

Any Reasonable Cost Function Can Be Used for a Posteriori Probability Approximation

Marco Saerens*, Patrice Latinne[†] & Christine Decaestecker[‡]
Université Catholique de Louvain and Université Libre de Bruxelles
Belgium

October 18, 2002

Abstract

In this paper, we provide a straightforward proof of an important, but nevertheless little known, result obtained by Lindley in the framework of subjective probability theory. This result, once interpreted in the machine learning/pattern recognition context, puts new lights on the probabilistic interpretation of the output of a trained classifier. A learning machine, or more generally a model, is usually trained by minimizing a criterion – the expectation of the cost function – measuring the discrepancy between the model output and the desired output. In this letter, we first show that, for the binary classification case, training the model with any ‘reasonable cost function’ can lead to Bayesian a posteriori probability estimation. Indeed, after having trained the model by minimizing the criterion, there always exists a computable transformation that maps the output of the model to the Bayesian a posteriori probability of the class membership given the input. Then, necessary conditions allowing the computation of the transformation mapping the outputs of the model to the a posteriori probabilities are derived for the multi-output case. Finally, these theoretical results are illustrated through some simulation examples involving various cost functions.

*Marco Saerens is with the Information Systems Research Unit, IAG, Université Catholique de Louvain, 1 Place des Doyens, B-1348 Louvain-la-Neuve, Belgium. Email: saerens@isys.ucl.ac.be

[†]Patrice Latinne is with the IRIDIA Laboratory (Artificial Intelligence Laboratory), cp 194/6, Université Libre de Bruxelles, 50 avenue Franklin Roosevelt, B-1050 Brussels, Belgium. Email: platinne@ulb.ac.be.

[‡]Christine Decaestecker is a Research Associate with the Belgian Research Funds (F.N.R.S.) at the Laboratory of Histopathology, cp 620, Université Libre de Bruxelles, 808 route de Lennik, B-1070 Brussels, Belgium. Email: cdecaes@ulb.ac.be.

1. Introduction

An important problem concerns the probabilistic interpretation to be given to the output of a learning machine, or more generally a model, after training. It appears that this probabilistic interpretation depends on the cost function used for training. Classification models are almost always trained by minimizing a given criterion, the expectation of the cost function. It is therefore of fundamental importance to have a precise idea of what can be achieved with the choice of this criterion.

Consequently, there has been considerable interest in analyzing the properties of the mean square error criterion – the most commonly used criterion. It is for instance well-known that artificial neural nets (or more generally any model), when trained using the mean square error criterion, produce as output an approximation of the expected value of the desired output conditional on the explanatory input variables if ‘perfect training’ is achieved (see for instance [1],[5]). We say that **perfect training** is achieved if

- A minimum of the criterion is indeed reached after training, and
- The learning machine is a ‘sufficiently powerful model’ that is able to approximate the optimal estimator to any degree of accuracy (perfect model matching property).

It has also been shown that other cost functions, for instance the cross-entropy between the desired output and the model output in the case of pattern classification, lead to the same property of approximating the conditional expectation of the desired output as well. We may therefore wonder what conditions a cost function should satisfy in order that the model output has this property. In 1991, following the results of Hampshire & Pearlmutter [3], Miller, Goodman & Smyth [7], [8] answered to this question by providing conditions on the cost function ensuring that the output of the model approximates the conditional expectation of the desired output given the input, in the case of perfect training. These results were rederived by Saelens by using the calculus of variations [9], and were then extended to the conditional median [10]. Also, in [10], a close relationship between the conditions on the cost function ensuring that the output of the model approximates the conditional probability of the desired output given the input, when the performance criterion is minimized, and the quasi-likelihood functions used in the context of applied statistics (generalized linear models; see [6]) was pointed out.

In this work, we focus on **classification**, in which case the model will be called a classifier. In this framework, we show that, for the binary classification case, training the classifier with **any reasonable cost function** leads to a posteriori probability estimation. Indeed, after having trained the model by minimizing

the criterion, there always exists a **computable transformation** that maps the output of the model to the a posteriori probability of the class label. This means that we are free to choose any reasonable cost function we want, and train the classifier with it. We can always remap the output of the model afterwards to the a posteriori probability, for Bayesian decision making. We will see that this property generalizes to a certain extent to the multi-output case.

This important result was proved by Lindley in 1982, in the context of subjective probability theory [4]. Briefly, Lindley considered the case where a person expresses his uncertainty about an event E , conditional upon an event F , by assigning a number, x (we use Lindley's notations). For example, consider a physician who, after the medical examination of a patient, has to express his uncertainty about the diagnosis of a given disease (E), conditional on the result (F) of the examination. This physician then receives a score $f(x, I_E)$ which is function of x and the truth or falsity of E when F is true (where I_E is an indicator variables, i.e. $I_E = 1$ (0) if the event E is true (false)). The score function $f(x, I_E)$ can be interpreted as assigning a penalty or reward depending of the discrepancy between the person's response and the true state of the event E . It is assumed that the person wishes to reduce his expected score. Under a number of reasonable assumptions on the score function f and the possible values x which can be chosen by the person, Linsley proved that there exists a simple transform of the values x which map them on probabilities. This transform is a function of the values $f'(x, 0)$ and $f'(x, 1)$, the differentials of f .

In the present paper, we show that Lindley's approach can be applied in the machine learning/pattern recognition context in the case of pattern classification problems, leading to an interesting result concerning the cost functions used to train a classifier. Lindley's derivation was based on geometrical facts and reasoning, while our proof relies on standard differential calculus, and partly extends to the multiple class problem.

In the following sections, we first introduce the binary output problem from an estimation theory perspective (section 2). Then, we derive the transformation that must be applied to the output of the model to obtain the a posteriori probability of the desired output, given the input (section 3). Some results for the multi-output case are provided in section 4. Finally, examples of cost functions and corresponding mappings to a posteriori probabilities are presented in section 5. We conclude in section 6.

2. Statement of the two-class problem

Let us consider that we are given a sequence of N independent m -dimensional training patterns $\mathbf{x}_k = [x_1(k), x_2(k), \dots, x_m(k)]^T$, with $k = 1, 2, \dots, N$, as well as corresponding scalar desired outputs $y_k \in \{0, 1\}$ providing information about the class label of the pattern. If the observation $\mathbf{x}_k$ is assigned to the class label ω_0 ,

then $y_k = 0$; if it is assigned to the class label ω_1 , then $y_k = 1$. The $\mathbf{x}_k$ and the y_k are realizations of the random variables $\mathbf{x}$ and y . We hope that the random vector $\mathbf{x}$ provides some useful information that allows to predict the class label y with a certain accuracy. The objective is to train a model, say a neural network, in order to supply outputs $\hat{y}_k$ (we assume $0 \leq \hat{y}_k \leq 1$) that are 'accurate' (in some predefined manner; see below) estimations or predictions of the desired outputs y_k :

$$\hat{y}_k = N[\mathbf{x}_k, \mathbf{w}] \text{ with } 0 \leq \hat{y}_k \leq 1 \quad (2.1)$$

where $N[.,.]$ is the function provided by the model, $\mathbf{x}_k$ the input vector (the vector of explanatory variables) supplied to the model, and $\mathbf{w}$ is the parameter vector of the model. In order to measure how 'accurate' is the estimation (2.1), we define a **cost function** (or loss function, penalty function, objective function, empirical risk measure, scoring rule) that provides us a measure of the discrepancy between the **predicted value** $\hat{y}_k$ and the **desired value** y_k : $\mathcal{L}[\hat{y}_k; y_k]$. The purpose of the training is, of course, to estimate the parameters that minimize this cost.

