



Original Article

Best-classifier feedback in diagnostic classification training

Corey J. Bohil*, Andrew J. Wismer, Troy A. Schiebel, Sarah E. Williams

Department of Psychology, University of Central Florida, United States



ARTICLE INFO

Article history:

Received 22 March 2015

Received in revised form 28 July 2015

Accepted 29 July 2015

Available online 7 August 2015

Keywords:

Feedback

Diagnosis

Training

Classification

Optimal-classifier

ABSTRACT

Diagnostic classification training requires viewing many examples along with category membership feedback. “Objective” feedback based on category membership suggests that perfect accuracy is attainable when it may not be (e.g., with confusable categories). Previous work shows that feedback based on an “optimal” responder (that sometimes makes classification errors) leads to higher long-run reward, especially in unequal category payoff conditions. In the current study, participants learned to classify normal or cancerous mammography images, earning more points for correct “cancer” than “normal” responses. Feedback was either objective or based on performance of an empirically determined “best” classifier. This approach is necessary because theoretically optimal responses cannot be determined with complex real-world stimuli with unknown perceptual distributions. Replicating earlier work that used simple artificial stimuli, we found that best-classifier performance led to decision-criterion values (β) closer to the reward-maximizing criterion, along with higher point totals and a slight reduction (as predicted) in overall accuracy.

© 2015 Society for Applied Research in Memory and Cognition. Published by Elsevier Inc. All rights reserved.

In many diagnostic domains, such as radiology, there exists substantial variability among experts in decision rule adoption despite similar overall accuracy rates (Swets, 1998). Beyond individual differences, diagnostic criterion placement can vary with circumstances. For example, limiting unnecessary biopsies requires a conservative (requiring strong evidence for an “abnormal” judgment) criterion for routine mammograms but a more lenient criterion following a doctor’s referral. To maximize consistency across practitioners, and to afford policymakers some level of control over long run diagnostic accuracy rates and outcomes, it would be beneficial if “optimal” decision criterion placement could be trained.

Diagnostic classification judgments such as those made by dermatologists or radiologists require expertise gained from hundreds of hours of practice (Gunderman, Nyce, & Steele, 2002). Training for this type of classification involves viewing numerous examples from diagnostic categories along with some type of feedback (e.g., correct category membership). This enables trainees to learn a classification decision-rule that maximizes either long-run accuracy or some measure of reward (Maddox, 2002). Although perfect accuracy is desirable, it is often impossible due to the

confusability of diagnostic categories. Achieving the best possible long-run performance requires sensitivity to category payoffs (benefits and costs for different outcomes) and base-rates (relative prevalence of diagnostic alternatives). Although base-rates are of critical importance, the current research focuses on category payoffs. We consider base-rates further in Section 5.

Many studies have applied signal detection theoretic (SDT) analysis (or related computational model-based analysis) to examine payoff influence over classification performance (Bohil & Maddox, 2003a; Busemeyer & Myung, 1992; Busemeyer & Rappaport, 1988; Erev, 1998; Maddox & Bohil, 2000; Maddox & Bohil, 2004; Stevenson, Busemeyer, & Naylor, 1991). SDT assumes confusable categories (typically sampled from overlapping Gaussian distributions; see Fig. 1). The category payoff values determine the decision criterion (β) used by an “optimal” classifier (Green & Swets, 1966) that maximizes long-run reward (β_{rew} in Fig. 1), which serves as a benchmark for understanding human performance. Because categories overlap, classification errors are inevitable even for the optimal classifier.

1. Competition between reward and accuracy (COBRA)

With unequal category payoffs, the reward (β_{rew}) and accuracy (β_{acc}) maximizing criterion values are not the same. In fact, using the reward-maximizing criterion actually leads to a reduction in overall accuracy. If responding is biased to favor a higher payoff category then accuracy increases for that category at the expense of

* Corresponding author at: Department of Psychology, University of Central Florida, 4111 Pictor Lane Building 99, Orlando, FL 32816, United States. Tel.: +1 407 823 2755; fax: +1 407 823 5862.

E-mail address: corey.bohil@ucf.edu (C.J. Bohil).

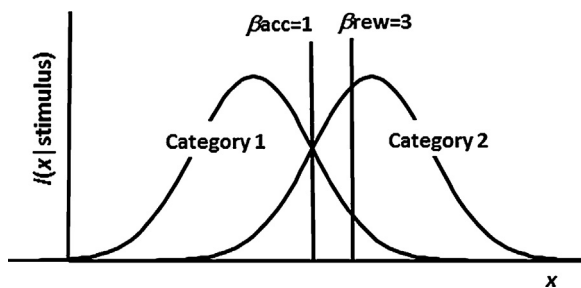


Fig. 1. β_{acc} = accuracy maximizing criterion; β_{rew} = reward maximizing criterion; x = perceptual effect created by presentation of a stimulus; $f(x|\text{stimulus})$ = likelihood of perceptual effect x given the stimulus.

the lower payoff category. For example, assuming an equal number of stimuli from each category, overall accuracy is maximized using a decision criterion $\beta = 1$ (β_{acc} in Fig. 1). Given a 3:1 Category 1 to Category 2 payoff ratio, $\beta = 3$ (β_{rew} in Fig. 1) maximizes long-run reward.

