

Metrics that Matter: Objective Functions and Their Impact on Signature Representation in Conceptual Hydrological Models

Peter Wagener^{1,2}, Wouter Knoben², Niels Schütze¹, and Diana Spieler^{1,2}

¹Institute of Hydrology and Meteorology, TUD Dresden University of Technology, Dresden, Germany

²Schulich School of Engineering, Department of Civil Engineering, University of Calgary, Canada

Correspondence: Peter Wagener (peter.wagener@ucalgary.ca)

Abstract. Although objective functions (OFs) are widely discussed in the literature, many modelling studies still default to a few common metrics, without much consideration of their relative strengths and weaknesses. This paper systematically investigates the impact of OF choice on the representation of various streamflow characteristics across 47 conceptual models and 10 hydro-climatically diverse catchments selected from the CARAVAN dataset. We use eight different OFs for calibration,

5 including the Kling–Gupta efficiency (KGE), Nash–Sutcliffe efficiency (NSE), ~~and their respective logarithmic variants, as well as a variant of each that uses a streamflow transformation, and~~ four more recently proposed metrics. We further account for parameter uncertainty by retaining the top 500 parameter sets for each combination of model, basin and OF. We evaluate the representation of 15 hydrological signatures that capture a relevant selection of streamflow characteristics to determine generalizable strengths and weaknesses of individual OFs across different models and ~~catchments. Results show that the choice~~
10 ~~of OF can significantly affect a model's capability to simulate different hydrological signatures such as runoff ratios, extreme flow percentiles, and certain baseflow characteristics~~ hydroclimatic conditions. A random forest importance analysis showed that OF choice is often more influential than model structure choice in determining signature representation, with catchment characteristics remaining the dominant control. While certain signatures, particularly those related to flow variability, are relatively insensitive to OF choice, ~~others~~ other signatures exhibit large performance shifts across different OFs. ~~Generally As~~
15 ~~other studies have already noted,~~ no single OF simultaneously achieved high performance across all tested signatures, highlighting that a single-objective calibration is unlikely to lead to an all-purpose model. Our results reinforce calls to choose objective functions deliberately and in line with the objectives of a study. They also provide initial performs best across all signatures. Our consistent multi-model and multi-catchment experiment now provides generalizable guidance on which ~~metrics highlight particular facets~~ OF best reproduces which aspect of streamflow behaviour. This helps to better select objective functions most
20 appropriate for a specific modelling purpose and provides insights into which metric combinations can cover a broader range of flow characteristics.

1 Introduction

Setting up and running a hydrological model requires modellers to make many decisions. These ~~typically may~~ include the choice of a suitable model (structure), sensible parameter boundaries, an appropriate calibration period, and which method to

25 use for forcing data interpolation. Selecting which performance metric(s) to employ as the objective function during model calibration is another critical decision. ~~Different choices of~~ where research has shown that different performance metrics can lead to substantially different model outcomes. ~~For example, Mai (2023)~~ Mai (2023), for example, tested six alternative ~~metrics for calibration~~ calibration metrics and found marked differences in the resulting hydrographs. ~~Similarly, Garcia et al. (2017)~~ Garcia et al. (2017) compared variants of the KGE metric for low-flow simulation, observing large differences in annual runoff

30 estimates. Studies such as ~~Pool et al. (2017) and Vis et al. (2015)~~ Pool et al. (2017) and Vis et al. (2015) further demonstrated that the choice of performance metric for calibration influences how well the model reproduces key hydrological signatures (i.e., streamflow characteristics). Similarly, Mendoza et al. (2016) and Seiller et al. (2017) show the impact calibration choices can have on the portrayal of climate change impacts.