Since it is generally not possible to minimize the cost function for each k because of the presence of noise or disturbances (for a given value of the input $\mathbf{x}$, the desired output is distributed with a probability density function $p(y|\mathbf{x})$), the best we can do is to minimize this cost 'on average'. This leads to the definition of the **performance criterion** $C[\hat{y}]$:

$$C[\hat{y}] = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{k=1}^N \mathcal{L}[\hat{y}_k; y_k] \quad (2.2)$$

$$= \iint \mathcal{L}[\hat{y}; y] p(\mathbf{x}, y) d\mathbf{x} dy = E_{\mathbf{x}y} \{ \mathcal{L}[\hat{y}; y] \} \quad (2.3)$$

where the integral is defined on the Euclidean space $\mathbb{R}^m \times \mathbb{R}^1$ and we assume that there are enough samples so that we can rely on the asymptotic form of the performance criterion. $E_{\mathbf{x}y}\{.\}$ is defined as the standard expectation.

It is convenient to rewrite (2.3)

$$C[\hat{y}] = \int \left\{ \int \mathcal{L}[\hat{y}; y] p(y|\mathbf{x}) dy \right\} p(\mathbf{x}) d\mathbf{x} \quad (2.4)$$

If we minimize the inner integral of (2.4) for every possible value of $\mathbf{x}$, then $C[\hat{y}]$ will also be minimized, since $p(\mathbf{x})$ is non negative. We therefore select $\hat{y}(\mathbf{x})$ in order to minimize the **conditional criterion**

$$C[\hat{y}|\mathbf{x}] = \int \mathcal{L}[\hat{y}; y] p(y|\mathbf{x}) dy = E_y \{ \mathcal{L}[\hat{y}; y] | \mathbf{x} \} \quad (2.5)$$

for every $\mathbf{x}$, where $C[\hat{y}|\mathbf{x}]$ is a function of both $\hat{y}$ and $\mathbf{x}$, and $E_y\{.\}|\mathbf{x}$ is the conditional expectation, given $\mathbf{x}$. This means that the minimization of (2.5) can

be performed independently for every $\mathbf{x}$. Moreover, since $\hat{y}$ is chosen in order to minimize (2.5) for every value of $\mathbf{x}$, this $\hat{y}$ will be a function of $\mathbf{x}$. The function of $\mathbf{x}$ that minimizes (2.5) will be called the best or **optimal** estimator, and will be denoted by $\hat{y}^*(\mathbf{x})$.

We assume that this optimal estimator can be approximated to any degree of accuracy by the model, $\hat{y} = N[\mathbf{x}, \mathbf{w}]$, for some optimal value of the parameters $\mathbf{w} = \mathbf{w}^*$ (perfect parameters tuning: $\hat{y}^*(\mathbf{x}) = N[\mathbf{x}, \mathbf{w}^*]$). In other words, we are making a 'perfect model matching' assumption. In the Miller, Goodman & Smyth terminology [7], [8], such a model is called a **sufficiently powerful model** that is able to produce the optimal estimator.

Notice that in the case of **binary classification** ($y \in \{0, 1\}$), the probability density $p(y|\mathbf{x})$ in (2.5) reduces to

$$p(y|\mathbf{x}) = p(y = 0|\mathbf{x}) \delta(y - 0) + p(y = 1|\mathbf{x}) \delta(y - 1) \quad (2.6)$$

Where $\delta(x)$ is the Dirac delta distribution. The conditional criterion (2.5) can therefore be rewritten as

$$C[\hat{y}|\mathbf{x}] = p(y = 0|\mathbf{x}) \mathcal{L}[\hat{y}; 0] + p(y = 1|\mathbf{x}) \mathcal{L}[\hat{y}; 1] \quad (2.7)$$

In the next section, we define a class of 'reasonable cost functions', and we derive the transformation that maps the output of the trained model $\hat{y}^*$ to the a posteriori probability $p(\omega_1|\mathbf{x}) = p(y = 1|\mathbf{x}) = E_y\{y|\mathbf{x}\}$.

3. Mapping the output of the trained model to the a posteriori probability (binary output case)

3.1. A class of reasonable cost functions

For training our classifier, we must choose a cost function that measures the discrepancy between the model's output and the observed desired output. For this purpose, we will consider the class of cost functions $\mathcal{L}[\hat{y}; y]$ of the type

$$\begin{cases} \mathcal{L}[\hat{y}; y] &= 0 \text{ if and only if } \hat{y} = y \\ \mathcal{L}[\hat{y}; y] &> 0 \text{ if } \hat{y} \neq y \\ \mathcal{L}[\hat{y}; y] &\text{is twice continuously differentiable} \\ &\text{in terms of all its arguments} \end{cases} \quad (3.1)$$

We also make the natural requirement that when the predicted value $\hat{y}$ moves away from the desired value y , the cost $\mathcal{L}[\hat{y}; y]$ increases. Symmetrically, the cost $\mathcal{L}[\hat{y}; y]$ should decrease when the predicted value $\hat{y}$ approaches the desired value y . This implies that

$$\frac{\partial \mathcal{L}[\hat{y}; y]}{\partial \hat{y}} \text{ is } \begin{cases} > 0 & \text{if } \hat{y} > y \\ < 0 & \text{if } \hat{y} < y \end{cases} \quad (3.2)$$

and, together with (3.1), that

$$\left. \frac{\partial \mathcal{L}[\hat{y}; y]}{\partial \hat{y}} \right|_{\hat{y}=y} = 0 \quad (3.3)$$

Finally, we also assume that $\mathcal{L}[\hat{y}; y]$ depends on $\mathbf{x}$ only through the variable $\hat{y}$.

Equations (3.1), (3.2) and (3.3) define the class of **reasonable cost functions** we will be working with. Some examples of such cost functions are provided in section 5.

3.2. Minimizing the criterion

Suppose now that we choose to train a sufficiently powerful model with one of these reasonable cost functions. This means that we pick up the model parameters $\mathbf{w}^*$ that minimize the performance criterion defined by equation (2.3), or equivalently the conditional criterion (2.7), in the binary classification case. The conditional criterion $C[\hat{y}|\mathbf{x}]$ is therefore minimized for some optimal value $\hat{y}^*(\mathbf{x}) = N[\mathbf{x}, \mathbf{w}^*]$ – since we assume that the model is "perfect", optimizing with respect to $\mathbf{w}$ is equivalent to optimizing with respect to $\hat{y}$. This value, $\hat{y}^*$, is the optimal output with respect to the criterion $C[\hat{y}|\mathbf{x}]$ defined by (2.7). This means that the following standard optimality conditions must hold:

$$\left. \frac{\partial C[\hat{y}|\mathbf{x}]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} = 0 \quad (3.4)$$

$$\left. \frac{\partial^2 C[\hat{y}|\mathbf{x}]}{\partial \hat{y}^2} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} > 0 \quad (3.5)$$

for every $\mathbf{x}$. The model therefore supplies, after training, the output $\hat{y}^*(\mathbf{x})$ representing some 'degree of plausibility' of the event $y = 1$.

