, but just some unknown process. Yet this game is a useful way to look at the problem of updating our probabilities even if we do not believe that the coarsening mechanism is chosen adversarially: if we simply do not know how 'nature' chooses which message to give us and do not want to make any assumptions about this, then choosing the worst-case (or *minimax*) optimal probability update as defined here guarantees that we get at most some fixed expected loss, while any other probability update may lead to a larger expected loss depending on the unknown coarsening mechanism.

We will need a loss function to measure how well the quizmaster and the contestant are doing at this game. Our results apply to a wide variety of loss functions. For an analysis of the Monty Hall game, 0-1 loss would be appropriate, as the contestant must choose a single door; this is the approach used in Gill (2011); Gnedin (2011). Other loss functions, such as logarithmic loss and Brier loss, also allow the contestant to formulate their prediction of where the prize is hidden as an arbitrary probability distribution over the outcomes. For the Monty Hall game with one of these two loss functions, the worst-case optimal answer for the contestant is to put probability $1/3$ on his initially chosen door and $2/3$ on the other door. (These probabilities agree with the literature on the Monty Hall game.) Surprisingly, we will see (in Example 6.D on page 129) that for similar games, logarithmic and Brier loss may lead to two different answers!

We will find in this chapter that for finite outcome spaces, both players in our game have worst-case optimal strategies for many loss functions: the quizmaster has a strategy that makes the contestant's prediction task as hard as possible, and the contestant has a strategy that is guaranteed to give good predictions no matter what the quizmaster does. We give characterizations that allow us to recognize such strategies, for different conditions on the loss

functions. Interestingly, Theorem 6.10 shows that the worst-case optimal prediction under logarithmic loss satisfies a property that has a striking similarity with the CAR condition, but switches the roles of outcome and message. By Lemma 6.14, if a betting game is played repeatedly and the contestant is allowed to distribute investments over different outcomes and to reinvest all capital gained so far in each round, then the same strategy is optimal, *regardless of the pay-offs!*

Example 6.A (continued). For logarithmic loss, the worst-case optimal prediction of the die roll conditional on the revealed subset is found with the help of Theorem 6.10. The worst-case optimal prediction given that you observe the set $\{1, 2, 3, 4\}$ is: predict outcomes 1 and 2 each with probability $1/3$, and predict 3 and 4 each with probability $1/6$. Given that you observe the set $\{3, 4, 5, 6\}$, the worst-case optimal prediction is: 3 and 4 with probability $1/6$, and 5 and 6 with probability $1/3$.

These probabilities correspond with the coarsening mechanism given earlier. However, it is a good prediction even if you do not know what coarsening mechanism I am using. An intuitive argument for this is the following: If I wanted, I could use a very extreme coarsening mechanism, always choosing to reveal the set $\{1, 2, 3, 4\}$ when the die comes up 3 or 4. But this is balanced by the possibility that I might be using the opposite coarsening mechanism, which always reveals $\{3, 4, 5, 6\}$ if the result is 3 or 4. The worst-case optimal prediction given above hedges against both possibilities.

6.1.1 Caveats on the use of the word ‘conditioning’

Since this chapter is concerned with making a worst-case optimal prediction conditional on a set of outcomes, we want to highlight the use of the word *conditioning*. Above, we used the word *standard* conditioning for using the conditional information in the standard way: with random variables X the true outcome and Y the coarse observation (a set of outcomes), computing $P(X = x \mid Y = y)$ by $P(X = x \mid X \in y) = P(X = x) / P(X \in y)$.

In the two examples, we saw that this does not always give the correct probabilities given a coarse observation. For instance in Example 6.A, $P(X \mid X \in y)$ is the uniform distribution, while $P(X \mid Y = y)$ cannot simultaneously be uniform for both $y = \{1, 2, 3, 4\}$ and $y = \{3, 4, 5, 6\}$. In such situations, we call this computation *naïve conditioning*. The formula for conditional probabilities does give the correct result if instead of on the outcome space, we work in a larger space: the space of all combinations of an outcome and a set. The problem we face in this chapter is that we do not know the probabilities of all these combinations, as they depend on the unknown coarsening mechanism.

In the case of Bayesian probabilities, the question is how to extend our prior on the outcomes to a prior on the larger space. The worst-case optimal coarsening mechanism we characterize in our theorems can be seen as a recommendation for such a prior.

6.1.2 Contents

In Section 6.2, we will give a precise definition of the ‘conditioning game’ we described. In Section 6.3, we find general conditions on the loss function under which worst-case optimal strategies for the quizmaster exist, and we characterize such strategies. Section 6.4 does the same for worst-case optimal strategies for the contestant. (See Figure 6.4 for a visual illustration of the concepts used in these sections.) If stronger conditions hold, worst-case optimal strategies for both players may be easier to recognize. This is explored for two classes of loss functions in Section 6.5; in particular, we find in Section 6.5.2 that for a class of loss functions including logarithmic loss, worst-case optimal strategies for the quizmaster are characterized by a simple condition on their probabilities: the *RCAR* (*reverse CAR*) condition. An overview of the theorems and the conditions under which they apply is given in Table 6.1 on page 125. Many examples are included to illustrate (the limits of) the theoretical results. Section 6.6 gives some concluding remarks.

All proofs are given in Appendix 6.A at the end of this chapter.

This work is an extension of Feenstra (2012) to loss functions other than logarithmic loss, and to the case where the worst-case optimal strategy for the quizmaster assigns probability 0 to some combinations of outcomes x and messages y with $x \in y$. It can also be seen as a concrete application of the ideas in Grünwald and Dawid (2004) about minimax optimal decision making and its relation to entropy.

6.2 Definitions and problem formulation

A (*probability updating*) game $\mathcal{G}$ is defined as a quadruple $(\mathcal{X}, \mathcal{Y}, p, L)$, where $\mathcal{X}$ is a finite set, $\mathcal{Y}$ is a family of distinct subsets of $\mathcal{X}$ with $\bigcup_{y \in \mathcal{Y}} y = \mathcal{X}$, p is a nowhere-zero probability mass function on $\mathcal{X}$, and L is a function $L : \mathcal{X} \times \Delta_{\mathcal{X}} \rightarrow [0, \infty]$, where $\Delta_{\mathcal{X}}$ is the set of all probability mass functions on $\mathcal{X}$. We call $\mathcal{X}$ the *outcome space*, $\mathcal{Y}$ the *message structure*, p the *marginal distribution*, and L the *loss function*. We shall discuss choices we made in this definition in Section 6.2.3, and first work out the definition in detail.

Example 6.B (continued). We assume the car is hidden uniformly at random behind one of the three doors. With this assumption, we can abstract away the initial choice of a door by the contestant: by symmetry, we can assume without loss of generality that he always picks door 2. Then the probability updating game starts with the quizmaster opening door 1 or 3, thereby giving the message “the car is behind door 2 or 3” or “the car is behind door 1 or 2”, respectively. This can be expressed as follows in our formalization:

- outcome space $\mathcal{X} = \{x_1, x_2, x_3\}$;
- message space $\mathcal{Y} = \{y_1, y_2\}$ with $y_1 = \{x_1, x_2\}$ and $y_2 = \{x_2, x_3\}$;
- marginal distribution p uniform on $\mathcal{X}$.

If a loss function L is also given, this fully specifies a game. One example is randomized 0-1 loss, which is given by $L(x, Q) = 1 - Q(x)$. Here x is the true outcome, and Q is the contestant's prediction of the true outcome in the form of a probability distribution. Thus the prediction is awarded a smaller loss if it assigned a larger probability $Q(x)$ to the outcome that actually obtained. We will see other examples of loss functions in Section 6.2.2.

A function from some finite set S to $\mathbf{R}$ corresponds to an $|S|$ -dimensional vector when we fix an order on the elements of S . We write $\mathbf{R}^S$ for the set of such functions/vectors. Even if no order on S is specified, this allows us to apply concepts from linear algebra to $\mathbf{R}^S$ without ambiguity. For example, we may say that some set is an affine subspace of $\mathbf{R}^S$. (This identification and the resulting notation are also used by Schrijver (2003a).)

Using this correspondence, we identify the elements of $\Delta_{\mathcal{X}}$ with the $|\mathcal{X}|$ -dimensional vectors in the unit simplex, though we use ordinary function notation $P(x)$ for its elements. The probability mass function p that is part of a game's definition is also a vector in $\Delta_{\mathcal{X}}$. Vector notation p_x will be used to refer to its elements to set p apart from P , which will denote distributions chosen by the quizmaster rather than fixed by the game.

For any message $y \subseteq \mathcal{X}$, we define $\Delta_y = \{P \in \Delta_{\mathcal{X}} \mid P(x) = 0 \text{ for } x \notin y\}$. Note that these are vectors of the same length as those in $\Delta_{\mathcal{X}}$, though contained within a lower-dimensional affine subspace.

A loss function L is called *proper* if $\arg \min_{Q \in \Delta_{\mathcal{X}}} \mathbf{E}_{X \sim P} L(X, Q) = P$ for all $P \in \Delta_{\mathcal{X}}$, and *strictly proper* if this minimizer is unique (this is standard terminology; see for instance Gneiting and Raftery (2007)). Thus if a predicting agent believes the true distribution of an outcome to be given by some P , such a loss function will encourage him to report $Q = P$ as his prediction.

6.2.1 Strategies

Strategies for the players are specified by conditional distributions: a strategy P for the quizmaster consists of distributions on $\mathcal{Y}$, one for each possible $x \in \mathcal{X}$, and a strategy Q for the contestant consists of distributions on $\mathcal{X}$, one for each possible $y \in \mathcal{Y}$. These strategies define how the two players act in any situation: the quizmaster's strategy defines how he chooses a message containing the true outcome (the coarsening mechanism), and the contestant's strategy defines his prediction for each message he might receive.

We write $P(\cdot \mid x)$ for the distribution on $\mathcal{Y}$ the quizmaster plays when the true outcome is $x \in \mathcal{X}$. Because $p_x > 0$, this conditional distribution can be recovered from the joint $P(x, y) := P(y \mid x)p_x$; we will use this joint distribution to specify a strategy for the quizmaster. If $P(y) := \sum_{x \in \mathcal{X}} P(x, y) > 0$, we may also write $P(\cdot \mid y)$ for the vector in Δ_y given by $P(x \mid y) := P(x, y) / P(y)$. No such rewrites can be made for Q , as no marginal $Q(y)$ is specified by the game or by the strategy Q . To shorten notation and to emphasize that Q is not a joint distribution, we write $Q_{|y}$ rather than $Q(\cdot \mid y)$ for the distribution that the contestant plays in response to message y .

We restrict the quizmaster to conditional distributions P for which $P(y \mid x) = 0$ if $x \notin y$; that is, he may not ‘lie’ to the contestant. We make no similar requirement on the contestant’s choice of Q , though for proper loss functions, and in fact all other loss functions we will consider in our examples, the contestant can gain nothing from using a strategy Q for which $Q_{|y}(x) > 0$ where $x \notin y$.

Example 6.B (continued). The table below specifies all aspects of a game except for its loss function: its outcome space (here, for the Monty Hall game, $\mathcal{X} = \{x_1, x_2, x_3\}$), message space ($\mathcal{Y} = \{y_1, y_2\}$ with $y_1 = \{x_1, x_2\}$ and $y_2 = \{x_2, x_3\}$) and marginal distribution (p uniform on $\mathcal{X}$).

P	x_1	x_2	x_3
y_1	1/3	1/6	—
y_2	—	1/6	1/3
p_x	1/3	1/3	1/3

(6.1)

In this table we have filled in a strategy P for the quizmaster in the form of a joint distribution on pairs of x and y . The cells in the table where $x \notin y$ are marked with a dash to indicate that P may not assign positive probability there. The probabilities in each column sum to the marginal probabilities at the bottom, so this joint distribution P has the correct marginal distribution on the outcomes. For this particular strategy, if the true outcome is x_2 , the quizmaster will give message y_1 or y_2 to the contestant with equal probability.

More formally, write $\mathcal{R}(\mathcal{X}, \mathcal{Y})$ as an abbreviation for the set of pairs $\{(x, y) \mid x \in y \in \mathcal{Y}\}$. In the case of the Monty Hall game, there are four such pairs: $\mathcal{R}(\mathcal{X}, \mathcal{Y}) = \{(x_1, y_1), (x_2, y_1), (x_2, y_2), (x_3, y_2)\}$. The notation $\mathbf{R}_{\geq 0}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$ represents the set of all functions from $\mathcal{R}(\mathcal{X}, \mathcal{Y})$ to $\mathbf{R}_{\geq 0}$. If P is an element of this set and $(x, y) \in \mathcal{R}(\mathcal{X}, \mathcal{Y})$, the value of P at (x, y) is denoted by $P(x, y)$. For (x, y) with $x \notin y$, the notation $P(x, y)$ does not correspond to a value of the function, but is taken to be 0.