35 ~~Among~~ The choice of calibration metric, therefore, plays a crucial role among the various decisions involved in model calibration, evaluation, and diagnostic analysis; ~~the choice of objective function is often critical.~~ It directly determines how model parameters are optimized and, consequently, how well the resulting simulations reproduce both observed data and underlying hydrological processes. Despite the well-established influence of objective function selection, and an extensive body of research discussing the individual strengths and weaknesses of various performance metrics (e.g. Thirel et al., 2024; Althoff

40 and Rodrigues, 2021; Cinkus et al., 2023; Knoben et al., 2019b; Mizukami et al., 2019; Bennett et al., 2013; Gupta et al., 2009; Krause et al., 2005), Jackson et al. (2019) conclude that: “in most hydrologic modelling studies error metrics are chosen based on familiarity, [and] without consideration of the relative strengths and weaknesses” in their review of more than 60 different performance metrics. And indeed, a quick and informal review of 60 modelling studies in the existing literature library of the authors ~~–~~ supports this claim by showing that most studies provide little to no reasoning for their choice of objective function

45 (Figure 1; details on review in Supplement S1). The papers that gave reasoning mainly referred to general characteristics and perceptions of their metric of choice. Reasons we documented (see Supplement S1) can roughly be classified into various forms of “the KGE is a more balanced metric than the NSE”, “NSE is a common metric for high flows”, “logNSE/logKGE focus on low flows”, “NSE/KGE are widely used metrics” or “we use this metric because it is comparable with previous studies”.

50 We believe that part of the prevailing tendency to define objective functions in an often ad-hoc manner is the broad but scattered and often site- and/or model-specific research on the impacts of different performance metrics used as objective functions. Most good modelling practice guidelines (e.g. Wavereen et al., 1999; Jakeman et al., 2006, 2018) simply advise to select an appropriate objective function that aligns with the purpose of the modelling study, but do not provide guidance on how to do so. The question of how to link-connect the choice for a suitable calibration metric to a specific modelling purpose thus

55 remains unclear (Jacobs et al., 2024). While several ~~authors studies~~ have investigated connections between objective function choice and how the calibrated models represent certain streamflow characteristics, they are often limited in the number of models they use (1 model in e.g. Hallouin et al., 2020; Melsen et al., 2019; Pool et al., 2017; Vis et al., 2015) or signatures they consider (less than 5 in e.g. Araya et al., 2023; Mizukami et al., 2019; Seiller et al., 2017). As Table 1 highlights, studies with a similar focus to our study also mostly consider catchments of only one specific region or country and no study the

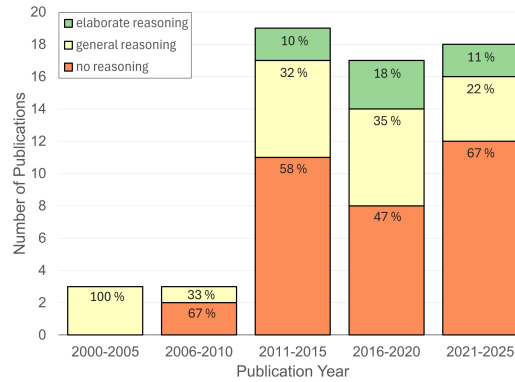


Figure 1. a) Number of analyzed studies published during different five-year periods b) Percentage and the percentage of these studies that give no - general - or elaborate reasoning on their choice of objective function. This data stems from an informal/non-exhaustive literature review of 60 studies taken from the authors' literature database. Details regarding the review are given in Supplement S1.

60 authors are aware of tested more than 12 models while also considering several metrics and signatures. ~~Consequently, few findings generalize beyond the familiar dictum that squared metrics emphasize the larger errors commonly associated with higher flows, and transformations help to emphasize low flows. The available site or model specific knowledge make a robust transfer to new studies uncertain and generalizable trade-offs of specific metric choices seem underdiscussed.~~