We can easily show that the minimum of $C[\hat{y}|\mathbf{x}]$ lies in the interval $[0, 1]$ ($\in [0, 1]$). Indeed, from (2.7),

$$\frac{\partial C[\hat{y}|\mathbf{x}]}{\partial \hat{y}} = p(y = 0|\mathbf{x}) \frac{\partial \mathcal{L}[\hat{y}; 0]}{\partial \hat{y}} + p(y = 1|\mathbf{x}) \frac{\partial \mathcal{L}[\hat{y}; 1]}{\partial \hat{y}}$$

and since $\partial \mathcal{L}[\hat{y}|\mathbf{x}]/\partial \hat{y} > 0$ when $\hat{y} > y$ (3.2), for $\hat{y} > 1$,

$$\begin{aligned} \frac{\partial C[\hat{y}|\mathbf{x}]}{\partial \hat{y}} &= p(y = 0|\mathbf{x}) \frac{\partial \mathcal{L}[\hat{y}; 0]}{\partial \hat{y}} + p(y = 1|\mathbf{x}) \frac{\partial \mathcal{L}[\hat{y}; 1]}{\partial \hat{y}} \\ &> 0 \text{ for } \hat{y} > 1 \end{aligned}$$

so that $C[\hat{y}|\mathbf{x}]$ is continuously increasing when $\hat{y} > 1$ and $\hat{y}$ increases above 1.

Symmetrically, we can show in a similar manner that $C[\hat{y}|\mathbf{x}]$ is continuously increasing when $\hat{y} < 0$ and $\hat{y}$ decreases below 0 ($\partial C[\hat{y}|\mathbf{x}]/\partial \hat{y} < 0$ for $\hat{y} < 0$).

The minimum of $C[\hat{y}|\mathbf{x}]$ therefore $\in [0, 1]$, and the fact that the output of the model $\in [0, 1]$ ($0 \leq \hat{y} \leq 1$; see (2.1)) is not a restriction at all, since the minimum is always attainable (it lies in $[0, 1]$).

3.3. The mapping to a posteriori probabilities

Now that we have trained our model by optimizing the criterion $C[\hat{y}|\mathbf{x}]$, the model provides as output $\hat{y}^*(\mathbf{x})$ verifying (3.4).

In the appendix, we show that there always exists a transformation $f(\hat{y}^*)$ that maps the model's optimal output $\hat{y}^*(\mathbf{x})$ to the a posteriori probability $p(\omega_1|\mathbf{x}) = p(y = 1|\mathbf{x})$. This transformation is

$$f(\hat{y}^*) = \frac{\mathcal{L}'[\hat{y}^*; 0]}{\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]} \quad (3.6)$$

where $\mathcal{L}'[\hat{y}^*; 0] = \partial \mathcal{L}[\hat{y}; 0]/\partial \hat{y}|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$ and $\mathcal{L}'[\hat{y}^*; 1] = \partial \mathcal{L}[\hat{y}; 1]/\partial \hat{y}|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$. (3.6) can be rewritten in symmetric form:

$$f(\hat{y}^*) = \frac{1}{1 - \frac{\mathcal{L}'[\hat{y}^*; 1]}{\mathcal{L}'[\hat{y}^*; 0]}} \quad (3.7)$$

This mapping transforms the optimal output of the model to the a posteriori probability

$$f(\hat{y}^*(\mathbf{x})) = p(y = 1|\mathbf{x}) = p(\omega_1|\mathbf{x}) \quad (3.8)$$

Moreover, we also show in the appendix that if (3.5) holds for every $\hat{y}^* \in [0, 1]$, the mapping is one-to-one.

More precisely, in the Appendix A, we show that if the model has been trained by optimizing the criterion (it supplies optimal values $\hat{y}^*$ verifying (3.4)) and if there exists a mapping that transforms the output of the model to the a posteriori probabilities (3.8), then this mapping is given by (3.6).

In the Appendix B, we show that if the model has been trained by optimizing the criterion (3.4) and we transform the model's output $\hat{y}^*$ by (3.6), then the result of the mapping is $p(y = 1|\mathbf{x})$, the a posteriori probability of observing $y = 1$ conditional on $\mathbf{x}$ (equation 3.8).

Finally, in the Appendix C, we show that a second-order condition (3.5) holding for every $\hat{y}^* \in [0, 1]$ is equivalent to a strictly monotonic increasing $f(\hat{y}^*)$ on $[0, 1]$. In this case, the mapping is one-to-one, and the conditional criterion has only one global minimum (no local minimum). On the contrary, a non-monotonic

increasing mapping (i.e. the function $f(\hat{y}^*)$ is stationary or decreasing on some interval $\in]0, 1[$) is associated with multiple local minima of the conditional criterion, for some value of $p(y = 1|\mathbf{x})$. We should therefore restrict the class of reasonable cost functions to those that have a strictly monotonic increasing mapping $f(\hat{y}^*)$.

It is easy to verify that (3.6) is a function that maps the interval $[0, 1]$ on $[0, 1]$. Indeed, by examining (3.6), from (3.2) and the fact that $0 \leq \hat{y}^* \leq 1$ (2.1), we easily find that $(\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]) > 0$ and $\mathcal{L}'[\hat{y}^*; 0] > 0$, so that $f(\hat{y}^*) \geq 0$ for $\hat{y}^* \in [0, 1]$. Moreover, from (3.2), (3.3), and the fact that $f(\hat{y}^*)$ is continuous, we deduce that $f(0) = 0$, $f(1) = 1$ and that $0 \leq f(\hat{y}^*) \leq 1$ (the equation $f(\hat{y}^*) = 0$ has only one solution, $\hat{y}^* = 0$, on $[0, 1]$; similarly, $f(\hat{y}^*) = 1$ has only one solution, $\hat{y}^* = 1$, so that $f(\hat{y}^*)$ remains in $[0, 1]$). The transformation $f(\hat{y}^*)$ is therefore a function that maps the interval $[0, 1]$ on $[0, 1]$ (see section 5 for examples of mappings).

A remarkable property of (3.6) is the fact that the mapping only depends on the cost function $\mathcal{L}$ and, in particular, does not depend on $p(y|\mathbf{x})$. Moreover, we can easily show that if the cost function verifies the conditions that lead to the estimation of the a posteriori probability (stated in [3] and reproduced in [9]), the mapping reduces to $f(\hat{y}^*) = \hat{y}^*$.