We again identify the elements of $\mathbf{R}_{\geq 0}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$ with vectors. Thus the mass function P shown in (6.1) is identified with a four-element vector $(1/3, 1/6, 1/6, 1/3)$. (We could have chosen a different ordering instead.)

We define the set $\mathcal{P}$ of strategies for the quizmaster as $\{P \in \mathbf{R}_{\geq 0}^{\{x \in y\}} \mid \sum_{y \ni x} P(x, y) = p_x \text{ for all } x\}$; this is a convex set. The set of strategies for the contestant is $\mathcal{Q} := \Delta_{\mathcal{X}}^{\mathcal{Y}} = \{(Q_{|y})_{y \in \mathcal{Y}} \mid Q_{|y} \in \Delta_{\mathcal{X}} \text{ for each } y \in \mathcal{Y}\}$.

For given strategies P and Q , the expected loss the contestant incurs is

$$\begin{aligned} \sum_{x \in \mathcal{X}} p_x \sum_{\substack{y: \\ x \in y \in \mathcal{Y}}} P(y \mid x) L(x, Q_{|y}) &= \mathbf{E}_{X \sim p} \mathbf{E}_{Y \sim P(\cdot \mid X)} L(X, Q_{|Y}) \\ &= \mathbf{E}_{(X, Y) \sim P} L(X, Q_{|Y}). \end{aligned} \quad (6.2)$$

We allowed L to take the value ∞ ; if this value occurs with positive probability, then the contestant’s expected loss is infinite. However, for terms where the

probability is zero, we define $0 \cdot \infty = 0$, as is consistent with measure-theoretic probability.

We approach this problem as a zero-sum game between two players: the quizmaster chooses $P \in \mathcal{P}$ to maximize (6.2), while simultaneously (that is, without knowing P) the contestant chooses $Q \in \mathcal{Q}$ to minimize that quantity. The game $(\mathcal{X}, \mathcal{Y}, p, L)$ is common knowledge for the two players.

If the contestant knew the quizmaster's strategy, he would pick a strategy Q that for each y minimizes the expected loss of predicting x given y . When the contestant receives a message and knows the distribution $P \in \Delta_{\mathcal{X}}$ over the outcomes given that message, this expected loss is written as

$$H_L(P) := \inf_{Q \in \Delta_{\mathcal{X}}} \sum_x P(x) L(x, Q) = \inf_{Q \in \Delta_{\mathcal{X}}} \mathbf{E}_{X \sim P} L(X, Q). \quad (6.3)$$

This is the *generalized entropy* of P for loss function L (Grünwald and Dawid, 2004). (Note that in the preceding display, P and Q are not strategies but simply distributions over $\mathcal{X}$.) If the contestant picks his strategy Q this way, (6.2) becomes the *expected generalized entropy* of the quizmaster's strategy $P \in \mathcal{P}$:

$$\sum_{y \in \mathcal{Y}} P(y) H_L(P(\cdot | y)), \quad (6.4)$$

where we again define terms with $P(y) = 0$ as 0. We say a strategy P is *worst-case optimal for the quizmaster* if it maximizes this expected generalized entropy over all $P \in \mathcal{P}$. We call the version of the game where the quizmaster has to play first the *maximin* game, where the order of the words 'max' and 'min' reflects the order in which they appear in the expression for the value of this game as well as the order in which the maximizing and minimizing players take their turns.

Similarly, if the contestant were to play first (the *minimax* game), his goal might be to find a strategy Q that minimizes his worst-case expected loss

$$\max_{P \in \mathcal{P}} \sum_{x \in \mathcal{X}} P(x, y) L(x, Q_{|y}) = \max_{P \in \mathcal{P}} \mathbf{E}_{(X, Y) \sim P} L(X, Q_{|Y}). \quad (6.5)$$

(In this case, the maximum is always achieved so we can write max rather than sup: for each x , the quizmaster can choose P that puts all mass on a $y \ni x$ with the maximum loss.) We call a strategy *worst-case optimal for the contestant* if it achieves the minimum of (6.5).

It is an elementary result from game theory that if worst-case optimal strategies P^* and Q^* exist for the two players, their expected losses are related by

$$\sum_{y \in \mathcal{Y}} P^*(y) H_L(P^*(\cdot | y)) \leq \max_{P \in \mathcal{P}} \sum_{x \in \mathcal{X}} P(x, y) L(x, Q^*_{|y}) \quad (6.6)$$

(Rockafellar, 1970, Lemma 36.1: "maximin $\leq$ minimax"). The inequality expresses that in a sequential game where one of the players knows the other's strategy before choosing his own, the player to move second may have an advantage.

In the next section, we will see that in many probability updating games, worst-case optimal strategies for both players exist (but may not be unique), and the maximum expected generalized entropy *equals* the minimum worst-case expected loss:

$$\sum_{y \in \mathcal{Y}} P^*(y) H_L(P^*(\cdot | y)) = \max_{P \in \mathcal{P}} \sum_{x \in \mathcal{X}} P(x, y) L(x, Q^*_{|y}). \quad (6.7)$$

When this is the case, we say that the minimax theorem holds (von Neumann, 1928; Ferguson, 1967). We remark here that our setting, while a zero-sum game, differs from the usual setting of zero-sum games in some respects: We consider possibly infinite loss and (in general) infinite sets of strategies available to the players, but do not allow the players to randomize over these strategies. Randomizing over $\mathcal{P}$ would not give the quizmaster an advantage, as $\mathcal{P}$ is convex and he could just play the corresponding convex combination directly; because (6.2) is linear in P , this results in the same expected loss. (Another way to view this is that, essentially, the quizmaster *is* randomizing, over a finite set of strategies.) For the contestant, $\mathcal{Q}$ is also convex, but in general (depending on L), playing a convex combination of strategies does not correspond to randomizing over those strategies. The two do correspond in the case of randomized 0-1 loss, where L is linear. If L is convex, then playing the convex combination is at least as good for him as randomizing (and if L is strictly convex, better), so allowing randomization would again not give an advantage.

When (6.7) holds, any pair of worst-case optimal strategies (P^*, Q^*) forms a (*pure strategy*) *Nash equilibrium*, a concept introduced by Nash (1951): neither player can benefit from deviating from their worst-case optimal strategy if the other player leaves his strategy unchanged. This means that the definitions of worst-case optimality given above are also meaningful in the game we are actually interested in, where the players move simultaneously in the sense that neither knows the other's strategy when choosing his own.

6.2.2 Three standard loss functions

Three commonly used loss functions are logarithmic loss, Brier loss, and randomized 0-1 loss. These are defined as follows (Grünwald and Dawid, 2004):

Logarithmic loss is a strictly proper loss function, given by

$$L(x, Q) = -\log Q(x).$$

Its entropy is the Shannon entropy $H_L(P) = \sum_x -P(x) \log P(x)$. The functions L and H_L are displayed in Figure 6.1 for the case of a binary prediction (i.e. a prediction between two possible outcomes). The (three-dimensional) graph of H_L for the case of three outcomes will appear in Figure 6.4 on page 128.

Brier loss is another strictly proper loss function, corresponding to squared Euclidean distance:

$$L(x, Q) = \sum_{x' \in \mathcal{X}} (\mathbf{1}_{x'=x} - Q(x'))^2 = (1 - Q(x))^2 + \sum_{x' \in \mathcal{X}, x' \neq x} Q(x')^2.$$

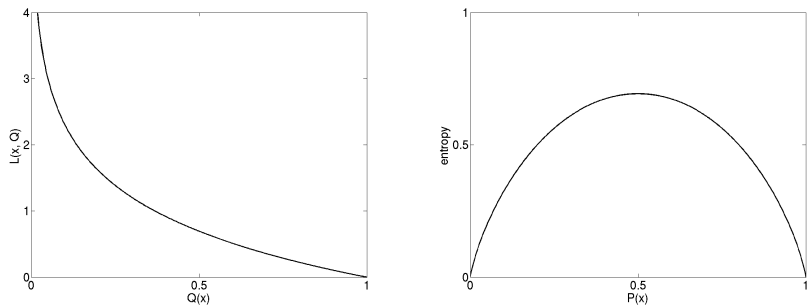


Figure 6.1: Logarithmic loss and entropy (natural base) on a binary prediction

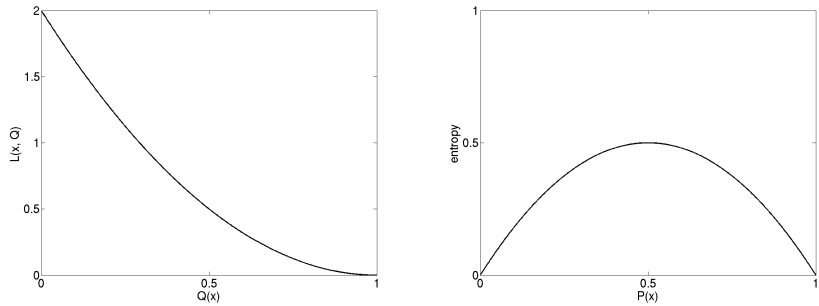


Figure 6.2: Brier loss and entropy on a binary prediction

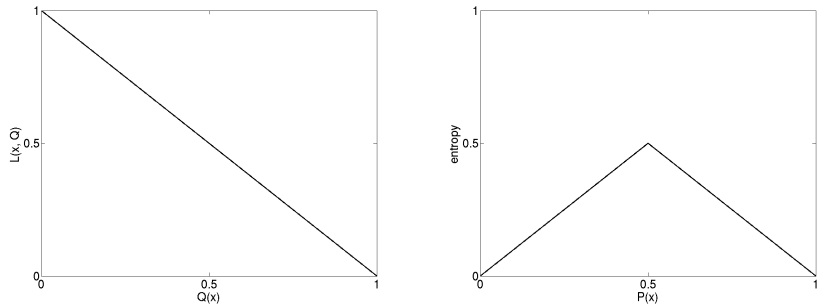


Figure 6.3: Randomized 0-1 loss and entropy on a binary prediction

Its entropy function is $H_L(P) = 1 - \sum_{x \in \mathcal{X}} P(x)^2$; L and H_L are displayed in Figure 6.2 for a binary prediction. Because L is a function of the entire distribution Q and not only of $Q(x)$, the graph for L does not fully capture the behaviour of Brier loss on non-binary predictions.

The third loss function we will often refer to is *randomized 0-1 loss*, given by

$$L(x, Q) = 1 - Q(x).$$

It is improper: an optimal response Q to some distribution P puts all mass on outcome(s) with maximum $P(x)$. Its entropy function is $H_L(P) = 1 - \max_{x \in \mathcal{X}} P(x)$ (see Figure 6.3). It is related to *hard 0-1 loss*, which requires the contestant to pick a single outcome x' and gives loss 0 if $x' = x$ and 1 otherwise. Randomized 0-1 loss essentially allows the contestant to randomize his prediction: $L(x, Q)$ equals the expected value of hard 0-1 loss when x' is distributed according to Q . An important difference between games with hard and randomized 0-1 loss will be shown later in Example 6.F.

6.2.3 Notes on our definition

Our definition of a game rules out duplicate messages in $\mathcal{Y}$, which would not meaningfully change the options of either player as the two messages represent the same move for the quizmaster; this will be made precise in Lemma 6.2. The definition does allow duplicate outcomes: pairs of outcomes $x_1, x_2 \in \mathcal{X}$ such that $x_1 \in y$ if and only if $x_2 \in y$ for all $y \in \mathcal{Y}$. We will see later (in Example 6.D) that games with such outcomes cannot generally be solved in terms of games without, and thus we must analyse them in their own right.

We also ruled out the existence of outcomes with marginal probability zero. A game with such a marginal is equivalent to a game with the corresponding outcomes removed from the outcome space and from all messages; this modified game does have positive marginal probability on all outcomes.

6.3 Worst-case optimal strategies for the quizmaster

In this section, we formulate the problem of finding a worst-case optimal strategy for the quizmaster as a convex optimization problem. Worst-case optimal strategies for the contestant will be discussed in Section 6.4. In order to be applicable to a wide range of loss functions, these two sections are rather technical, and the characterizations of worst-case optimal strategies we find here are not always easy to use (though the abstract results in these sections are illustrated by concrete examples in Sections 6.3.1 and 6.4.3). We will find simpler characterizations for smaller classes of loss functions in Section 6.5. An overview of these results is given in Table 6.1.