65 Some of the existing studies show recurring tendencies across regions and models. For instance, Pool et al. (2017) and Hallouin et al. (2020) found that traditional objective functions such as NSE or KGE can perform as well or better than tailored signature-based metrics for predicting flow characteristics not explicitly included in calibration. Mizukami et al. (2019) similarly noted that application-specific metrics may excel at their target but can degrade performance on other metrics. Generally, most studies can identify a preferred metric or metric combination. However, that selection is strongly conditioned to the modelling
 70 purpose, the flow conditions of interest, and the hydroclimatic setting of the individual study (e.g. seasonal forecasting in snow-influenced catchments for Araya et al. (2023) or tropical low-flow applications for Althoff and Rodrigues (2021)). A robust transfer of findings to new contexts is therefore difficult and further complicated by inconsistent experimental designs across studies (differing catchment sets, model selections, and signature choices), which prevent a straightforward synthesis on how metric choices influence signature representation.

75

~~Building on this motivation~~ To address this problem, this study aims to develop a more generalized understanding of how the choice of performance ~~metries-metric~~ used for model calibration influences the representation of the hydrological ~~regime-~~ behavior. We focus on single-objective calibration because it remains the calibration standard and can provide insights into metric combinations that might serve as complementary multi-objective calibration targets. By systematically analyzing catchments
 80 ~~spanning that span~~ a broad range of climatic conditions and employing a large set of conceptual model structures, we seek to

Publication	Main Focus	Catchments	Location	Models	OFs	Signatures
Althoff and Rodrigues (2021)	OF impact on streamflow simulations	179	Brazil	1	11*	7
Araya et al. (2023)	OF impact on streamflow forecast	22	Chile	3	12*	5
Hallouin et al. (2020)	OF impact on signatures	33	Ireland	1	6*	18
Hernandez-Suarez et al. (2018)	OF impact on signatures	1	ACF-Basin, USA	1	8*	167
Melsen et al. (2019)	calibration impact on flood/drought	9	Thur Basin, Switzerland	1	2	6
Mendoza et al. (2016)	calibration impact on climate projections	3	Colorado Basin, USA	4	4	6
Mizukami et al. (2019)	OF impact on high flows	492	USA	2	5	2
Pool et al. (2017)	OF impact on signatures	25	Tennessee, USA	1	29*	13
Seiller et al. (2017)	OF impact on climate projections	37	Quebec, Canada	12	3	4
Vis et al. (2015)	OF impact on signatures	27	Tennessee, USA	1	9*	12
This study	OF impact on signatures	10	worldwide	47	8	15

Summary

of similar studies that investigate the impact of objective function choice on the representation of hydrological signatures. We compare the number of investigated models, catchments, metrics and signatures to this study. Note that studies with an *asterisk also tested combinations of different metrics (i.e. composite criteria or multi-metric objective functions) in addition to individual calibration metrics.

Table 1. Summary of similar studies that investigate the impact of objective function choice on the representation of hydrological signatures. We compare the number of investigated models, catchments, metrics and signatures to this study. Note that studies with an *asterisk also tested combinations of different metrics (i.e. composite criteria or multi-metric objective functions) in addition to individual calibration metrics.

identify general benefits and limitations associated with different calibration metrics across diverse hydrological contexts. The selected set of catchments, though limited in number, was chosen based on differing climate indices ~~such as~~ (aridity, seasonality, and fraction of snow, ~~ensuring~~), to ensure a representative spread of hydroclimatic conditions (cf. the 12 MOPEX basins selected in Duan et al. (2006); van Werkhoven et al. (2008); Carrillo et al. (2011)). ~~This design~~. This enables us to examine