A consequence of these results is that we are free to choose any reasonable cost function in order to train the classification model. If we need the a posteriori probabilities, we compute the mapping (3.6) in order to obtain an approximation of the Bayesian a posteriori probabilities.

Notice, however, that all our results are essentially asymptotic, and that issues regarding estimation from finite data sets are not addressed.

4. Some results for the multi-output case

All the previously derived results concern the binary output case. In this section, we will discuss the multi-output case, for which necessary conditions for obtaining a mapping to the a posteriori probabilities will be derived. However, the obtained results will be far less general than for the binary case.

In the multi-output case, we will consider that, for each training pattern $\mathbf{x}_k$, there is a corresponding **desired output vector** $\mathbf{y}_k$, where each $\mathbf{y}_k$ is associated to one of n mutually exclusive classes. That is, $\mathbf{y}_k$ will indicate the **class label** $\in \{\omega_1, \dots, \omega_n\}$ of the observation $\mathbf{x}_k$. Each class label ω_i will be represented numerically by an **indicator vector** $\mathbf{e}_i$: if the observation $\mathbf{x}_k$ of the training set is assigned to the class label ω_i , then $\mathbf{y}_k = \mathbf{e}_i = [0, \dots, 0, 1, 0, \dots, 0]^T$.

Correspondingly, the neural network provides a predicted value vector as output:

$$\hat{\mathbf{y}}_k = N[\mathbf{x}_k, \mathbf{w}] \quad (4.1)$$

with $\hat{\mathbf{y}}_k = [\hat{y}_1(k), \hat{y}_2(k), \dots, \hat{y}_n(k)]^T$. We will assume that the outputs of the neural network sum to one ($\sum_{i=1}^n \hat{y}_i = 1$), as it is often the case for classification models (see for example the case of a softmax nonlinearity [1], or a logistic regression model [2]). This means that the output vector $\hat{\mathbf{y}}$ has only $n - 1$ degrees of freedom, and can be represented by $\hat{\mathbf{y}} = [\hat{y}_1, \hat{y}_2, \dots, \hat{y}_{n-1}, 1 - \sum_{i=1}^{n-1} \hat{y}_i]^T$.

Now, notice, as a particular case, that the mapping (3.6) can be applied to multi-output classifiers, provided that they are trained with a cost function which is a sum of individual scores, each score depending only on one output.

In full generality, for training the model, we will consider the class of cost functions $\mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}]$ of the type

$$\begin{cases} \mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}] = 0 & \text{if and only if } \hat{\mathbf{y}} = \mathbf{y} \\ \mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}] > 0 & \text{if } \hat{\mathbf{y}} \neq \mathbf{y} \\ \mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}] & \text{is twice continuously differentiable} \\ & \text{in terms of all its arguments} \end{cases} \quad (4.2)$$

By following the same steps as in section 2, the conditional criterion can be written as

$$C[\hat{\mathbf{y}}|\mathbf{x}] = \int \mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}] p(\mathbf{y}|\mathbf{x}) d\mathbf{y} = E_{\mathbf{y}}\{\mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}]|\mathbf{x}\} \quad (4.3)$$

In the classification case, the conditional criterion reduces to

$$C[\hat{\mathbf{y}}|\mathbf{x}] = \sum_{j=1}^n p(\mathbf{y} = \mathbf{e}_j|\mathbf{x}) \mathcal{L}[\hat{\mathbf{y}}; \mathbf{e}_j] \quad (4.4)$$

A necessary set of equations for $\hat{\mathbf{y}}^*$ to be an optimum of the criterion is given by

$$\left. \frac{\partial C[\hat{\mathbf{y}}|\mathbf{x}]}{\partial \hat{y}_i} \right|_{\hat{\mathbf{y}}(\mathbf{x})=\hat{\mathbf{y}}^*(\mathbf{x})} = 0, \text{ for } i = 1 \dots n - 1 \quad (4.5)$$

Notice that there are only $n - 1$ equations since we replaced $\hat{y}_n$ by $(1 - \sum_{i=1}^{n-1} \hat{y}_i)$.

In the Appendix D, we show that if there exists a mapping of the outputs of the model to the a posteriori probabilities

$$f_i(\hat{\mathbf{y}}^*(\mathbf{x})) = p(\mathbf{y} = \mathbf{e}_i|\mathbf{x}) = p(\omega_i|\mathbf{x}), \text{ for } i = 1 \dots n - 1, \quad (4.6)$$

this mapping is provided by solving the following system of $n - 1$ equations in terms of the $f_j(\hat{\mathbf{y}}^*)$

$$\boxed{\sum_{j=1}^{n-1} \left[1 - \frac{\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_j]}{\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_n]} \right] f_j(\hat{\mathbf{y}}^*) = 1, \text{ for } i = 1 \dots n - 1} \quad (4.7)$$

where $\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_j] = \partial \mathcal{L}[\hat{\mathbf{y}}; \mathbf{e}_j] / \partial \hat{y}_i|_{\hat{\mathbf{y}}(\mathbf{x}) = \hat{\mathbf{y}}^*(\mathbf{x})}$, and $f_n(\hat{\mathbf{y}}^*) = 1 - \sum_{i=1}^{n-1} f_i(\hat{\mathbf{y}}^*)$.

However, we were not able to provide sufficient conditions for the multi-output case. Indeed, several conditions should be checked before being able to state that these transformations exist and map the outputs to the a posteriori probabilities:

- After having minimized the criterion, we cannot be sure that the output values $\hat{y}_i^* \in [0, 1]$;
- We should check that the system of equations (4.7) has indeed a solution;
- For $\hat{\mathbf{y}}^*$ to be a minimum of $C[\hat{\mathbf{y}}|\mathbf{x}]$, the matrix of second-order derivatives should be definite positive.

For the rather general cost function definition that we defined, these conditions are quite difficult to assess, and should be verified on a case-by-case basis, for the cost function being used.

5. Some examples

In this section, we provide examples of mappings to a posteriori probabilities. We consider six different cost functions, plot the corresponding mapping (3.6), and examine the effect of the mapping on the optimal output.

The six cost functions are:

$$\mathcal{L}[\hat{y}; y] = \exp[y] (y - \hat{y} - 1) + \exp[\hat{y}] \quad (5.1)$$

$$\mathcal{L}[\hat{y}; y] = (\hat{y} - y)^4 \quad (5.2)$$

$$\mathcal{L}[\hat{y}; y] = 1 - \exp \left[-(\hat{y} - y)^2 \right] \quad (5.3)$$

$$\mathcal{L}[\hat{y}; y] = \log \left[1 + (\hat{y} - y)^2 \right] \quad (5.4)$$

$$\mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}] = \log \left[1 + \|\hat{\mathbf{y}} - \mathbf{y}\|^2 \right] \quad (5.5)$$

$$\mathcal{L}[\hat{\mathbf{y}}; \mathbf{y}] = \exp \left[\|\hat{\mathbf{y}} - \mathbf{y}\|^2 \right] + \exp \left[-\|\hat{\mathbf{y}} - \mathbf{y}\|^2 \right] - 2 \quad (5.6)$$

These cost functions are displayed in table 5.1 and the corresponding mappings $f(\hat{y}^*)$ provided by equation (3.6) and (4.7) are displayed in table 5.2. The first four cost functions ((5.1)-(5.4)) illustrate the binary output case; the two last cost functions ((5.5), (5.6)) illustrate a 3-output problem. In the later case (two last graphs of table 5.2), we show the mapping $f_1(\hat{y}_1^*, \hat{y}_2^*, \hat{y}_3^*)$ with $\hat{y}_1^* \in [0, 0.8]$, $\hat{y}_2^* = 0.2$ and $\hat{y}_3^* = (1 - \hat{y}_2^* - \hat{y}_3^*)$.

