



A review of efficiency criteria suitable for evaluating low-flow simulations

Raji Pushpalatha^a, Charles Perrin^{a,*}, Nicolas Le Moine^b, Vazken Andréassian^a

^a Cemagref, UR HBAN, 1 rue Pierre-Gilles de Gennes, CS 10030, 92761 Antony Cedex, France

^b Université Pierre et Marie Curie, UMR 7619 Sisyphe, 4 Place Jussieu, 75252 Paris Cedex 05, France

ARTICLE INFO

Article history:

Received 5 May 2011

Received in revised form 27 September 2011

Accepted 29 November 2011

Available online 8 December 2011

This manuscript was handled by Andras Bardossy, Editor-in-Chief, with the assistance of Ezio Todini, Associate Editor

Keywords:

Low flows

Simulation efficiency

Model evaluation

Benchmark model

Criteria

SUMMARY

Low flows are seasonal phenomena and an integral component of the flow regime of any river. Because of increased competition between water uses, the demand for forecasts of low-flow periods is rising. But how low-flow predictions should be evaluated? This article focuses on the criteria able to evaluate the efficiency of hydrological models in simulating low flows. Indeed, a variety of criteria have been proposed, but their suitability for the evaluation of low-flow simulations has not been systematically assessed.

Here a range of efficiency criteria advised for low flows is analysed. The analysis mainly concentrates on criteria computed on continuous simulations that include all model errors. The criteria were evaluated using two rainfall–runoff models and a set of 940 catchments located throughout France. In order to evaluate the capacity of each criterion to discriminate low-flow errors specifically, we looked for the part of the hydrograph that carries most of the weight in the criterion computation.

Contrary to what was expected, our analysis revealed that, in most of the existing criteria advised for low flows, high flows still make a significant contribution to the criterion's value. We therefore recommend using the Nash–Sutcliffe efficiency criterion calculated on inverse flow values, a valuable alternative to the classically used criteria, in that on average it allows focusing on the lowest 20% of flows over the study period.

© 2011 Elsevier B.V. All rights reserved.

1. Introduction

1.1. In the jungle of efficiency criteria

Hydrological modelling aims at understanding and interpreting catchment hydrological behaviour. It is also used to address a number of practical issues, ranging from flood estimation to water resources management and low-flow forecasting. Whatever model is applied, the model user needs appropriate and meaningful indicators informing on the actual capacity of the model.

However, the evaluation of goodness-of-fit is not as straightforward as it may seem at first glance. Of course, model performance can first be evaluated by the visual comparison of the observed and simulated flow hydrographs, but this remains extremely dependent on the evaluator's experience (Chiew and McMahon, 1993; Houghton-Carr, 1999). A more objective way to evaluate model performance is to use numerical criteria, but then the user may get lost in a jungle of potential criteria. Why is the choice so difficult? Several reasons can be put forward:

1. flows vary by several orders of magnitude that may not be equally useful for the modeller;

2. hydrological models often produce heteroscedastic errors, i.e. their variance is not independent of the flow value;
3. the range of target flows may vary significantly between evaluation periods;
4. the model may be used for different applications, which may require specific criteria.

For these reasons, a large variety of criteria have been proposed and used over the years in hydrological modelling, as shown for example by the lists of criteria given by the ASCE (1993), Dawson et al. (2007), Moriasi et al. (2007) and Reusser et al. (2009). Among these criteria, some are *absolute* criteria such as the widely used root mean square error, while others are *relative* criteria (i.e. normalized) such as the Nash and Sutcliffe (1970) efficiency criterion (NSE). In the latter, model errors are compared to the errors of a reference or benchmark model (Seibert, 2001; Perrin et al., 2006). This provides a useful quantification of model performance in that it indicates to which extent the model is better (or worse) than the benchmark. It also facilitates the comparison of performance between catchments.

However, the choice of a benchmark is difficult: different benchmarks are more or less demanding (and thus the comparison more or less informative) depending on the type of hydrological regime or the type of model application. This sometimes makes it difficult to interpret the relative criteria and a bit puzzling for an inexperienced

* Corresponding author. Tel.: +33 140966086; fax: +33 140966199.

E-mail address: charles.perrin@irstea.fr (C. Perrin).

Download English Version:

<https://daneshyari.com/en/article/4577163>

Download Persian Version:

<https://daneshyari.com/article/4577163>

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