85 how models calibrated with different performance metrics behave under varying hydrological regimes while also balancing depth with breadth of analysis (Gupta et al., 2014). ~~In this study, we~~ We evaluate how 47 conceptual model structures, each calibrated using alternative performance metrics and an ensemble of high-performing parameter sets to account for parameter uncertainty, reproduce 15 selected hydrological signatures that describe distinct aspects of the flow regime, such as flow variability, baseflow contribution, and seasonal timing runoff dynamics. This multi-dimensional setup ~~— combining diverse~~
90 ~~catchments, model structures, and metrics~~ — contributes to identifying generalizable strengths and weaknesses of individual performance metrics that extend beyond their known sensitivities to particular flow conditions, offering new insights into how calibration choices influence process realism in conceptual hydrological modelling. ~~We will uncover potential drawbacks of specific metric choices and provide a more comprehensive guideline for modellers to decide which objective function may be most appropriate for their specific modelling goal.~~ The following sections will introduce the catchments, models and signatures
95 we used (Section 2), present ~~and discuss~~ the results (Section 3) and provide an overarching discussion of the implications metric choice can have on the representation of the hydrological regime behavior and well as limitations and future steps (Section 4). Conclusions are presented in Section 5.

2 Methodology

Figure 2 gives an overview of the methodology used in this study. We calibrate 47 conceptual hydrological models from the ~~MARRMoT-Toolbox~~ Modular Assessment of Rainfall-Runoff Models Toolbox (MARRMoT) (Knoben et al., 2019a; Trotter et al., 2022) using eight different calibration metrics. The models and metrics used are introduced in Section 2.1.1 and 2.1.2 respectively. We calibrate all of these models (as described in Section 2.2) for a subset of 10 hydro-climatically differing catchments from the CARAVAN dataset (Kratzert et al., 2023) as introduced in Section 2.1.3. For each calibration run, we retain the top 500 parameter sets for further investigation. After selecting only the well-performing model structures according to a benchmark procedure (Section 2.3) we use the simulated discharge timeseries to calculate 15 selected signatures that represent varying aspects of the hydrological ~~regime-behavior~~ (as introduced in Section 2.4). Through a comparison of the simulated and observed signature values we are able to identify the strengths and weaknesses of the tested metrics in representing the hydrological regime over a broad range of hydro-climatic conditions and different model complexities.

2.1 Experiment Components

2.1.1 Models

We obtained the 47 models used in this study from the MARRMoT toolbox (Knoben et al., 2019a; Trotter and Knoben, 2022), ~~partly~~ because the toolbox allows easy and consistent use of the different models. These models are based on the scientific literature and mimic some widely used and well-known models such as HBV, GR4J, TOPMODEL, VIC and HYMOD. A full list of all models included in this analysis can be found in the Supplementary Material S2.1. Of note, only 12 of the 47 models possess a snow module. Details of these modules vary depending on the model under consideration. Figure 2b shows how the model structures differ in their number of considered stores (i.e., state variables) and parameters and therefore gives an idea of the different models' complexity. The number of parameters ranges from one to 24 and the number of considered stores ranges from one to eight. The simplest model has one store and one parameter, whereas the most complex models have eight stores and 12 parameters, or three stores and 24 parameters. By using such a wide range of different complexities in model structures, we ensure that the strengths and weaknesses identified generalize across different model structures.

2.1.2 Objective Functions

We calibrate the 47 hydrological models with eight objective functions ~~each using different performance metrics, each based on a different performance metric~~, and analyze how well the calibrated models simulate 15 streamflow signatures. We differentiate between model accuracy (~~objective function value~~), expressed by the objective-function value, and model adequacy (~~signature representation~~) and aim to investigate their relationship. We selected four widely known and frequently used metrics and four metrics that have only recently been introduced in attempts to improve on known weaknesses of the more commonly used metrics, expressed by signature representation. This selection of eight metrics is necessarily non-exhaustive, many further metrics and transformations exist, but was chosen to span widely used and recently proposed formulations while keeping