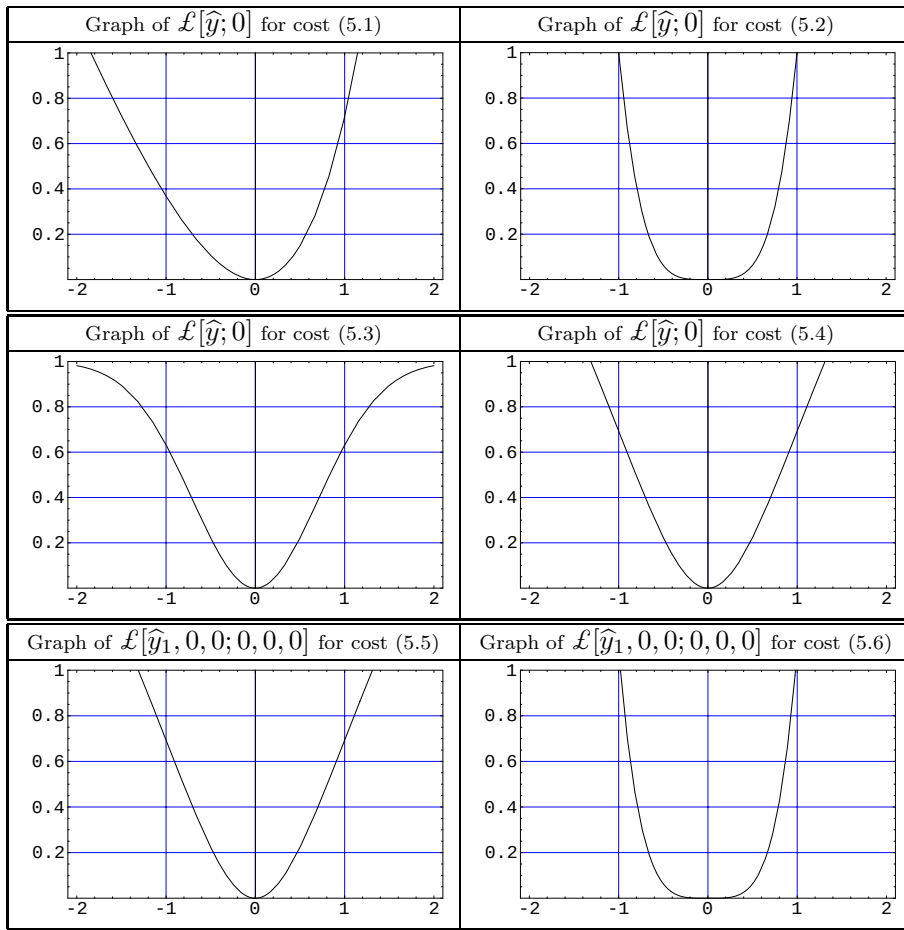


Table 5.1: Graph of the six cost functions. The first four functions illustrate the binary case; the two last ones illustrate the multi-output case (3 outputs).

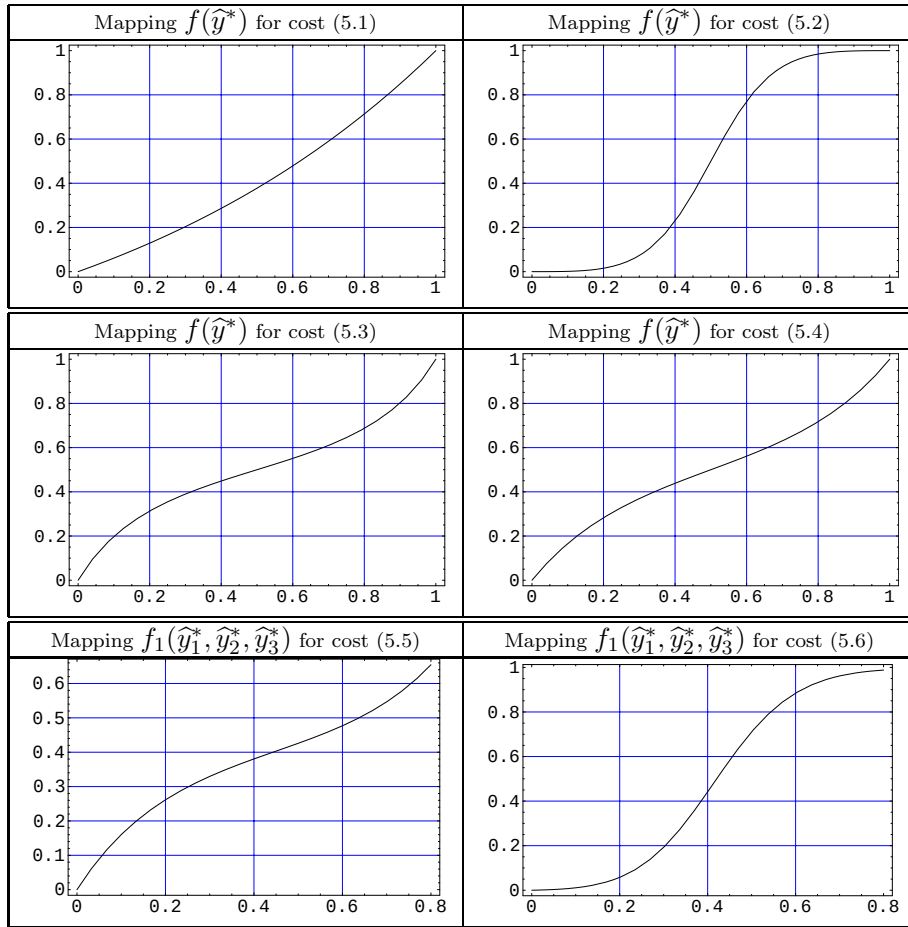


Table 5.2: Graph of the mapping to the a posteriori probabilities, for the six cost functions, as provided by equation (3.6) (binary case) and (4.7) (multi-output case).

By using the Mathematica software [11], for each of the first four cost functions ((5.1)-(5.4); binary output case), we compute the minimum $\hat{y}^*$ of the criterion

$$\begin{aligned} C[\hat{y}|\mathbf{x}] &= p(y=0|\mathbf{x}) \mathcal{L}[\hat{y}; 0] + p(y=1|\mathbf{x}) \mathcal{L}[\hat{y}; 1] \\ &= (1 - p(y=1|\mathbf{x})) \mathcal{L}[\hat{y}; 0] + p(y=1|\mathbf{x}) \mathcal{L}[\hat{y}; 1] \end{aligned} \quad (5.7)$$

for different values of $p(y=1|\mathbf{x})$ ranging from 0 to 1, illustrating all the potential situations that can occur (table 5.3, plain line). These are the optimal outputs of the model corresponding to different class distributions $p(y=1|\mathbf{x})$ that can be encountered in a binary classification problem. Notice that $\mathbf{x}$ does not play any role here since all our probability densities are conditioned on $\mathbf{x}$.