For the convex optimization problem we will discuss in this section, it will prove convenient to allow P to range over $\mathbf{R}_{\geq 0}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$. That is, we allow arbitrary vectors of nonnegative reals that may disobey the marginal constraint,

Table 6.1: Results on worst-case optimal strategies for different loss functions

Conditions on L	Results	Example
H_L finite and continuous	P^* exists and is characterized by Theorem 6.3	hard 0-1 loss
H_L finite and continuous; all minimal supporting hyperplanes realizable	Q^* exists, a Nash equilibrium exists (Theorem 6.5); Q^* characterized by Theorem 6.7	randomized 0-1 loss
L proper and continuous; H_L finite and continuous	all the above simplified by Theorem 6.9	Brier loss
L local and proper; H_L finite and continuous	characterization of P^* simplified further by Theorem 6.10 (RCAR condition)	logarithmic loss

and might not even sum to one where we would normally expect a probability distribution (though we do still have $P(x, y) = 0$ for $x \notin y$). However, when $P(y) > 0$ for some y , $P(\cdot \mid y)$ defined by $P(x \mid y) := P(x, y)/P(y)$ as in Section 6.2.1 still defines a probability distribution, because the scale factor cancels out. The constraints in the optimization problem will then enforce that our solution P lies in $\mathcal{P}$; such a P is called a *feasible* solution.

The following function extends the quizmaster's objective function (6.4) (the expected generalized entropy of $P \in \mathcal{P}$) to the domain $\mathbf{R}_{\geq 0}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$:

$$f_0(P) := \inf_{Q \in \mathcal{Q}} \sum_{y \in \mathcal{Y}, x \in y} P(x, y) L(x, Q|_y). \quad (6.8)$$

Note that while the quizmaster is given more freedom in that we allow $P \notin \mathcal{P}$ as argument to f_0 , the minimization (representing the contestant's choice) is still restricted to the set of strategies $\mathcal{Q}$ defined in the previous section. Because the conditional distributions that make up $\mathcal{Q}$ can be chosen independently for each y , we can rewrite f_0 in terms of ordinary generalized entropies as follows:

$$f_0(P) = \sum_{\substack{y \in \mathcal{Y}: \\ P(y) > 0}} \inf_{\substack{Q|_y \in \Delta_{\mathcal{X}} \\ P(y) > 0}} \sum_{x \in y} P(x, y) L(x, Q|_y) = \sum_{\substack{y \in \mathcal{Y}: \\ P(y) > 0}} P(y) H_L(P(\cdot \mid y)).$$

We will need the following properties of H_L throughout our theory:

Lemma 6.1. *For all loss functions L , if H_L is finite, then it is also concave and lower semi-continuous. If L is finite everywhere, then H_L is finite, concave, and continuous.*

(When we talk about (semi-)continuity, this is always with respect to the extended real line topology of losses, as in Rockafellar (1970, Section 7).)

Using just the concavity of the objective f_0 (which is a linear combination of concave generalized entropies), we can prove the following intuitive result.

Lemma 6.2 (Message subsumption). *Suppose that for $P \in \mathcal{P}$ there are two messages $y_1, y_2 \in \mathcal{Y}$ such that any outcome $x \in y_2$ with $P(x, y_2) > 0$ is also in y_1 . Then if H_L is concave, the quizmaster can do at least as well without using y_2 . More precisely, P' given by*

$$P'(x, y) = \begin{cases} P(x, y_1) + P(x, y_2) & \text{for } y = y_1; \\ 0 & \text{for } y = y_2; \\ P(x, y) & \text{otherwise.} \end{cases}$$

is also in $\mathcal{P}$ and its expected generalized entropy is at least as large as that of P . In particular, if P is worst-case optimal, then so is P' .

In particular, if $y_1 \supset y_2$, any strategy P can be replaced by a strategy P' with $P'(y_2) = 0$ without making things worse for the quizmaster. Thus the quizmaster, who wants to maximize the contestant's expected loss, never needs to use a message that is contained in another.

A *dominating hyperplane* to a function f from $D \subseteq \mathbf{R}^{\mathcal{X}}$ to $\mathbf{R}$ is a hyperplane in $\mathbf{R}^{\mathcal{X}} \times \mathbf{R}$ that is nowhere below f . A *supporting hyperplane* to f (at P) is a dominating hyperplane that touches f at some point P .² A concave function has at least one supporting hyperplane at every point (Rockafellar, 1970, Theorem 11.6), but it may be vertical. A nonvertical hyperplane can be described by a linear function $\ell: \mathbf{R}^{\mathcal{X}} \rightarrow \mathbf{R}$: $\ell(P) = \alpha + \sum_x P(x)\lambda_x$, where $\alpha \in \mathbf{R}$ and $\lambda \in \mathbf{R}^{\mathcal{X}}$.

While H_L is defined as a function of $\Delta_{\mathcal{X}}$, we will often need to talk about supporting hyperplanes to the function H_L restricted to Δ_y for some message $y \in \mathcal{Y}$. We use the notation $H_L \upharpoonright \Delta_y$ for the restriction of H_L to the domain Δ_y . (Recall that we defined Δ_y as a subset of $\Delta_{\mathcal{X}}$.) A supporting hyperplane to $H_L \upharpoonright \Delta_y$ is not a supporting hyperplane to H_L itself if it goes below H_L at some $P \in \Delta_{\mathcal{X}} \setminus \Delta_y$.

A *supergradient* is a generalization of the gradient: a supergradient of a concave function at a point is the gradient of a supporting hyperplane. If $H_L \upharpoonright \Delta_y$ is finite and continuous (and thus concave by Lemma 6.1), then for any vector $\lambda \in \mathbf{R}^{\mathcal{X}}$, a unique supporting hyperplane to $H_L \upharpoonright \Delta_y$ can be found having that vector as its gradient, by choosing α appropriately in $\ell(P) = \alpha + \sum_x P(x)\lambda_x$ (Rockafellar, 1970, Theorem 27.3). It will often be convenient in our discussion to talk about supporting hyperplanes rather than supergradients because they fix this choice of α .

Theorem 6.3 (Existence and characterization of P^*). *For H_L finite and upper semi-continuous (thus continuous), a worst-case optimal strategy for the quizmaster (that is, a $P \in \mathcal{P}$ maximizing (6.4)) exists, and P^* is such a strategy if and only if there exists a $\lambda^* \in \mathbf{R}^{\mathcal{X}}$ such that*

$$H_L(P') \leq \sum_{x \in y} P'(x)\lambda_x^* \quad \text{for all } y \in \mathcal{Y} \text{ and } P' \in \Delta_y,$$

²We deviate slightly from standard terminology here: what we call a supporting hyperplane to a concave function f is usually called a supporting hyperplane to $\{(u, v) \in \mathbf{R}^{\mathcal{X}} \times \mathbf{R} \mid v \leq f(u)\}$, the hypograph of f .

with equality if $P^*(y) > 0$ and $P' = P^*(\cdot | y)$. That is, for y with $P^*(y) > 0$, the linear function $\sum_{x \in \mathcal{Y}} P(x) \lambda_x^*$ defines a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at $P^*(\cdot | y)$, and a dominating hyperplane for other y .

A vector $\lambda^* \in \mathbf{R}^{\mathcal{X}}$ that satisfies the above for some worst-case optimal P^* satisfies it for all worst-case optimal P^* and is called a Kuhn-Tucker vector (or KT-vector).

Several examples illustrating the application of Theorem 6.3 are given in Section 6.3.1; a graphical illustration of the theorem is also included there (Figure 6.4). We will see in Section 6.4 that KT-vectors form the bridge between worst-case optimal strategies for the quizmaster and for the contestant.

Note that at any P with $P(y) = 0$ for some $y \in \mathcal{Y}$, the objective function f_0 defined in (6.8) is not differentiable for most L , so there may be multiple supporting hyperplanes at P . (This is why to formulate the preceding theorem we needed the theory of Rockafellar (1970), where no differentiability is assumed.)

6.3.1 Application to standard loss functions

The generalized entropy for logarithmic loss has only vertical supporting hyperplanes at the boundary of Δ_y for any $y \in \mathcal{Y}$. These hyperplanes do not correspond to any KT-vector $\lambda^* \in \mathbf{R}^{\mathcal{X}}$, from which it follows that for any y with $P^*(y) > 0$, the worst-case optimal strategy will not have $P^*(\cdot | y)$ at the boundary of Δ_y . The same is not generally true: we will see below how for randomized 0-1 loss (in Example 6.B on page 127, and Example 6.D) and Brier loss (in Example 6.E), games may have a worst-case optimal strategy for the quizmaster that has $P^*(y) > 0$, yet $P^*(x | y) = 0$ for some $x \in \mathcal{Y}$.

Of the three loss functions we saw earlier, Brier loss and 0-1 loss are finite, so by Lemma 6.1, all conditions of Theorem 6.3 are satisfied for them. Logarithmic loss is infinite when the obtained outcome was predicted to have probability zero. The generalized entropy is still finite, because for any true distribution P , there exist predictions Q that give finite expected loss (in particular, $Q = P$ does this). The entropy is also continuous: $-P(x) \log P(x)$ is continuous as a function of $P(x)$ with our convention that $0 \cdot \infty = 0$, and H_L is the sum of such continuous functions. Thus we can apply Theorem 6.3 to analyse the Monty Hall problem for each of these three loss functions.

Example 6.B (continued). For Monty Hall, the strategy P^* of choosing a message uniformly when the true outcome is x_2 is worst-case optimal for the quizmaster, for all three loss functions. It is easy to verify that the theorem is satisfied by this strategy combined with the appropriate KT-vector:

$$\text{for logarithmic loss: } \lambda^* = \left(-\log \frac{2}{3}, -\log \frac{1}{3}, -\log \frac{2}{3} \right);$$

$$\text{for Brier loss: } \lambda^* = \left(\frac{2}{9}, \frac{8}{9}, \frac{2}{9} \right);$$

$$\text{for randomized 0-1 loss: } \lambda^* = (0, 1, 0).$$

The situation for logarithmic loss is illustrated in Figure 6.4.

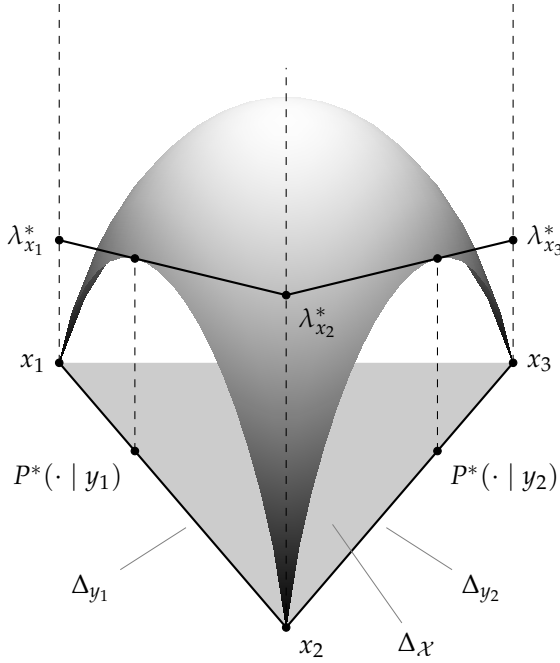


Figure 6.4: The worst-case optimal strategy for the quizmaster in the Monty Hall game with logarithmic loss, as characterized by Theorem 6.3. The triangular base is the full simplex $\Delta_{\mathcal{X}}$, on which the entropy function H_L is defined (this is the grey dome); the points labelled x_1 , x_2 and x_3 are the elements of this simplex putting all mass on that single outcome; and the line segments Δ_{y_1} and Δ_{y_2} are the subsets of $\Delta_{\mathcal{X}}$ consisting of all distributions on y_1 and y_2 respectively. Restricted to the domain Δ_{y_1} , the vector λ^* defines a linear function (having height λ_x^* at each $x \in \mathcal{X}$) that is a supporting hyperplane to H_L at $P^*(\cdot | y_1)$ (and similar for y_2). Note that when the linear function defined by λ^* is extended to all of $\Delta_{\mathcal{X}}$, it may go below H_L there, but not in Δ_{y_1} or Δ_{y_2} .

We also find that for logarithmic loss and Brier loss, P^* is the unique worst-case optimal strategy, as the hyperplanes specified by λ^* touch the generalized entropy functions at only one point each. For randomized 0-1 loss, on the other hand, all quizmaster strategies are worst-case optimal, as the hyperplane specified by λ^* touches $H_L \upharpoonright \Delta_{y_1}$ at all $P(\cdot \mid y_1)$ with $P(x_1 \mid y_1) \geq 1/2$.