Maddox and Bohil (2003, 2004); (see also Bohil & Maddox, 2003a) tested the hypothesis that there is competition between reward and accuracy (COBRA) maximization goals. They argued that learners attempt to maximize payoff on each trial, but erroneously believe that maximizing accuracy is the way to achieve this goal. Observed β values were generally greater than one, indicating sensitivity to the payoff ratio. However, β values were typically “conservative” relative to the reward maximizing criterion. Participants consistently failed to adjust criterion values far enough to achieve maximum long-run reward (this result has been found in many similar studies, e.g., von Winterfeldt & Edwards, 1982).

Maddox and Bohil (2001) noted that the feedback used in classification training often implies a level of performance that is actually unattainable even for the optimal classifier. Due to category overlap, stimuli are often misclassified when they fall into the wrong response region relative to the reward maximizing rule. For example, a stimulus sampled from Category 1’s upper tail may look like a good example of Category 2, and if its percept falls to the right of β_{rew} it will elicit a “Category 2” response. Even though the response would be correct relative to the optimal reward-maximizing criterion, the actual category membership of this stimulus would lead to feedback indicating an incorrect response.

We refer to feedback based on actual category membership as “objective” feedback. Objective category membership is the basis for learning feedback in virtually all classification training. When categories are confusable this feedback is misleading, as it suggests a better decision rule might be found when in fact the learner may already be using the optimal rule.

Maddox and Bohil (2001, 2005); (see also Bohil & Maddox, 2003b) conducted several studies comparing category learning with “objective” versus “optimal” classifier feedback. Optimal classifier feedback following each trial revealed the optimal classifier’s response (see Fig. 3; note that “optimal” is replaced with “best” in Fig. 3 – details below). The “optimal” classifier uses the reward-maximizing decision criterion but still gets many classification responses wrong from the standpoint of objective (actual) category membership. If participants attend to the payoff ratio and attempt to maximize reward (i.e., move their criterion in the β_{rew} direction), then they must learn to sacrifice accuracy, which is maximized by β_{acc} . Feedback (e.g., objective category membership) that contradicts this goal should limit their ability to adjust in the optimal direction. Maddox and Bohil consistently found that optimal-classifier feedback led to performance that was closer to optimal (β values were larger while overall accuracy was reduced).

Optimal classifier feedback decouples the goal of maximizing reward from the strategy of attempting to achieve perfect

accuracy (in effort to maximize reward). By deemphasizing accuracy on every trial, this decoupling helps learners adopt a criterion that is closer to the reward maximizing value than is typically achieved after feedback training. We believe this finding may have important practical applications.

2. Best-classifier feedback for mammography training

In Maddox and Bohil’s (2001, 2005); (Bohil & Maddox, 2003b) studies, the categorization stimulus (bar graph height) values were sampled from overlapping normal distributions whose means and variances were determined a priori. These controlled conditions make it possible to determine the optimal classifier’s response to each stimulus. Our goal in the current research was to determine whether this training feedback manipulation could translate to a more complex (i.e., high dimensional; less controlled) stimulus set based on images used in real diagnostic classification training. We examined classification learning with unequal category payoffs, and our stimuli were mammography images showing either presence or absence of cancer. Participants were trained with feedback based either on objective category membership or on performance of an empirically determined “best” classifier. This approach reflects the fact that in real-world diagnostic training the responses of a best performer (however defined) at the task can be used as a training signal for others.

Because we used stimuli drawn from populations with unknown characteristics, we could not derive the performance of an “optimal” classifier for feedback. We instead conducted a small pilot-study using objective feedback and selected the participant whose β was closest to optimal ($\beta_{opt} = 3$ in the pre-study) as a “best” classifier. The study presented in this article used this person’s responses as “best-classifier” feedback during training.

We expected to replicate Maddox and Bohil’s (2001, 2005); (Bohil & Maddox, 2003b) earlier findings. Decision criterion values (β) should be closer to the reward maximizing value after training with best-classifier feedback. This should be met with a corresponding decline in overall accuracy rate but an increase in total points.

3. Methods

3.1. Design

We manipulated the type of corrective feedback following each classification response in a mammography screening task. *Objective-classifier* feedback displayed the number of points that could have been earned had the objectively correct response been given (e.g., feedback based solely on category membership of the stimulus, implying perfect accuracy is possible). *Best-classifier* feedback displayed the number of points earned by the best performing classifier (described below) after each trial. Critically, this “best” classifier sometimes gives the incorrect response despite using the correct strategy (i.e., a single decision bound that divides the perceptual space into two response regions). Following an error the best-classifier earns 0 points. This feedback should help participants learn that 100% accuracy is not necessarily their goal and that they can sacrifice some accuracy in order to adopt a criterion that maximizes long-run reward. Because our participants were novices at mammography classification, half were shown a preview image prior to beginning the task (see Fig. 2) to help orient them to the stimulus types. This manipulation had no significant effect on any of our measures and is not considered further.

Each participant completed a single session of classification training with a 3.6:1 payoff ratio (3.6 points for correct “cancer” responses, 1 point for correct “normal” responses). This ratio was

Download English Version:

<https://daneshyari.com/en/article/881559>

Download Persian Version:

<https://daneshyari.com/article/881559>

[Daneshyari.com](https://daneshyari.com)