Then, we transform the output $\hat{y}^*$ by using the mapping $f(\hat{y}^*)$ (3.6) and plot the results in terms of $p(y=1|\mathbf{x})$ (table 5.3, dash line). We clearly observe that the transformed output is mapped on the a posteriori probability ($f(\hat{y}^*) = p(y=1|\mathbf{x})$).

For the multi-output case (two last cost functions (5.5), (5.6)), we plot the output $\hat{y}_1^*$ before remapping (y-axis, plain line) and after remapping by $f_1(\hat{y}_1^*, \hat{y}_2^*, \hat{y}_3^*)$ (see (4.7)) (y-axis, dash line), in function of the a posteriori probability $p(\mathbf{y} = \mathbf{e}_1|\mathbf{x})$ (x-axis), for values of $p(\mathbf{y} = \mathbf{e}_1|\mathbf{x}) \in [0, 0.8]$, $p(\mathbf{y} = \mathbf{e}_2|\mathbf{x}) = 0.2$, $p(\mathbf{y} = \mathbf{e}_3|\mathbf{x}) = 1 - p(\mathbf{y} = \mathbf{e}_1|\mathbf{x}) - p(\mathbf{y} = \mathbf{e}_2|\mathbf{x})$ (see table 5.3).

6. Conclusion

In this paper, we provide a straightforward proof of an important, but nevertheless little known, result that was published in 1982 by Lindley [4] in the framework of subjective probability theory. Lindley's result, when reformulated in the machine learning/pattern recognition context, puts new lights on the probabilistic interpretation of the outputs of a trained classifier.

Roughly speaking, it says that, when training a classification model by minimizing a cost function, it is always possible to map the output of the model to the Bayesian a posteriori probabilities of the classes.

However, we must keep in mind that the results obtained in this paper are only valid if

- A minimum of the criterion is indeed reached after training, and
- The neural network is a 'sufficiently powerful model' that is able to approximate the optimal estimator to any degree of accuracy (perfect model matching).

Notice also that the results presented here are essentially asymptotic, and that issues regarding estimation from finite data sets are not addressed.

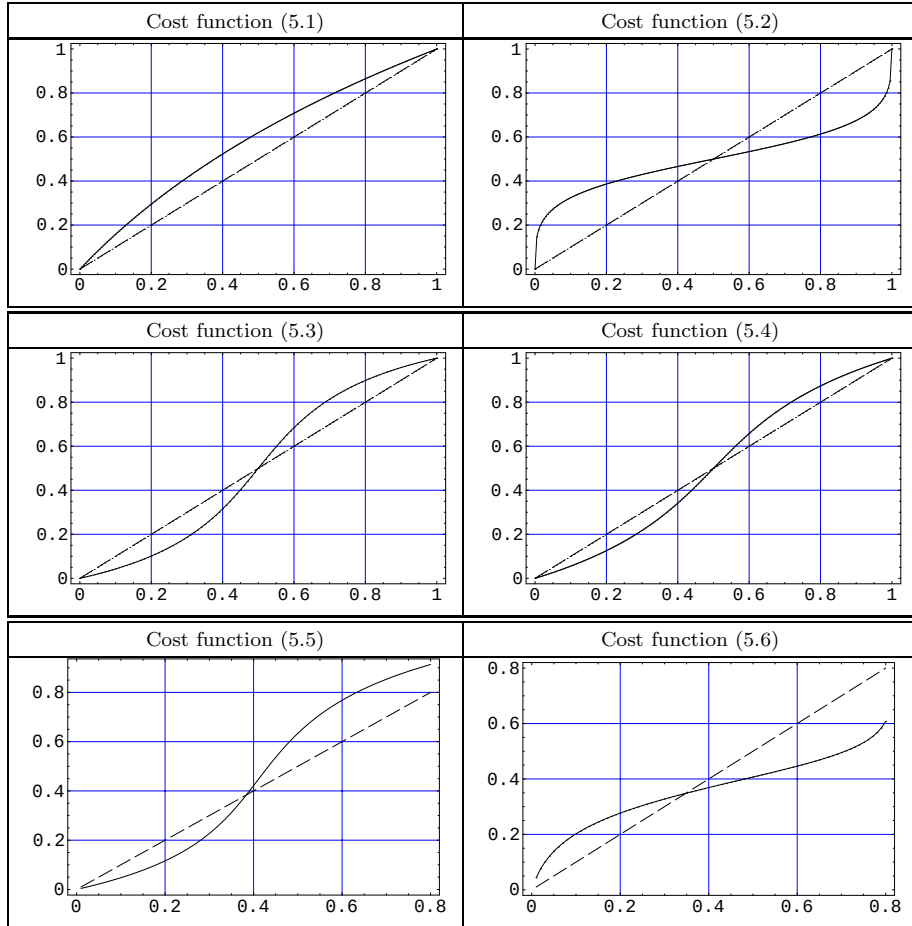


Table 5.3: Graphs of the optimal output obtained after minimization of the criterion before remapping (y-axis, plain line) and after remapping the output (y-axis, dash line), plotted in relation with different a posteriori probabilities (x-axis), for the six different cost functions. The applied mappings are shown in the table (5.3). The first four graphs are for binary models; the last two are for 3-output models. We clearly observe that the transformation maps the output of the model on the a posteriori probability of the class.

Acknowledgments

This work was partially supported by the project RBC-BR 216/4041 from the 'Région de Bruxelles-Capitale', and funding from the SmalS-MvM. Patrice Latinne is supported by a grant under an ARC (Action de Recherche Concertée) program of the Communauté Française de Belgique. We also thank the two anonymous reviewers for their pertinent and constructive remarks.

References

- [1] Bishop C. (1995). "Neural networks for pattern recognition". Oxford University Press.
- [2] Fomby T., Carter Hill R. & Johnson S. (1984). "Advanced econometric methods". Springer-Verlag.
- [3] Hampshire J.B. & Pearlmutter B. (1990). "Equivalence proofs for multi-layer perceptron classifiers and the Bayesian discriminant function". In *Proceedings of the 1990 Connectionist Models Summer School*, Touretzky D., Elman J., Sejnowski T. & Hinton G. (editors), Morgan Kaufmann, pp. 159-172.
- [4] Lindley D. (1982). "Scoring rules and the inevitability of probability (with discussions)". *International Statistical Review*, 50, pp. 1-26.
- [5] Richard M.D. & Lippmann R.P. (1991). "Neural network classifiers estimate Bayesian a posteriori probabilities". *Neural Computation*, 3, pp. 461-483.
- [6] McCullagh P. & Nelder J.A. (1990) "Generalized linear models, 2nd ed". Chapman and Hall.
- [7] Miller J.W., Goodman R. & Smyth P. (1991). "Objective functions for probability estimation". *Proceedings of the IEEE International Joint Conference on Neural Networks*, San Diego, pp. I-881-886.
- [8] Miller J.W., Goodman R. & Smyth P. (1993). "On loss functions which minimize to conditional expected values and posterior probabilities". *IEEE Transactions on Information Theory*, IT-39 (4), pp. 1404-1408.
- [9] Saerens M. (1996). "Non mean square error criteria for the training of learning machines". *Proceedings of the 13th International Conference on Machine Learning (ICML)*, July 1996, Bari (Italy), pp. 427-434.
- [10] Saerens M. (2000). "Building cost functions minimizing to some summary statistics". *IEEE Transactions on Neural Networks*, NN-11 (6), pp. 1263-1271.