Example 6.C (The quizmaster discards a message). Consider a different game, with $\mathcal{X} = \{x_1, x_2, x_3, x_4\}$, $\mathcal{Y} = \{\{x_1, x_2\}, \{x_2, x_3\}, \{x_3, x_4\}\}$, p given by $p_{x_4} = 2/5$ and $p_x = 1/5$ elsewhere, and L logarithmic loss. In the terminology of the Monty Hall puzzle, there is no initial choice by the contestant that determines what moves are available to the quizmaster, but the quizmaster will again leave two doors closed: the one hiding the car, and another adjacent to it. Then one strategy for the quizmaster is to never give message y_2 to the contestant; i.e. to pick the following strategy $P \in \mathcal{P}$ with $P(y_2) = 0$:

P	x_1	x_2	x_3	x_4
y_1	1/5	1/5	—	—
y_2	—	0	0	—
y_3	—	—	1/5	2/5
p_x	1/5	1/5	1/5	2/5

The depicted strategy P is worst-case optimal: When applying the theorem, we see that the KT-vector $\lambda^* = (\log 2, \log 2, \log 3, -\log(2/3))$ gives supporting hyperplanes to $H_L \upharpoonright \Delta_{y_1}$ and $H_L \upharpoonright \Delta_{y_3}$, but a non-supporting dominating hyperplane to $H_L \upharpoonright \Delta_{y_2}$. This strategy can be seen to be intuitively reasonable because when the contestant receives message $y_3 = \{x_3, x_4\}$, he knows that the probability of the true outcome being x_4 is at least twice as large as the probability of it being x_3 . By always giving message y_3 when the true outcome is x_3 , the quizmaster can keep this difference from becoming larger.

P is also the unique worst-case optimal strategy for Brier loss (as shown by the same analysis) and for randomized 0-1 loss (where the KT-vector is not unique: $(a, 1-a, 1, 0)$ for any $a \in [0, 1]$ is a KT-vector).

In the previous examples, the worst-case optimal strategies P coincided for logarithmic and Brier loss. The following example shows that this is not always the case.

Example 6.D (Dependence on loss function). Consider the family of games with $\mathcal{X} = \{x_1, x_2, x_3, x_4\}$, $\mathcal{Y} = \{\{x_1, x_2\}, \{x_2, x_3, x_4\}\}$, $p_{x_1} = p_{x_2} = 1/3$, and $p_{x_3} = p_{x_4} = 1/6$:

P	x_1	x_2	x_3	x_4
y_1	1/3	1/6	—	—
y_2	—	1/6	1/6	1/6
p_x	1/3	1/3	1/6	1/6

This game is also similar to Monty Hall, but now one door has been ‘split in two’: the quizmaster will either open door 1, or doors 3 and 4. The strategy P shown above is worst-case optimal for logarithmic loss, but not for Brier loss:

for both loss functions, there is a unique supporting hyperplane for both y that touches $H_L \upharpoonright \Delta_y$ at $P(\cdot \mid y)$ for P as shown in the table, but for Brier loss, these two hyperplanes do not have the same height at the common outcome x_2 . (Using Theorem 6.9 from page 136, we can find the worst-case optimal strategy for the quizmaster under Brier loss by solving a quadratic equation with one unknown; this strategy has $P(x_2, y_1) = 11/3 - 2\sqrt{3} \approx 0.20$ and $P(x_2, y_2) = 2\sqrt{3} - 10/3 \approx 0.13$.)

For randomized loss, neither the worst-case optimal strategy nor the KT-vector are unique: the KT-vectors are $(0, 1, a, 1 - a)$ for any $a \in [0, 1]$; the worst-case optimal strategies are the P given above, the strategy that always gives message y_1 when the true outcome is x_2 , and all convex combinations of these.

Example 6.E (The quizmaster discards a message-outcome pair). Again consider the game from the previous example, but now with a different marginal:

P	x_1	x_2	x_3	x_4
y_1	0.45	0.05	—	—
y_2	—	0	0.25	0.25
p_x	0.45	0.05	0.25	0.25

The strategy P is worst-case optimal for Brier loss, with KT-vector

$$\lambda^* = (0.02, 1.62, 0.5, 0.5).$$

P displays another curious property (that we also saw for randomized 0-1 loss in the previous example): while the quizmaster uses message y_2 for some outcomes, he does not use it in combination with outcome x_2 . In the theorem, the hyperplane on Δ_{y_2} is supporting at $P(\cdot \mid y_2)$, but is not a tangent plane: compared to the tangent plane, it has been ‘lifted up’ at the opposite vertex of the simplex $\Delta_{y_2}(x_2)$ to the same height as the supporting hyperplane on Δ_{y_1} .

This behaviour cannot occur in games with logarithmic loss: as we observed at the beginning of Section 6.3.1, if a worst-case optimal strategy P^* has $P^*(y) > 0$ for some $y \in \mathcal{Y}$, then it must assign positive probability to $P^*(x \mid y)$ for all $x \in y$.

6.4 Worst-case optimal strategies for the contestant

We now turn our attention to worst-case optimal strategies for the contestant. To this end, we look at the relation between the KT-vectors that appeared in Theorem 6.3 and the set of strategies $\mathcal{Q}$ the contestant can choose from.

6.4.1 Realizable hyperplanes

For any $y \in \mathcal{Y}$, Δ_y is defined in Section 6.2 as a $(|y| - 1)$ -dimensional subset of $\mathbf{R}_{\geq 0}^{\mathcal{X}}$. Thus a linear function $\ell : \Delta_y \rightarrow \mathbf{R}$ can be extended to a linear function $\bar{\ell}$ on the domain $\mathbf{R}_{\geq 0}^{\mathcal{X}}$ in different ways. Hence many different vectors $\lambda \in \mathbf{R}^{\mathcal{X}}$

representing supporting hyperplanes will correspond to what we can view as a single supergradient, because the hyperplanes agree on Δ_y . We can make the extension unique by requiring $\bar{\ell}$ to be zero at the origin and at the vertices of the simplex $\Delta_{\mathcal{X} \setminus y}$. Because such a normalized function $\bar{\ell} : \mathbf{R}_{\geq 0}^{\mathcal{X}} \rightarrow \mathbf{R}$ obeys $\bar{\ell}(0) = 0$, it can be written as $\bar{\ell}(P) = P^\top \lambda$ for some λ . These functions are thus uniquely identified by their gradients λ , allowing us to refer to them using ‘the (supporting) hyperplane λ' ’. Let Λ_y be the set of all gradients of such normalized functions that represent dominating hyperplanes to $H_L \upharpoonright \Delta_y$; in formula, let

$$\Lambda_y = \{\lambda \in \mathbf{R}^{\mathcal{X}} \mid \lambda_x = 0 \text{ for } x \notin y, \text{ and } \forall P \in \Delta_y : P^\top \lambda \geq H_L(P)\}.$$

For each nonvertical supporting hyperplane of $H_L \upharpoonright \Delta_y$, clearly the gradient is in Λ_y ; that is, all finite supergradients of this restricted function have a normalized representative in Λ_y . The set also includes vectors λ for which $P^\top \lambda > H_L(P)$ for all $P \in \Delta_y$, which do not correspond to supporting hyperplanes.

Not all vectors $\lambda \in \Lambda_y$ may be available to the contestant as responses to a play of $y \in \mathcal{Y}$ by the quizmaster. As a trivial example, consider logarithmic loss and a vector λ with $\sum_{x \in y} e^{-\lambda_x} < 1$ and $\lambda_x = 0$ for $x \notin y$. Then $\lambda \in \Lambda_y$ because the hyperplane defined by λ is dominating to $H_L \upharpoonright \Delta_y$ (thus the expected loss from λ is *larger* than what the contestant could achieve), but clearly there is no distribution $Q \in \Delta_{\mathcal{X}}$ that results in these losses on $x \in y$. We say that a vector $\lambda \in \Lambda_y$ is *realizable on y* if there exists a $Q \in \Delta_{\mathcal{X}}$ such that $L(x, Q) = \lambda_x$ for all $x \in y$, and then we say that such a Q *realizes* λ .

A partial order on vectors $\lambda, \lambda' \in \mathbf{R}^{\mathcal{X}}$ is given by: $\lambda \leq \lambda'$ if and only if $\lambda_x \leq \lambda'_x$ for all $x \in \mathcal{X}$. We write $\lambda < \lambda'$ when $\lambda \leq \lambda'$ and $\lambda \neq \lambda'$. For all $y \in \mathcal{Y}$, this partial order has the following property: For $\lambda, \lambda' \in \Lambda_y$, we have $\lambda \leq \lambda'$ if and only if for all $P \in \Delta_y$, $P^\top \lambda \leq P^\top \lambda'$. Therefore if $Q, Q' \in \Delta_{\mathcal{X}}$ realize $\lambda, \lambda' \in \Lambda_y$ respectively and $\lambda \leq \lambda'$, the contestant is never hurt by using Q instead of Q' as a prediction given the message y .

Any minimal element with respect to this partial order defines a supporting hyperplane to $H_L \upharpoonright \Delta_y$. For P in the relative interior of Δ_y , the converse also holds: all supporting hyperplanes at P are minimal elements. This is not the case for P at the relative boundary of Δ_y , where some supporting hyperplanes (the ones that ‘tip over’ the boundary) are not minimal.

Lemma 6.4. *If H_L is finite and continuous on Δ_y , then the following hold:*

1. *If $\lambda \in \Lambda_y$ is not a supporting hyperplane to $H_L \upharpoonright \Delta_y$, then there exists a supporting hyperplane $\lambda' \in \Lambda_y$ with $\lambda' < \lambda$;*
2. *If $\lambda \in \Lambda_y$ is a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at P but is not minimal in Λ_y , then there exists a minimal $\lambda' < \lambda$ in Λ_y ;*
3. *If $\lambda \in \Lambda_y$ is a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at P , then any $\lambda' \leq \lambda$ in Λ_y is a supporting hyperplane at P and obeys $\lambda'_x = \lambda_x$ for all $x \in y$ with $P(x) > 0$.*

Thus the contestant never needs to play a $Q_{|y}$ realizing a non-minimal element of Λ_y .

6.4.2 Existence

With the help of Lemma 6.4, we can formulate sufficient conditions for the existence of a worst-case optimal strategy Q^* for the contestant that, together with P^* for the quizmaster, forms a Nash equilibrium.

Theorem 6.5 (Existence of Q^*). *Suppose that the conditions of Theorem 6.3 hold, and that for all $y \in \mathcal{Y}$, all minimal supporting hyperplanes $\lambda \in \Lambda_y$ to $H_L \upharpoonright \Delta_y$ are realizable on y . Then there exists a worst-case optimal strategy $Q^* \in \mathcal{Q}$ for the contestant that achieves the same expected loss in the minimax game as P^* achieves in the maximin game: (P^*, Q^*) is a Nash equilibrium.*

We will see in Theorem 6.9 that a (or rather, at least one) Nash equilibrium exists for logarithmic loss and Brier loss. The existence of a Nash equilibrium in games with randomized 0-1 loss is shown by the following consequence of Theorem 6.5.

Proposition 6.6. *In games with randomized 0-1 loss, a Nash equilibrium exists.*

The following example shows what may go wrong if some supporting hyperplanes are not realizable.

Example 6.F (Hard 0-1 loss).

P^*	x_1	x_2	x_3
y_1	1/6	1/6	—
y_2	—	1/6	1/6
y_3	1/6	—	1/6
p_x	1/3	1/3	1/3

Consider the game with $\mathcal{X}$, $\mathcal{Y}$ and p as shown in the table, and with hard 0-1 loss (so that the contestant is not allowed to randomize):

$$L(x, Q) = \begin{cases} 0 & \text{if } Q(x) = 1; \\ 1 & \text{otherwise.} \end{cases}$$

This loss function has the same entropy function as randomized 0-1 loss, so the two loss functions are the same from the quizmaster's perspective. The table shows the unique worst-case optimal strategy for the quizmaster, with KT-vector $\lambda^* = (1/2, 1/2, 1/2)$ and expected loss $1/2$. For randomized 0-1 loss, the (as we will see below: unique) worst-case optimal strategy for the contestant would be to respond to any message y with the uniform distribution on y . However, for all $y \in \mathcal{Y}$, λ given by $\lambda_x = \mathbf{1}_{x \in y} \lambda_x^*$ is not realizable on y under hard 0-1 loss, so Theorem 6.5 does not apply. In fact, for any strategy Q the contestant might use, there exists a strategy P for the quizmaster that gives

expected loss $2/3$ or larger (because for at least two outcomes x , there must be a $y \ni x$ such that $L(x, Q|_y) = 1$). Thus the inequality (6.6) is strict: there is no Nash equilibrium, and a worst-case optimal strategy for either player is optimal only in the minimax/maximin sense.

This example also shows that the condition on existence of supporting hyperplanes in Theorem 6.5 cannot be replaced by the weaker condition that the infimum appearing in the definition (6.3) of H_L is always attained.

Games without Nash equilibria We will now briefly go into the situation seen in the preceding example, where Theorem 6.5 does not apply.

