

No hay necesidad de usar sustitutos ad-hoc: el costo esperado es una métrica de clasificación basada en principios y de propósito general.

Luciana Ferrer

Instituto de Ciencias de la Computación, CONICET-UBA, Argentina,
ORCID:0000-0002-0426-8683, lferrer@dc.uba.ar

Resumen. El costo esperado (EC, por sus siglas en inglés) es una de las principales métricas de clasificación introducidas en libros de estadística y aprendizaje automático. Se basa en el supuesto de que, para una aplicación de interés determinada, cada decisión tomada por el sistema tiene un costo correspondiente que depende de la clase verdadera de la muestra. Una métrica de evaluación puede entonces definirse tomando la esperanza del costo sobre los datos. Dos casos especiales del EC son ampliamente utilizados en la literatura de aprendizaje automático: la tasa de error y la tasa de error balanceada. Otras instancias del EC pueden ser útiles para aplicaciones en las que algunos tipos de errores son más graves que otros, o cuando las probabilidades a priori de las clases difieren entre los datos de evaluación y el escenario de uso. Sorprendentemente, la forma general del EC rara vez se utiliza en la literatura de aprendizaje automático. En su lugar, se emplean métricas alternativas ad-hoc como el score F-beta y el coeficiente de correlación de Matthews (MCC). En este trabajo, argumentamos que el EC es superior a estas métricas alternativas, ya que es más general, interpretable y adaptable a cualquier escenario de aplicación. Proporcionamos tanto discusiones con base teórica como ejemplos para ilustrar el comportamiento de las distintas métricas.

Palabras clave: Costo esperado, métricas de clasificación, score F-beta

No need for ad-hoc substitutes: the expected cost is a principled all-purpose classification metric

Abstract. The expected cost (EC) is one of the main classification metrics introduced in statistical and machine learning books. It is based on the assumption that, for a given application of interest, each decision made by the system has a corresponding cost which depends on the true class of the sample. An evaluation metric can then be defined by taking the expectation of the cost over the data. Two special cases of the EC

are widely used in the machine learning literature: the error rate (one minus the accuracy) and the balanced error rate (one minus the balanced accuracy or unweighted average recall). Other instances of the EC can be useful for applications in which some types of errors are more severe than others, or when the prior probabilities of the classes differ between the evaluation data and the use-case scenario. Surprisingly, the general form for the EC is rarely used in the machine learning literature. Instead, alternative ad-hoc metrics like the F-beta score and the Matthews correlation coefficient (MCC) are used for many applications. In this work, we argue that the EC is superior to these alternative metrics, being more general, interpretable, and adaptable to any application scenario. We provide both theoretically-motivated discussions as well as examples to illustrate the behavior of the different metrics.

Keywords: Expected cost, classification metrics, F-beta score

DOI: El journal no genera números DOI.

Publicado en: Transactions on Machine Learning Research, 2025

Factor de impacto: No disponible aún.