- [11] Wolfram S. (1999). "*The Mathematica Book, 4th ed.*". Wolfram Media & Cambridge University Press.
-

APPENDIX: PROOF OF THE MAIN RESULTS

A. Appendix: If the model is trained by optimizing $C[\hat{y}|\mathbf{x}]$ (equation 3.4), and if there exists a mapping that transforms the output of the model $\hat{y}^*$ to the a posteriori probabilities (equation 3.8), then this mapping is provided by (3.6)

Let us recall the different hypothesis. After training, the criterion attains its optimal value at $\hat{y}^*(\mathbf{x})$. Thus, from (2.7) and (3.4), we obtain

$$\begin{aligned} \left. \frac{\partial C[\hat{y}|\mathbf{x}]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} &= \left. \frac{\partial \mathcal{L}[\hat{y}; 0]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} p(y=0|\mathbf{x}) \\ &+ \left. \frac{\partial \mathcal{L}[\hat{y}; 1]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} p(y=1|\mathbf{x}) = 0 \end{aligned} \quad (\text{A.1})$$

Moreover, let us suppose that there exists a mapping that transforms the optimal output $\hat{y}^*(\mathbf{x})$ to the a posteriori probabilities:

$$f(\hat{y}^*) = p(y=1|\mathbf{x}) = p(\omega_1|\mathbf{x}) \quad (\text{A.2})$$

with

$$p(y=0|\mathbf{x}) + p(y=1|\mathbf{x}) = 1 \quad (\text{A.3})$$

By developing (A.1) and using (A.2)–(A.3), we easily obtain

$$\left. \frac{\partial \mathcal{L}[\hat{y}; 0]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} (1 - f(\hat{y}^*)) + \left. \frac{\partial \mathcal{L}[\hat{y}; 1]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} f(\hat{y}^*) = 0 \quad (\text{A.4})$$

from which we compute $f(\hat{y}^*)$

$$f(\hat{y}^*) = \frac{\mathcal{L}'[\hat{y}^*; 0]}{\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]} \quad (\text{A.5})$$

where $\mathcal{L}'[\hat{y}^*; 0] = \partial \mathcal{L}[\hat{y}; 0] / \partial \hat{y}|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$ and $\mathcal{L}'[\hat{y}^*; 1] = \partial \mathcal{L}[\hat{y}; 1] / \partial \hat{y}|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$.

This shows that if there exists a mapping to the a posteriori probabilities, this mapping is provided by (3.6). ■

B. Appendix: If the model is trained by optimizing $C[\hat{y}|\mathbf{x}]$ (equation 3.4), and we transform the model's output $\hat{y}^*$ by (3.6), then the result of the mapping is the a posteriori probability defined by (3.8)

As in appendix A, let us consider a trained model (equation (3.4) is verified). From (2.7),

$$\left. \frac{\partial C[\hat{y}|\mathbf{x}]}{\partial \hat{y}} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} = \mathcal{L}'[\hat{y}^*; 0] p(y=0|\mathbf{x}) + \mathcal{L}'[\hat{y}^*; 1] p(y=1|\mathbf{x}) = 0 \quad (\text{B.1})$$

where $\mathcal{L}'[\hat{y}^*; 0] = \partial \mathcal{L}[\hat{y}; 0] / \partial \hat{y}|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$ and $\mathcal{L}'[\hat{y}^*; 1] = \partial \mathcal{L}[\hat{y}; 1] / \partial \hat{y}|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$. From (B.1), we obtain

$$\frac{\mathcal{L}'[\hat{y}^*; 0]}{\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]} = p(y=1|\mathbf{x}) \quad (\text{B.2})$$

If we apply the mapping

$$f(\hat{y}^*) = \frac{\mathcal{L}'[\hat{y}^*; 0]}{\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]} \quad (\text{B.3})$$

we find

$$f(\hat{y}^*) = p(y=1|\mathbf{x}) \quad (\text{B.4})$$

Since we require that the cost function is twice differentiable (3.1), the mapping (B.3) always exists; it transforms the optimal output $\hat{y}^*$ to the a posteriori probability $p(y=1|\mathbf{x})$. $\blacksquare$

C. Appendix: A conditional criterion $C[\hat{y}|\mathbf{x}]$ (2.4) having only one global minimum (no local minimum) for every possible $p(y=1|\mathbf{x})$ is equivalent to a strictly monotonic increasing mapping $f(\hat{y}^*)$ (3.6)

Notice that the requirements on the cost function (3.1)–(3.3) do not guarantee that the criterion has only one global minimum (no local minimum). Let us consider that $C[\hat{y}|\mathbf{x}]$ is already optimized, and therefore (3.4) is verified. From appendix A and B, this means that the optimum of $C[\hat{y}|\mathbf{x}]$, denoted by $\hat{y}^*$, is such that

$$p(y=1|\mathbf{x}) = \frac{\mathcal{L}'[\hat{y}^*; 0]}{\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]} \quad (\text{C.1})$$

In this appendix, we are interested in the second-order properties of the criterion. For $\hat{y}^*$ to be a minimum, the second-order condition (3.5) should be verified

in addition to (3.4). Let us compute the second-order derivative of $C[\hat{y}|\mathbf{x}]$. From (2.7), we have

$$\left. \frac{\partial^2 C[\hat{y}|\mathbf{x}]}{\partial \hat{y}^2} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} = (1 - p(y = 1|\mathbf{x})) \mathcal{L}''[\hat{y}^*; 0] + p(y = 1|\mathbf{x}) \mathcal{L}''[\hat{y}^*; 1] \quad (\text{C.2})$$

Where $\mathcal{L}''[\hat{y}^*; 0] = \partial^2 \mathcal{L}[\hat{y}; 0] / \partial \hat{y}^2|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$ and $\mathcal{L}''[\hat{y}^*; 1] = \partial^2 \mathcal{L}[\hat{y}; 1] / \partial \hat{y}^2|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})}$. Since we are at an optimum of $C[\hat{y}|\mathbf{x}]$, we can substitute $p(y = 1|\mathbf{x})$ by (C.1) in (C.2). We obtain

$$\left. \frac{\partial^2 C[\hat{y}|\mathbf{x}]}{\partial \hat{y}^2} \right|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} = \frac{\mathcal{L}'[\hat{y}^*; 0] \mathcal{L}''[\hat{y}^*; 1] - \mathcal{L}'[\hat{y}^*; 1] \mathcal{L}''[\hat{y}^*; 0]}{(\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1])^2} \quad (\text{C.3})$$