While for some games with hard 0-1 loss, no Nash equilibrium may exist, worst-case optimal strategies for the contestant do exist, and can be characterized using stable sets of a graph. A *stable set* is a set of nodes no two of which are adjacent (Schrijver, 2003b, Chapter 64). Consider the graph with node set $\mathcal{X}$ and with an edge between two nodes if and only if they occur together in some message. A set $S \subseteq \mathcal{X}$ is stable in this graph if and only if there exists a strategy $Q \in \mathcal{Q}$ for the contestant such that for all $x \in S$, $\max_{y: x \in y \in \mathcal{Y}} L(x, Q|_y) = 0$, and equal to 1 otherwise. The worst-case loss obtained by this strategy is $1 - \sum_{x \in S} p_x$. Thus finding the worst-case optimal strategy Q for the contestant is equivalent to finding a stable set S with maximum weight. Algorithmically, this is an NP-hard problem in general, though polynomial-time algorithms exist for certain classes of graphs, including perfect (this includes bipartite) graphs and claw-free graphs (Schrijver, 2003b).

With the exception of two examples in Section 6.5.1 illustrating the limits of our theory, we will not look at games without Nash equilibria any more from now on.

6.4.3 Characterization and nonuniqueness

The concept of a KT-vector, which helped characterize worst-case optimal strategies for the quizmaster in Theorem 6.3, now returns for a similar role in the characterization of worst-case optimal strategies for the contestant.

Theorem 6.7 (Characterization of Q^*). *Under the conditions of Theorems 6.3 and 6.5, a strategy $Q^* \in \mathcal{Q}$ is worst-case optimal for the contestant if and only if the vector given by $\lambda_x := \max_{y \ni x} L(x, Q^*|_y)$ is a KT-vector.*

If the loss $L(x, Q^*|_y)$ equals λ_x for all $x \in y$, then the worst-case optimal strategy Q^* is an equalizer strategy (Ferguson, 1967): the expected loss of Q^* does not depend on the quizmaster's strategy. Not all games have an equalizer strategy as worst-case optimal strategy, as Example 6.H below shows.

When constructing a worst-case optimal strategy for the contestant, there are three points where different options are available, so that a worst-case optimal Q is in general not unique. The following examples demonstrate these three points.

Example 6.G (λ^* not unique).

	x_1	x_2	x_3	x_4
y_1	*	*	—	—
y_2	—	*	*	—
y_3	—	—	*	*
y_4	*	—	—	*
p_x	1/4	1/4	1/4	1/4

Consider the game with $\mathcal{X}$, $\mathcal{Y}$ and p as in the table above, and with randomized 0-1 loss. For the quizmaster, any P^* that is uniform given each y is worst-case optimal, and any $\lambda_a = (a, 1-a, a, 1-a)$ with $a \in [0, 1]$ is a KT-vector. To each λ_a corresponds a unique worst-case optimal Q^* , namely the strategy that puts conditional probability $1-a$ on outcome x_1 or x_3 (whichever is in the given message), and probability a on x_2 or x_4 .

Note that if we replace randomized 0-1 loss by a strictly proper loss function such as logarithmic or Brier loss, the KT-vector and the worst-case optimal strategy for the contestant become unique, while the same set of strategies as before continues to be worst-case optimal for the quizmaster. This shows that the freedom for the contestant we see here for randomized 0-1 loss is due to the nonuniqueness of the KT-vector, not due to the nonuniqueness of P^* .

Example 6.H (Minimal λ not unique).

P^*	x_1	x_2	x_3
y_1	1/5	3/10	—
y_2	—	3/10	1/5
y_3	0	—	0
p_x	1/5	3/5	1/5

Consider the game as shown in the table with logarithmic loss; the strategy P^* shown in this table is the unique worst-case optimal strategy for the quizmaster. Because logarithmic loss is proper, we know that $Q^*_{|y_1} = P^*(\cdot | y_1)$ and $Q^*_{|y_2} = P^*(\cdot | y_2)$ are optimal responses for the contestant, but this does not tell us what $Q^*_{|y_3}$ should be in a worst-case optimal strategy for the contestant.

We see that P^* assigns probability zero to message y_3 , and the KT-vector

$$\lambda^* = (-\log \frac{2}{5}, -\log \frac{3}{5}, -\log \frac{2}{5})$$

specifies a hyperplane that does not support H_L in Δ_{y_3} . Hence the construction of $Q^*_{|y_3}$ in the proof of Theorem 6.5 allows freedom in the choice of a minimal element $\lambda \in \Lambda_{y_3}$ less than $(-\log 2/5, 0, -\log 2/5)$: the valid choices are $(-\log q, 0, -\log(1-q))$ for any $q \in [2/5, 3/5]$; each of these is realized on y_3 by $Q_{|y_3} = (q, 0, 1-q)$. Using Theorem 6.7, we can verify that these choices of $Q^*_{|y_3}$ define worst-case optimal strategies: the vector λ defined there equals the KT-vector λ^* given above.

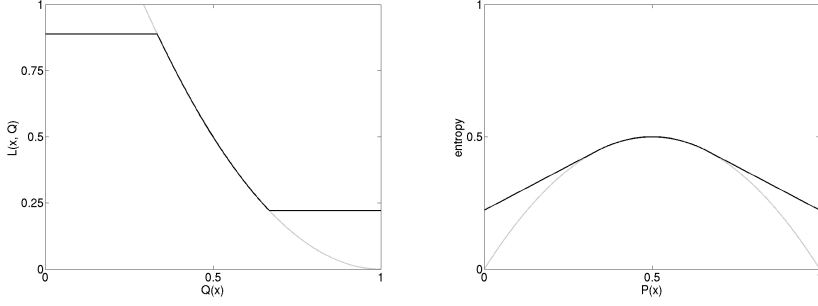


Figure 6.5: Loss and entropy on a binary prediction for the loss function in Example 6.I, with Brier loss and entropy shown in grey

This also shows that worst-case optimal strategies for the contestant cannot be characterized simply as ‘optimal responses to P^* ’: in this example, $P^*(\cdot | y_3)$ is undefined, yet there is a nontrivial constraint on $Q|_{y_3}$ in the worst-case optimal strategy Q for the contestant.

Example 6.I (Q realizing λ not unique). Take $\mathcal{X} = \{x_1, x_2, x_3\}$, $\mathcal{Y} = \{y_1 = \{x_1, x_2\}, y_2 = \{x_2, x_3\}\}$, and p uniform (as in Monty Hall), with loss function

$$L(x, Q) = \begin{cases} \frac{8}{9} & \text{for } Q(x) < \frac{1}{3}; \\ L_{\text{Brier}}(x, Q) & \text{for } Q(x) \in [\frac{1}{3}, \frac{2}{3}]; \\ \frac{2}{9} & \text{for } Q(x) > \frac{2}{3}; \end{cases}$$

(illustrated in Figure 6.5; L_{Brier} denotes the Brier loss function). This loss function is proper but not strictly proper: Because its generalized entropy is not strictly concave, a supporting hyperplane to $H_L \upharpoonright \Delta_{y_1}$ exists that supports H_L at all $P \in \Delta_{y_1}$ with $P(x_1) \leq 1/3$. Any Q in the same interval realizes this hyperplane and thus minimizes the expected loss.

6.5 Results for well-behaved loss functions

In the preceding sections, we have established characterization results for the worst-case optimal strategies of both players. While these results are applicable to many loss functions, they have the disadvantage of being complicated, involving supporting hyperplanes. For some common loss functions, simpler characterizations can be given.

6.5.1 Proper continuous loss functions

Recall from page 119 that for a *proper* loss function, the contestant’s expected loss for a given message is minimized if his predicted probabilities equal the

true probabilities. Such loss functions are natural to consider in our probability updating game, as our goal will often be to find these true probabilities. However, simplifying our theorems requires further restrictions on the class of loss functions. In this subsection, we consider loss functions that are both proper and continuous.

Lemma 6.8. *If the loss function $L(x, Q)$ is proper and continuous as a function of Q for all x and H_L is finite, then H_L is differentiable in the following sense: for all $y \in \mathcal{Y}$ and all $P \in \Delta_y$, there is at most one element of Δ_y that is a minimal supporting hyperplane to $H_L \upharpoonright \Delta_y$ at P ; if P is in the relative interior of Δ_y , there is exactly one. If it exists, the minimal supporting hyperplane at P is realized by $Q|_y = P$.*

The uniqueness of minimal supporting hyperplanes in Δ_y is equivalent to there being exactly one equivalence class of supergradients, where supergradients are taken to be equivalent if their corresponding supporting hyperplanes coincide on Δ_y . The property shown in the above lemma is then related to differentiability by Rockafellar (1970, Theorem 25.1), which says that for a finite, concave function such as H_L , uniqueness of the supergradient at P is equivalent to differentiability at P .

Theorem 6.9. *For L proper and continuous and H_L finite and continuous,*

1. *worst-case optimal strategies for both players exist and form a Nash equilibrium;*
2. *there is a unique KT-vector;*
3. *a strategy $P^* \in \mathcal{P}$ for the quizmaster is worst-case optimal if and only if there exists $\lambda^* \in \mathbf{R}^{\mathcal{X}}$ such that*

$$\begin{aligned} L(x, P^*(\cdot | y)) &= \lambda_x^* \quad \text{for all } x \in y \text{ with } P^*(x, y) > 0, \\ L(x, P^*(\cdot | y)) &\leq \lambda_x^* \quad \text{for all } x \in y \text{ with } P^*(x, y) = 0, P^*(y) > 0, \text{ and} \\ \exists Q|_y^* : L(x, Q|_y^*) &\leq \lambda_x^* \quad \text{for all } x \in y \text{ with } P^*(y) = 0; \end{aligned}$$

4. *a strategy Q^* for the contestant is worst-case optimal if and only if there exists a worst-case optimal P^* such that for all x ,*

$$\max_{y \ni x} L(x, Q|_y^*) = \max_{\substack{y \ni x, \\ P^*(y) > 0}} L(x, P^*(\cdot | y)), \quad (6.9)$$

which holds if and only if (6.9) holds for all worst-case optimal P^ .*

Using this theorem, many observations made about logarithmic loss and Brier loss in the examples we have seen so far can now be more easily verified. For instance, in the worst-case optimal strategy we saw in Example 6.E on page 130, we verify that $L(x_2, P^*(\cdot | y_2)) = 1.5 \leq 1.62 = \lambda_{x_2}^* = L(x_2, P^*(\cdot | y_1))$.

The preceding theory requires that L is both proper and continuous. If one of these conditions is removed, H_L may not be differentiable and the conclusions of Theorem 6.9 may fail to hold. This is illustrated by the following two examples, in which no Nash equilibria exist.

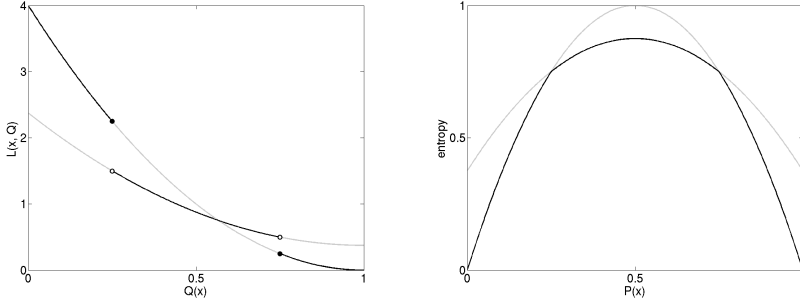


Figure 6.6: Loss and entropy on Δ_{y_1} for the loss function in Example 6.K. The functions are the same on Δ_{y_2} except for the value of L at the discontinuities. The grey lines show the two cases in the definition of L .

Example 6.J (Continuous but improper). The following loss function is continuous, but not proper:

$$L(x, Q) = 1 - Q(x)^2.$$

This loss function is actually hard 0-1 loss (see Example 6.F) in disguise: both have the same entropy function as randomized 0-1 loss (which is not differentiable), and for both, the hyperplane realized by $Q|_y$ will touch H_L on Δ_y only if $Q|_y$ puts mass 1 on some $x \in y$, so that only a small number of supporting hyperplanes is realizable. Example 6.F shows how for these loss functions, a Nash equilibrium may fail to exist.

(Another example of a continuous but improper loss function is randomized 0-1 loss; while a Nash equilibrium does exist there, parts 2, 3 and 4 of Theorem 6.9 generally do not hold.)

Example 6.K (Strictly proper but not continuous). Take $\mathcal{X} = \{x_1, x_2, x_3\}$ and $\mathcal{Y} = \{y_1 = \{x_1, x_2\}, y_2 = \{x_2, x_3\}\}$ as in Monty Hall, with loss function (we write L_{Brier} for the Brier loss function and H_{Brier} for its generalized entropy)

$$L(x, Q) = \begin{cases} L_{\text{Brier}}(x, Q) + \frac{3}{8} & \text{if } Q(x) \in (1/4, 3/4), \text{ or} \\ & Q(x) \in [1/4, 3/4] \text{ and } Q(x_1) = 0; \\ 2L_{\text{Brier}}(x, Q) & \text{otherwise} \end{cases}$$

(illustrated in Figure 6.6). For each $P \in \Delta_y$ (for both $y \in \mathcal{Y}$), $Q = P$ realizes a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at P , so that L is proper. For both $y \in \mathcal{Y}$, the entropy H_L as a function of $P \in \Delta_y$ is the pointwise minimum of $2H_{\text{Brier}}$ and $H_{\text{Brier}} + 3/8$ (shown in grey in Figure 6.6). These two entropies coincide when $P(x_2) \in \{1/4, 3/4\}$, where $H_L \upharpoonright \Delta_y$ is not differentiable but has multiple supporting hyperplanes. Only one of these hyperplanes is realizable. This means that L is strictly proper, but also that Theorem 6.5 does not apply.