Now, let us also compute the first derivative of the mapping $f(\hat{y}^*)$ (equation (3.6))

$$\frac{\partial f(\hat{y}^*)}{\partial \hat{y}^*} = \frac{\mathcal{L}'[\hat{y}^*; 0] \mathcal{L}''[\hat{y}^*; 1] - \mathcal{L}'[\hat{y}^*; 1] \mathcal{L}''[\hat{y}^*; 0]}{(\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1])^2} \quad (\text{C.4})$$

Since $\hat{y}^* \in [0, 1]$, from (3.2), $(\mathcal{L}'[\hat{y}^*; 0] - \mathcal{L}'[\hat{y}^*; 1]) > 0$. Therefore, by comparing (C.3) and (C.4), we observe that $\partial^2 C[\hat{y}|\mathbf{x}] / \partial \hat{y}^2|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} > 0$ is equivalent to $\partial f(\hat{y}^*) / \partial \hat{y}^* > 0$ (the numerators are the same and the denominators are positive). This means that a monotonic increasing $f(\hat{y}^*)$ ($\partial f(\hat{y}^*) / \partial \hat{y}^* > 0$) for every $\hat{y}^* \in [0, 1]$ is equivalent to the fact that the conditional criterion is a minimum at every $\hat{y}^* \in [0, 1]$ ($\partial^2 C[\hat{y}|\mathbf{x}] / \partial \hat{y}^2|_{\hat{y}(\mathbf{x})=\hat{y}^*(\mathbf{x})} > 0$).

Let us now examine what happens when the conditional criterion has a local minimum. Suppose that $C[\hat{y}|\mathbf{x}]$ has two minima at $\hat{y}^*$ and $\hat{y}'^*$ (with $\hat{y}^* < \hat{y}'^*$), for the same $p(y = 1|\mathbf{x})$. In this case, since $C[\hat{y}|\mathbf{x}]$ is differentiable, it must pass through a maximum $\hat{y}_{\max}$ located between the two minima $\hat{y}^* < \hat{y}_{\max} < \hat{y}'^*$. For this maximum, we have $\partial^2 C[\hat{y}|\mathbf{x}] / \partial \hat{y}^2|_{\hat{y}=\hat{y}_{\max}} < 0$ which is equivalent to $\partial f(\hat{y}) / \partial \hat{y}|_{\hat{y}=\hat{y}_{\max}} < 0$, and therefore a decreasing $f(\hat{y}^*)$ on some interval including $\hat{y}_{\max}$. This indicates that a decreasing $f(\hat{y}^*)$ on some interval is associated to local minima of the conditional criterion.

This shows that conditional criterion (2.4) having only one global minimum (no local minimum) for every possible $p(y = 1|\mathbf{x})$ is equivalent to a strictly monotonic increasing mapping (3.6). ■

D. Appendix: multi-output case. If the model is trained by optimizing $C[\hat{\mathbf{y}}|\mathbf{x}]$ (equation 4.5), and if there exists a mapping that transforms the output of the model $\hat{\mathbf{y}}^*$ to the a posteriori probabilities (equation 4.6), then this mapping is obtained by solving a system of $n - 1$ linear equations (4.7)

After training, the criterion attains its optimal value at $\hat{\mathbf{y}}^*(\mathbf{x})$. Thus, from (4.4) and (4.5), we obtain

$$\begin{aligned} \left. \frac{\partial C[\hat{\mathbf{y}}|\mathbf{x}]}{\partial \hat{y}_i} \right|_{\hat{\mathbf{y}}(\mathbf{x})=\hat{\mathbf{y}}^*(\mathbf{x})} &= \sum_{j=1}^n p(\mathbf{y} = \mathbf{e}_j|\mathbf{x}) \left. \frac{\partial \mathcal{L}[\hat{\mathbf{y}}; \mathbf{e}_j]}{\partial \hat{y}_i} \right|_{\hat{\mathbf{y}}(\mathbf{x})=\hat{\mathbf{y}}^*(\mathbf{x})} \\ &= 0, \text{ for } i = 1 \dots n - 1 \end{aligned} \quad (\text{D.1})$$

Moreover, let us suppose that there exists a mapping that transforms the optimal output vector $\hat{\mathbf{y}}^*(\mathbf{x})$ to the a posteriori probabilities:

$$f_i(\hat{\mathbf{y}}^*(\mathbf{x})) = p(\mathbf{y} = \mathbf{e}_i|\mathbf{x}) = p(\omega_i|\mathbf{x}) \quad (\text{D.2})$$

with

$$\sum_{i=1}^n f_i(\hat{\mathbf{y}}^*(\mathbf{x})) = \sum_{i=1}^n p(\mathbf{y} = \mathbf{e}_i|\mathbf{x}) = 1 \quad (\text{D.3})$$

By using (D.2) and (D.1), we easily obtain

$$\sum_{j=1}^n f_j(\hat{\mathbf{y}}^*) \left. \frac{\partial \mathcal{L}[\hat{\mathbf{y}}; \mathbf{e}_j]}{\partial \hat{y}_i} \right|_{\hat{\mathbf{y}}(\mathbf{x})=\hat{\mathbf{y}}^*(\mathbf{x})} = 0, \text{ for } i = 1 \dots n - 1 \quad (\text{D.4})$$

Let us define $\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_j] = \partial \mathcal{L}[\hat{\mathbf{y}}; \mathbf{e}_j] / \partial \hat{y}_i|_{\hat{\mathbf{y}}(\mathbf{x})=\hat{\mathbf{y}}^*(\mathbf{x})}$. By further using (D.3), we rewrite (D.4) as

$$\sum_{j=1}^{n-1} f_j(\hat{\mathbf{y}}^*) \mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_j] + \left[1 - \sum_{j=1}^{n-1} f_j(\hat{\mathbf{y}}^*) \right] \mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_n] = 0, \text{ for } i = 1 \dots n - 1 \quad (\text{D.5})$$

By rearranging the terms, we obtain

$$\sum_{j=1}^{n-1} [\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_n] - \mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_j]] f_j(\hat{\mathbf{y}}^*) = \mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_n], \text{ for } i = 1 \dots n - 1 \quad (\text{D.6})$$

Or equivalently

$$\sum_{j=1}^{n-1} \left[1 - \frac{\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_j]}{\mathcal{L}'_i[\hat{\mathbf{y}}^*; \mathbf{e}_n]} \right] f_j(\hat{\mathbf{y}}^*) = 1, \text{ for } i = 1 \dots n - 1 \quad (\text{D.7})$$

This shows that if there exists a mapping to the a posteriori probabilities, this mapping is provided by solving (4.7).

However, for such general cost function definitions, it is difficult to assess if this solution exists and if it is indeed a minimum. ■