Now consider the game with marginal $p = (1/8, 3/4, 1/8)$. The worst-case optimal strategy for the quizmaster has $P^*(x_2 | y_1) = P^*(x_2 | y_2) = 3/4$. Because L is strictly proper, the unique optimal responses $Q|_{y_1}$ and $Q|_{y_2}$ for the

contestant must coincide with these conditional distributions for (P^*, Q) to be a Nash equilibrium. However, $L(x_2, Q_{|y_1}) < L(x_2, Q_{|y_2})$, so that the quizmaster can increase the expected loss by choosing to give message y_2 more often when the true outcome is x_2 . This shows that Q is not a worst-case optimal strategy for the contestant. The contestant can do better in the minimax game by choosing $Q_{|y_2}$ with $Q_{|y_2}(x_2) = 3/4 + \epsilon > 3/4$, so that $L(x_2, Q_{|y_2}) < L(x_2, Q_{|y_1})$. The expected loss of this strategy can be made arbitrarily close, but not equal to, the worst-case expected loss achieved by the quizmaster in the maximin game.

While uniqueness of λ^* was established by the theorem, we do not have uniqueness of P^* or Q^* . Multiple worst-case optimal strategies Q^* for the contestant may exist as soon as a message is unused, as in Example 6.H. Multiple worst-case optimal strategies P^* for the quizmaster are also possible. For instance, the loss function in Example 6.I is proper and continuous (though not strictly proper) and thus H_L is differentiable. But if the marginal is $p = (2/5, 1/5, 2/5)$, all strategies for the quizmaster are worst-case optimal.

6.5.2 Local loss functions

Logarithmic loss is an example of a *local* loss function: a loss function where the loss $L(x, Q)$ depends on the probability assigned by the prediction Q to the obtained outcome x , but not on the probabilities assigned to outcomes that did not occur. The following theorem shows how for such loss functions, worst-case optimality of the quizmaster's strategy can be characterized purely in terms of probabilities, without converting them to losses.

Among loss functions that are 'smooth' for all x , logarithmic loss is, up to some transformations, the only proper local loss function (Bernardo, 1979). We do not know what non-smooth local proper loss functions may exist. In particular, it is conceivable (yet unlikely) that a discontinuous L exists satisfying the conditions of Theorem 6.10, but not those of Theorem 6.9.

Theorem 6.10 (Characterization of P^* for local L). *For L local and proper and H_L finite and continuous, $P^* \in \mathcal{P}$ is worst-case optimal if there exists a vector $q \in [0, 1]^{\mathcal{X}}$ such that*

$$\begin{aligned} q_x &= P^*(x | y) \text{ for all } x \in y \in \mathcal{Y} \text{ with } P^*(y > 0), \text{ and} \\ \sum_{x \in y} q_x &\leq 1 \text{ for all } y \in \mathcal{Y}. \end{aligned} \tag{6.10}$$

If additionally $H_L \upharpoonright \Delta_y$ is strictly concave for all $y \in \mathcal{Y}$, only such P^ are worst-case optimal for L .*

If additionally L is continuous, then Theorem 6.9 applies, and it follows that $Q^* \in \mathcal{Q}$ is a worst-case optimal strategy for the contestant if $Q_{|y}^*(x) \geq q_x$ for all $x \in y \in \mathcal{Y}$. For strictly proper loss functions such as logarithmic loss, this fully characterizes the worst-case optimal strategies for the contestant.

Example 6.H (continued). Consider again the following game:

P^*	x_1	x_2	x_3
y_1	1/5	3/10	—
y_2	—	3/10	1/5
y_3	0	—	0
p_x	1/5	3/5	1/5

with logarithmic loss. The conditionals $P^*(x | y)$ agree with the vector $q = (2/5, 3/5, 2/5)$. For all $y \in \mathcal{Y}$ with $P^*(y) > 0$, this implies that $\sum_{x \in \mathcal{X}} q_x = 1$; for y_3 , we see that this sum equals $4/5 \leq 1$. Thus P^* is verified to be worst-case optimal.

The equality of conditionals $P^*(x | y)$ with the same x in the statement of Theorem 6.10 is oddly similar to the CAR condition we saw in Section 6.1, but reversing the roles of outcomes and messages. We may say that a strategy P^* satisfying (6.10) is *RCAR* (sometimes *with vector q*), for ‘reverse CAR’. Note that whether a strategy is RCAR does not depend on the loss function.

A vector q is called an *RCAR vector* if a strategy $P^* \in \mathcal{P}$ exists such that P^* and q satisfy (6.10). This definition is also independent of the loss function. If q is an RCAR vector, then $q_x > 0$ for all $x \in \mathcal{X}$; otherwise we would get $P^*(x) = 0 < p_x$. Like the KT-vector λ^* in Theorem 6.9, the RCAR vector is unique:

Lemma 6.11. *Given $\mathcal{X}, \mathcal{Y}, p$, there exists a unique RCAR vector $q \in [0, 1]^{\mathcal{X}}$.*

If each message in $\mathcal{Y}$ contains an outcome x not contained in any other message, then any strategy P^* must have $P^*(y) > 0$ for all $y \in \mathcal{Y}$. Then the first line of (6.10) implies that $\sum_{x \in \mathcal{X}} q_x = 1$ for all y . Thus the second line is now satisfied automatically, allowing the theorem to be simplified for this case:

Corollary 6.12. *A strategy $P^* \in \mathcal{P}$ with $P^*(y) > 0$ for all $y \in \mathcal{Y}$ that satisfies*

$$P^*(x | y) = P^*(x | y') \text{ for all } y, y' \ni x \quad (6.11)$$

is worst-case optimal for the loss functions covered by Theorem 6.10

In this case, P^* is an equalizer strategy (Ferguson, 1967).

The symmetry between versions of CAR and RCAR is clearest in Corollary 6.12: the condition (6.11) is the mirror image of the definition of *strong CAR* in Jaeger (2005a). Thus we may call it *strong RCAR*. Ordinary RCAR (6.10) imposes an inequality on q for messages with probability 0, which has no analogue in the CAR literature that we know of: the definition of *weak CAR* in Jaeger puts no requirement at all on outcomes with probability 0.

Strict concavity of H_L occurred as a new condition in Theorem 6.10. The main loss function of interest here is logarithmic loss, and its entropy is strictly concave. For other loss functions, the following lemma relates strict concavity of H_L to conditions we have seen before.

Lemma 6.13. *If L is strictly proper and all minimal supporting hyperplanes $\lambda \in \Lambda_y$ to $H_L \upharpoonright \Delta_y$ are realizable on y for all $y \in \mathcal{Y}$, then H_L is strictly concave on $H_L \upharpoonright \Delta_y$ for all $y \in \mathcal{Y}$.*

Affine transformations of the loss function Above, we mentioned that logarithmic loss is the only local proper loss function up to some transformations. The transformations considered in Bernardo (1979) are affine transformations, of the form

$$L'(x, Q) = aL(x, Q) + b_x \quad (6.12)$$

for $a \in \mathbf{R}_{>0}$ and $b \in \mathbf{R}^{\mathcal{X}}$. (This transformation can result in a function L' that can take negative values, so that it does not satisfy our definition of a loss function. However, our results can easily be extended to loss functions bounded from below by an arbitrary real number, so we allow such transformations here.)

The following lemma shows that, for logarithmic loss as well as for other loss functions, the transformation (6.12) does not change how the players of the probability updating game should act.

Lemma 6.14. *Let L be a loss function for which H_L is finite and continuous, and let L' be an affine transformation of L as in (6.12). Then a strategy P^* is worst-case optimal for the quizmaster in the game $\mathcal{G}' := (\mathcal{X}, \mathcal{Y}, p, L')$ if and only if P^* is worst-case optimal in $\mathcal{G} := (\mathcal{X}, \mathcal{Y}, p, L)$. If $\mathcal{G}$ also satisfies the conditions of Theorem 6.5, then the same equivalence holds for worst-case optimal strategies Q^* for the contestant.*

Lemma 6.14 has highly important implications when applied to the logarithmic loss. While multiplying logarithmic loss by a constant $a \neq 1$ merely corresponds to changing the base of the logarithm, adding constants b_x allows the logarithmic loss to become the appropriate loss function for a very wide class of games. This means that the RCAR characterization of worst-case optimal strategies for logarithmic loss is also valid for all these games. We are referring to so-called *Kelly gambling games*, also known as *horse race games* (Cover and Thomas, 1991) in the literature. In such games (with terminology adapted to our setting), for any outcome x the contestant can buy a ticket which costs $\in 1$ and which pays off a positive amount $\in c_x$ if x actually obtains; if some $x' \neq x$ is realized, nothing is paid so the $\in 1$ is lost. The contestant is allowed to distribute his capital over several tickets (outcomes), and he is also allowed to buy a fractional nonnegative number of tickets. For example, if $\mathcal{X} = \{1, 2\}$ and $c_1 = c_2 = 2$, then the contestant is guaranteed to neither win nor lose any money if he splits his capital fifty-fifty over both outcomes.

Now consider a contestant with some initial capital (say, $\in 1$), who faces an i.i.d. sequence $(X_1, Y_1), (X_2, Y_2), \dots \sim P$ of outcomes in $\mathcal{X} \times \mathcal{Y}$. At each point in time i he observes ‘side information’ $Y_i = y_i$ and he distributes his capital gained so far over all $x \in \mathcal{X}$, putting some fraction $Q_{|y_i}(x)$ of his capital on outcome x . Then he is paid out according to the x_i that was actually realized. Here each $Q_{|y}$ is a probability distribution over $\mathcal{X}$, i.e. for all $y \in \mathcal{Y}$, all $x \in \mathcal{X}$, $Q_{|y}(x) \geq 0$ and $\sum_{x \in \mathcal{X}} Q_{|y}(x) = 1$. So if his capital was U_i before the i -th round, it will be $U_i \cdot Q_{|y_i}(x_i) c_{x_i}$ after the i -th round. By the law of large numbers, his capital will grow (or shrink, depending on the odds on offer) almost surely exponentially fast, with exponent $\mathbf{E}_{X, Y \sim P}[\log Q_{|Y}(X) c_X] = \mathbf{E}_{X, Y \sim P}[\log Q_{|Y}(X) - b_X]$, where $b_x = -\log c_x$ (Cover and Thomas, 1991, Chapter 6). Thus, the contestant’s capital will grow fastest, among all constant strategies and against

an adversarial distribution $P \in \mathcal{P}$, if he plays a worst-case optimal strategy for gains $\log Q(x) - b_x$, i.e. for loss function $L'(x, Q) = -\log Q(x) + b_x$. By Lemma 6.14 above, this worst-case optimal strategy Q^* is just the Q^* that is also worst-case optimal for logarithmic loss — *it does not depend on the pay-offs* ('odds' in the horse race interpretation) c_x . Clearly, if data are i.i.d. then this continues to hold even if the pay-offs are allowed to change over time, and even if the contestant is allowed to use different strategies at different time points: the worst-case optimal capital growth rate is always achieved by choosing Q^* at all time points.

The upshot is that whenever the probability updating game is played (a) repeatedly, and (b) the contestant is allowed to reinvest and redistribute his capital over all outcomes at each point in time, then his worst-case optimal strategy is equal to the worst-case optimal Q^* for logarithmic loss *irrespective of the pay-offs*. This makes the logarithmic loss, and hence the RCAR characterization, appropriate for a very wide class of settings.

6.6 Conclusion

We have seen many theorems in the last few subsections showing different properties of worst-case optimal strategies P^* and Q^* for the two players for different classes of loss functions. A summary of these theorems was given in Table 6.1 on page 125.

Worst-case optimal probability updating provides a robust new approach for dealing with underspecified distributions. There are many scenarios in which our results currently do not apply, but to which they might be extended. For example, the quizmaster's hard constraint $y \ni x$ could be replaced by some soft constraint, so that each message y still carries information about the true outcome, but no longer in the form of a subset of $\mathcal{X}$. One way to achieve this might be by affine transformations of the loss function as discussed at the end of Section 6.5.2, but allowing the constants to depend on both x and y . This could give a worst-case analogue to *Jeffrey conditioning* or *minimum relative entropy updating* (Grünwald and Halpern, 2003).

Another extension would be to infinite outcome and message spaces.

An online version of the probability updating game may also be considered (Cesa-Bianchi and Lugosi, 2006).

Other questions concern the comparison between different alternative approaches the contestant might use to update his probabilities. For example, can we bound the difference in expected loss between worst-case optimal and naive conditioning? What about ignoring the message and always predicting with the marginal, or ignoring the constraints imposed on the quizmaster by the marginal and predicting with the maximum entropy distribution on y ? (Both these strategies are overly pessimistic.) Conversely, we might wonder how much the contestant loses by playing a worst-case optimal strategy when the quizmaster is not adversarial, but for instance chooses from the available messages uniformly at random.

Appendix 6.A Proofs

Proof of Lemma 6.1. For finite H_L , concavity of H_L is shown by Grünwald and Dawid (2004, Proposition 3.2), and lower semi-continuity by Rockafellar (1970, Theorem 10.2) (using that the domain of H_L is a simplex). If L is finite, then picking any $Q \in \Delta_{\mathcal{X}}$ gives an upper bound to H_L , so that H_L is in particular finite. Concavity now follows by the first claim, and continuity by Grünwald and Dawid (2004, Corollary 3.3; an important condition is in Corollary 3.2). $\square$

Proof of Lemma 6.2. If $P(y_2) = 0$ then $P' = P$ and the result is trivial; if $P(y_2) > 0$ but $P(y_1) = 0$, then $P(\cdot | y_2) = P'(\cdot | y_1)$ so P and P' have the same expected generalized entropy. Otherwise $P(\cdot | y_1)$ and $P(\cdot | y_2)$ are well-defined, and $P'(\cdot | y_1) = (P(y_1)P(\cdot | y_1) + P(y_2)P(\cdot | y_2)) / (P(y_1) + P(y_2))$ is a convex combination of them. By concavity of H_L , $\sum_y P'(y)H_L(P'(\cdot | y)) \geq \sum_y P(y)H_L(P(\cdot | y))$. $\square$

Proof of Theorem 6.3. Rockafellar (1970, Theorem 27.3) provides conditions under which a convex minimization problem has a solution attaining the minimum. These are satisfied by $\mathcal{P}$ and $-H_L$: $\mathcal{P}$ is nonempty, closed, convex, and bounded (thus has no direction of recession), and $-H_L$ is convex (Lemma 6.1), finite for all $P \in \mathcal{P}$ (thus proper), and lower semi-continuous (thus closed).

By Rockafellar (1970, Corollary 28.2.2), a KT-vector λ^* exists, so that for the remaining claims of the theorem, it suffices to show that P^* is worst-case optimal and λ^* is a KT-vector if and only if the given conditions on (P^*, λ^*) hold.

We rewrite the maximin problem to

$$\begin{aligned} & \text{maximize} && f_0(P) \\ & \text{subject to} && \sum_{y \ni x} P(x, y) = p_x \quad \text{for all } x \in \mathcal{X}, \end{aligned}$$

with $P \in \mathbf{R}_{\geq 0}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$. By Rockafellar (1970, Theorem 28.3), $P^* \in \mathbf{R}_{\geq 0}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$ maximizes this and $\lambda^* \in \mathbf{R}^{\mathcal{X}}$ is a KT-vector if and only if $P^* \in \mathcal{P}$ and at P^* , the zero vector is a supergradient to

$$f_0(P^*) - \sum_{x \in \mathcal{X}} \lambda_x^* \left(\sum_{y \ni x} P^*(x, y) - p_x \right). \quad (6.13)$$

The term being subtracted is linear, with gradient $\bar{\lambda} \in \mathbf{R}^{\mathcal{R}(\mathcal{X}, \mathcal{Y})}$ given by

$$\bar{\lambda}_{x,y} := \frac{\partial}{\partial P^*(x, y)} \sum_{x \in \mathcal{X}} \lambda_x^* \left(\sum_{y \ni x} P^*(x, y) - p_x \right) = \lambda_x^*. \quad (6.14)$$

By Rockafellar (1970, Theorem 23.8), 0 is a supergradient to (6.13) if and only if $\bar{\lambda}$ is a supergradient to f_0 at P^* .

For any P^* that is not everywhere zero, we have for all $c \geq 0$ that $f_0(cP^*) = cf_0(P^*)$, so that a supporting hyperplane to f_0 at a feasible P^* must go through

the origin. Then the supporting hyperplane with gradient $\bar{\lambda}$ has as defining equation the linear expression $\sum_{x,y} P(x,y)\bar{\lambda}_{x,y}$.

If $\sum_{x,y} P(x,y)\bar{\lambda}_{x,y}$ defines a supporting hyperplane to f_0 at P^* , then

1. at every $y \in \mathcal{Y}$ with $P^*(y) > 0$, it is a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at $P^*(\cdot \mid y)$, and
2. for every y with $P^*(y) = 0$, $H_L(P') \leq \sum_x P'(x)\bar{\lambda}_{x,y}$ for all $P' \in \Delta_y$.

The converse also holds: we have for all $y \in \mathcal{Y}$ and $P' \in \Delta_y$ that $H_L(P') \leq \sum_{x \in y} P'\bar{\lambda}_{x,y}$, with equality if $P^*(y) > 0$ and $P' = P^*(\cdot \mid y)$; taking the convex combination with coefficients $P^*(y)$ shows that the hyperplane defined by $\sum_{x,y} P(x,y)\bar{\lambda}_{x,y}$ is nowhere below f_0 and touches it at $P = P^*$.

For $\bar{\lambda}$ of the required form (6.14), this is in turn equivalent to the characterization given in the statement of the theorem. $\square$

Proof of Lemma 6.4. The function $\sum_{x \in y} \lambda_x P(x) - H_L(P)$ attains its minimum d on Δ_y at some P (Rockafellar, 1970, Theorem 27.3). Let $\lambda' \in \Lambda_y$ be given by $\lambda'_x = \lambda_x - d$ for all $x \in y$: this defines a hyperplane to $H_L \upharpoonright \Delta_y$ that is supporting at the minimizing P , proving the first part of the lemma.

For the third part, let λ , λ' and P be as described. $\lambda' \leq \lambda$ implies $P^\top \lambda' \leq P^\top \lambda$. Neither can be smaller than $H_L(P)$, and since the right hand side must equal $H_L(P)$ because λ is supporting at P , so must the left hand side, showing that λ' is also supporting at P . If $P(x) = 1$ for some $x \in y$, then $P^\top \lambda'' = \lambda''_x$ for any λ'' , so in particular $\lambda'_x = \lambda_x = H_L(P)$. For other $P \in \Delta_y$, we use that two linear functions obeying an inequality on their domain $D := \Delta_{\{x \in y \mid P(x) > 0\}}$ and coinciding at a point in the relative interior of D must coincide everywhere on D , so that again $\lambda'_x = \lambda_x$ for $x \in y$ with $P(x) > 0$.

It remains to show that given such a λ , a minimal $\lambda' \leq \lambda$ exists in Λ_y . Consider the set $\Lambda' := \{\lambda' \in \Lambda_y \mid \lambda' \leq \lambda\}$. The set of supporting hyperplanes to $H_L \upharpoonright \Delta_y$ at P in Λ_y is closed (Rockafellar, 1970, Section 23, definition subdifferential (page 215)); Λ' is a subset of this set (as we just saw), obtained by adding further non-strict linear constraints, so it too is closed. It also has the property that if $\lambda' \in \Lambda'$ is minimal in that set, it is also minimal in Λ_y . Now fix any P' in the relative interior of Δ_y , and pick some $\lambda' \in \Lambda'$ that minimizes $P'^\top \lambda'$ (this minimum must be attained because the expression is bounded below and Λ' is closed). Such a λ' is also minimal in the partial order, so it is the element we are looking for. $\square$

Proof of Theorem 6.5. Take a worst-case optimal strategy P^* for the quizmaster and KT-vector λ^* . For each $y \in \mathcal{Y}$, define a vector

$$\lambda'_x = \begin{cases} \lambda^*_x & \text{for } x \in y \\ 0 & \text{for } x \notin y. \end{cases}$$

By the statement of Theorem 6.3, $\lambda' \in \Lambda_y$. Let λ be a minimal element of Λ_y with $\lambda \leq \lambda'$: such an element exists by parts 1 and 2 of Lemma 6.4. (If λ' is

itself minimal, $\lambda = \lambda'$). By assumption, λ is realizable on y . Let $Q_{|y}^*$ be given by this Q .

By playing this Q^* , the contestant will achieve expected loss (against any strategy $P \in \mathcal{P}$ for the quizmaster, for λ^* any KT-vector)

$$\sum_{x,y} P(x,y) L(x, Q_{|y}^*) \leq \sum_{x,y} P(x,y) \lambda_x^* = \sum_x p_x \lambda_x^*.$$

The right-hand side is the maximum loss the quizmaster can achieve in the maximin game. By (6.6), the reverse inequality also holds, so we find that the values of the minimax and maximin games must be equal. $\square$

Proof of Proposition 6.6. We first introduce some additional terminology in order to apply a corollary from Rockafellar (1970).

A nonvertical hyperplane defined by $\lambda \in \mathbf{R}^{\mathcal{X}}$ is geometrically a subset of $\mathbf{R}^{\mathcal{X}} \times \mathbf{R}$, namely $\{(P', z') \in \mathbf{R}^{\mathcal{X}} \times \mathbf{R} \mid z' = P'^{\top} \lambda\}$. This set is the boundary of the half-space $H_{\lambda} = \{(P', z') \in \mathbf{R}^{\mathcal{X}} \times \mathbf{R} \mid z' \leq P'^{\top} \lambda\}$. A hyperplane λ is supporting to a concave function $f : \mathbf{R}^{\mathcal{X}} \rightarrow \mathbf{R}$ at the point $P \in \mathbf{R}^{\mathcal{X}}$ with $f(P) = P^{\top} \lambda$ if and only if the hypograph of f is a subset of H_{λ} .

A column vector $(-\alpha \lambda, \alpha)$ is called *normal* to a convex set C at a point $(P, z) \in C$ if $(P' - P, z' - z)^{\top} (-\alpha \lambda, \alpha) \leq 0$ for all $(P', z') \in C$ (Rockafellar, 1970); that is, if $C \subseteq \{(P', z') \mid (P' - P, z' - z)^{\top} (-\alpha \lambda, \alpha) \leq 0\}$. This latter set is equal to H_{λ} if $\alpha > 0$ and $z = P^{\top} \lambda$. So if C is the hypograph of f and $f(P) = z = P^{\top} \lambda$, then λ is a supporting hyperplane to f at P if and only if $(-\lambda, 1)$ is normal to C .

The set of all vectors normal to C at (P, z) is called the *normal cone* at (P, z) . The normal cone to H_{λ} at given $(P, P^{\top} \lambda)$ is the half-line $\{(-\alpha \lambda, \alpha) \mid \alpha \in [0, \infty)\}$.

For L randomized 0-1 loss, let the function $f_0 : \mathbf{R}^{\mathcal{X}} \rightarrow \mathbf{R}$ be given by $f_0(P) = \min_{Q \in \Delta_{\mathcal{X}}} \sum_{x' \in \mathcal{X}} P(x') L(x', Q)$; note that $f_0 \upharpoonright \Delta_{\mathcal{X}} = H_L$, and that for all $y \in \mathcal{Y}$, any minimal supporting hyperplane $\lambda \in \Lambda_y$ to $H_L \upharpoonright \Delta_y$ can be extended to a supporting hyperplane λ' to f_0 with $\lambda'_x = \lambda_x$ for all $x \in y$.

The hypograph of f_0 is $C = \bigcap_{x \in \mathcal{X}} H_{\lambda^{(x)}}$, with $\lambda_{x'}^{(x)} = L(x', e_x)$ (where e_x is the distribution that puts all mass on x). By Rockafellar (1970, Corollary 23.8.1), for C of this form and (P, z) a point on the boundary of C , the normal cone of C at (P, z) is the sum of the individual normal cones. The normal cone of any set at a point in the interior of that set is just $\{0\}$, so we can ignore those halfspaces when determining the normal cone. Then the corollary says that any vector $(-\lambda, 1)$ normal to f_0 at $(P, f_0(P))$ can be written as $\sum_{x \in \mathcal{X}: P^{\top} \lambda^{(x)} = f_0(P)} (-\alpha_x \lambda^{(x)}, \alpha_x)$: λ is a convex combination of those $\lambda^{(x)}$.

Conclusion: any minimal supporting hyperplane λ to $H_L \upharpoonright \Delta_y$ at $P \in \Delta_y$ with randomized 0-1 loss is a convex combination of the hyperplanes realized by hard 0-1 loss that are supporting at P . Therefore, randomizing allows the contestant to realize λ . $\square$

Proof of Theorem 6.7. From Theorem 6.5, we know that a strategy exists for the contestant that achieves loss $\sum_x p_x \lambda_x^*$ where λ^* is any KT-vector, and that

this is worst-case optimal. Hence Q^* is worst-case optimal if and only if it achieves the same worst-case expected loss. The worst-case expected loss of a strategy $Q \in \mathcal{Q}$ is

$$\max_{P \in \mathcal{P}} \sum_{x,y} P(x,y) L(x, Q|_y) = \sum_x p_x \max_{y \ni x} L(x, Q|_y).$$

Therefore if for all x, y with $x \in y$, we have $L(x, Q|_y) \leq \lambda_x^*$ for some KT-vector λ^* , Q is worst-case optimal.

For the converse, pick any $Q \in \mathcal{Q}$ and suppose the vector given by $\lambda_x := \max_y L(x, Q|_y)$ is not a KT-vector. Then by Theorem 6.3, there is no $P \in \mathcal{P}$ such that P and λ satisfy the conditions of that theorem. Equivalently, for all $P \in \mathcal{P}$, there either is a message y such that the hyperplane defined by λ passes below H_L somewhere in Δ_y , or there is a message y with $P(y) > 0$ but the hyperplane lies strictly above H_L at $P(\cdot | y)$. The former contradicts the definition of H_L , so for λ not a KT-vector, the latter must be the case. But then against any $P \in \mathcal{P}$ (in particular against worst-case optimal P), there is a different strategy $Q' \in \mathcal{Q}$ that is equal to Q except for its response to the message y : Q'_y realizes a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at $P(\cdot | y)$. This strategy Q' obtains strictly smaller expected loss than Q , so Q is not worst-case optimal. (In other words: in a Nash equilibrium (P^*, Q^*) , the contestant can only do worse against P^* by changing strategy, but here he can do better.) $\square$

Proof of Lemma 6.8. For proper loss functions, L and H_L are related as follows: at all $y \in \mathcal{Y}$ and $P \in \Delta_y$, if the vector $\lambda = L(\cdot, P)$ is finite at all $x \in y$, it describes the nonvertical hyperplane realized by P , which is a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at P .

Now suppose that at some $P \in \Delta_y$, there exists a minimal supporting hyperplane $\lambda' \in \Lambda_y$ other than $\lambda := L(\cdot | P)$, the supporting hyperplane realized by P (here we allow vertical hyperplanes, for which λ may include infinities). Let $x \in y$ be an outcome where $\lambda'_x < \lambda_x$, which exists by minimality of λ' . Write e_x for the probability distribution that puts all mass on this outcome, and define $P_\alpha := (1 - \alpha)P + \alpha e_x$ for $\alpha \in (0, 1]$. For each of these points P_α , the hyperplane $L(\cdot, P_\alpha)$ realized by P_α is at most as high as λ' at P_α (because $L(\cdot, P_\alpha)$ is supporting there) and at least as high as λ' at P (where λ' is supporting), so $L(x, P_\alpha)$ is bounded away from λ_x by λ'_x : $L(x, P_\alpha) \leq \lambda'_x < \lambda_x$. Therefore $\lim_{\alpha \downarrow 0} L(x, P_\alpha) \neq L(x, P)$, and L is not continuous. For L proper and continuous, this proves the ‘at most one’ part of the lemma.

For the ‘exactly one’ part: by Rockafellar (1970, Theorem 23.3), a nonvertical supporting hyperplane may only fail to exist at P if there is a line segment through P falling inside Δ_y on one side of P and outside on the other; that is, for P on the relative boundary of Δ_y .

Finally, suppose that $\lambda = L(\cdot, P)$ (the hyperplane realized by P) is finite at all $x \in y$ but not minimal. Then by Lemma 6.4, a different minimal supporting hyperplane λ' exists at P , which by the above gives a contradiction. This shows that if λ is finite, it is the minimal supporting hyperplane. $\square$

Proof of Theorem 6.9. Theorem 6.3 applies, showing existence of a KT-vector and a worst-case optimal strategy for the quizmaster.

A worst-case optimal Q^* exists and forms a Nash equilibrium with P^* : For each $y \in \mathcal{Y}$ and each $P \in \Delta_y$, by Lemma 6.8 there is at most one minimal supporting hyperplane at P which is then realized by $Q = P$. So all minimal supporting hyperplanes are realizable on y , and Theorem 6.5 applies.

Next we show that the characterization of P^* and λ^* in Theorem 6.3 is equivalent to the one in this theorem. We consider y with $P^*(y) > 0$ first. If P^* is a worst-case optimal strategy for the quizmaster λ^* defines a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at $P^*(\cdot | y)$, then by Lemma 6.4 there exists a minimal $\lambda' \in \Delta_y$ which is also supporting at $P^*(\cdot | y)$ and which satisfies $\lambda'_x \leq \lambda_x$ for $x \in y$, with equality for $P^*(x | y) > 0$. By Lemma 6.8, $\lambda'_x = L(\cdot, P^*(\cdot | y))$ for all $x \in y$, showing that the conditions of this theorem hold. Conversely, if λ^* satisfies the equality in this theorem, then $\sum_{x \in y} P^*(x | y) L(x, P(\cdot | y)) = H_L(P^*(\cdot | y))$, so λ^* defines a supporting hyperplane at $P^*(\cdot | y)$.

For y with $P^*(y) = 0$, if the hyperplane defined by λ^* is nowhere below $H_L \upharpoonright \Delta_y$ as in Theorem 6.3, then using Lemma 6.4 it can be lowered to become a minimal supporting hyperplane, for which a realizing $Q^*_{|y}$ exists; conversely, the existence of a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at $Q^*_{|y}$ that is nowhere above λ^* implies that λ^* is itself nowhere below $H_L \upharpoonright \Delta_y$.

Uniqueness of the KT-vector: For any worst-case optimal strategy P^* , the characterization in this theorem puts an equality constraint on λ_x^* for each x , so only one vector can satisfy these conditions. We just saw that these conditions are equivalent to those in Theorem 6.3, so λ^* is the unique KT-vector.

Characterization of Q^* : By Theorem 6.7, Q^* is worst-case optimal for the contestant if and only if the (unique) KT-vector equals the left-hand side of (6.9). Similarly, if a strategy P^* is worst-case optimal for the quizmaster, then the KT-vector equals the right-hand side of (6.9). Therefore: if for given Q^* a worst-case optimal P^* exists for which (6.9) holds, then both sides equal the KT-vector and Q^* is worst-case optimal; if Q^* is worst-case optimal, then (6.9) holds for all worst-case optimal P^* ; and if, for given Q^* , (6.9) holds for all worst-case optimal P^* , then it holds for at least one worst-case optimal P^* by the existence of worst-case optimal P^* . $\square$

Proof of Theorem 6.10. For local L , by definition $L(x, Q) = f_x(Q(x))$ for some sequence of functions $f_x : [0, 1] \rightarrow [0, \infty]$. Given a point $P \in \Delta_y$, the vector λ given by $\lambda_x = f_x(P(x))$ defines a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at P , because L is proper. Each f_x is nonincreasing. (To see this, consider moving the point P along any line that goes through the vertex of the simplex $\Delta_{\mathcal{X}}$ which puts all mass on some x . Because H_L is concave along this line, the farther away P is from that vertex, the higher a supporting hyperplane to H_L at P will be at that vertex.)

Given P^* and q satisfying the conditions in this theorem, let λ_x^* be $f_x(q_x)$ for each $x \in \mathcal{X}$. We show that P^* is worst-case optimal by verifying that P^* and λ^* satisfy Theorem 6.3. For each y with $P^*(y) > 0$, λ^* defines a supporting hyperplane to $H_L \upharpoonright \Delta_y$ at $P^*(\cdot | y)$. For each other y , consider a support-

ing hyperplane to $H_L \upharpoonright \Delta_y$ at $Q_{|y}(x) = q_x / \sum_{x' \in y} q_{x'}$: because $Q_{|y}(x) \geq q_x$, $f_x(Q_{|y}(x)) \leq f_x(q_x) = \lambda_x^*$, the hyperplane defined by λ^* is everywhere at least as high as this supporting hyperplane, as required.

For the converse: For strictly concave H_L , f_x is strictly decreasing. Define functions g_x as follows: $g_x(\lambda_x) = \inf\{q \in [0, 1] \mid f_x(q) \leq \lambda_x\}$. If f_x is continuous, g_x is just the ordinary inverse of f_x , but if f_x has a jump discontinuity at q , there will be an interval where g_x is constantly equal to q . In either case, g_x satisfies $g_x(f_x(q)) = q$ for all $q \in [0, 1]$.

Take P^* some worst-case optimal strategy, and λ^* a KT-vector. Define $q \in [0, 1]^{\mathcal{X}}$ by $q_x = g_x(\lambda_x^*)$. We use Theorem 6.3 to show that q satisfies (6.10). For each $y \in \mathcal{Y}$, let $\lambda' \in \Lambda_y$ be a minimal supporting hyperplane to $H_L \upharpoonright \Delta_y$ that obeys $\lambda' \leq \lambda^*$; such a λ' exists by Lemma 6.4. Let $Q_{|y}$ be the (unique) point at which λ' supports $H_L \upharpoonright \Delta_y$. It satisfies $f_x(Q_{|y}(x)) = \lambda'_x$, from which it follows that $g_x(\lambda'_x) = g_x(f_x(Q_{|y}(x))) = Q_{|y}(x)$. Applying g_x to both sides of $\lambda'_x \leq \lambda_x^*$, we get $Q_{|y}(x) \geq q_x$ for all $x \in y$, so that $\sum_{x \in y} Q_{|y}(x) \leq \sum_{x \in y} q_x = 1$. If $P^*(y) > 0$, then the hyperplane defined by λ^* is itself a supporting hyperplane and $Q_{|y}$ is the point where it touches $H_L \upharpoonright \Delta_y$, namely the point $P^*(\cdot \mid y)$. Because $Q_{|y}(x)$ also satisfies $Q_{|y}(x) = g_x(\lambda_x^*) = q_x$, the equality $q_x = Q_{|y}(x) = P^*(x \mid y)$ follows. $\square$

Proof of Lemma 6.11. At least one vector q must exist because for the game with logarithmic loss and $\mathcal{X}, \mathcal{Y}, p$ as in the lemma, a worst-case optimal strategy P^* must exist by Theorem 6.3, and an associated RCAR vector q must exist for it by Theorem 6.10 using that H_L is strictly concave.

For logarithmic loss, the RCAR vector q and KT-vector λ^* are related by $\lambda_x^* = -\log q_x$. By Theorem 6.3, any strategy $P \in \mathcal{P}$ that does not agree with the KT-vector λ^* is not worst-case optimal, showing that q is unique. $\square$

Proof of Lemma 6.13. We will show that any nonvertical supporting hyperplane $\lambda \in \Lambda_y$ is supporting at no more than one point $P \in \Delta_y$. By Lemma 6.4, if a supporting hyperplane exists at P , then a minimum supporting hyperplane also exists at that point, so it suffices to restrict our attention to minimal $\lambda \in \Lambda_y$. We know that such a λ is realizable on y ; let Q be a distribution realizing it. Then Q minimizes the expected loss against any P at which λ supports $H_L \upharpoonright \Delta_y$. For strictly proper L , there can be at most one such P , proving strict concavity. $\square$

Proof of Lemma 6.14. The generalized entropy function of L' is given by

$$H_{L'}(P) = aH_L(P) + \sum_{x \in \mathcal{X}} b_x P(x),$$

where $a \in \mathbf{R}_{>0}$ and $b \in \mathbf{R}^{\mathcal{X}}$ are the constants in the affine transformation (6.12). $H_{L'}$ is again finite and continuous. If P^* is worst-case optimal for the quizmaster in game $\mathcal{G}$, then by Theorem 6.3 there exists a KT-vector λ^* satisfying that theorem's conditions. Define a transformed vector by $\lambda' = a\lambda^* + b$. This is a KT-vector for $\mathcal{G}'$, showing that P^* is also worst-case optimal in that game.

If the conditions of Theorem 6.5 hold for $\mathcal{G}$, then they also hold for $\mathcal{G}'$: If λ' is a minimal supporting hyperplane to $H_{L'} \upharpoonright \Delta_y$, then $\lambda = (1/a)(\lambda' - b)$ is a minimal supporting hyperplane to $H_L \upharpoonright \Delta_y$. By assumption, λ is realizable on y in game $\mathcal{G}$, say by $Q \in \Delta_{\mathcal{X}}$. Then the same Q also realizes λ' in game $\mathcal{G}'$.

If Q^* is worst-case optimal for the contestant in $\mathcal{G}$, then by Theorem 6.7, λ given by $\lambda_x := \max_{y \ni x} L(x, Q^*|_y)$ is a KT-vector. The transformed vector $\lambda' = a\lambda + b$ is then a KT-vector in $\mathcal{G}'$, so that by Theorem 6.7, Q^* is worst-case optimal in that game.

Because the affine transformation from L into L' can be reversed by a second affine transformation (with $a' = 1/a$ and $b' = -(1/a)b$), the reverse implications follow. $\square$