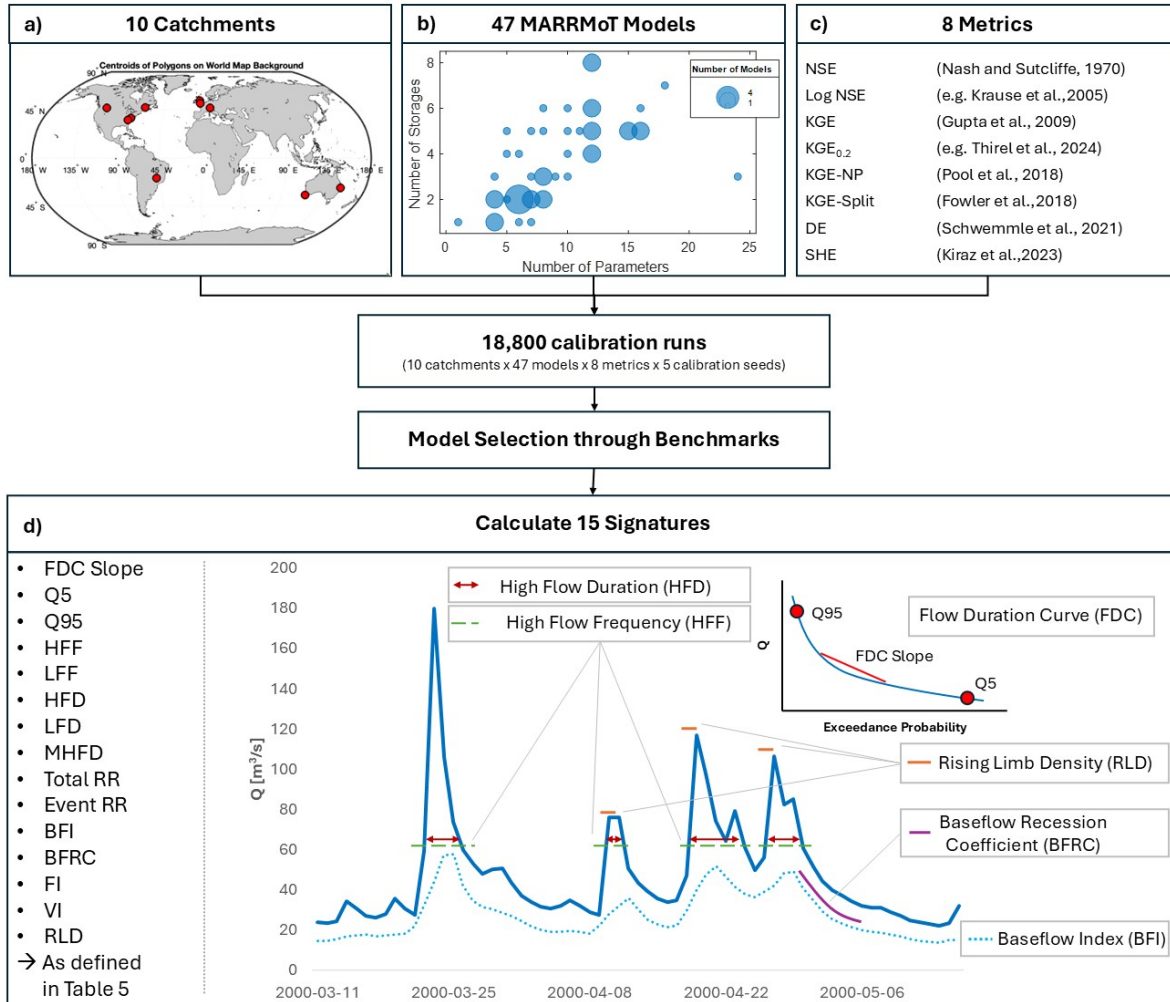


Figure 2. Overview of the experiment design of this study. a) Shows the geographical location of the study catchments. b) Visualizes the different complexities (number of storages vs number of parameters) of the investigated MARRMoT Models. c) Lists the eight objective functions used for calibration. d) Lists the 15 hydrological signatures investigated and shows a schematic visualization of some of them.

the analysis tractable. In general, "performance metrics" can be seen as any metric that quantifies performance metrics are numerical measures used to quantify model performance, while "objective function" indicates a more specific use of any performance metric as a target whereas objective functions refer to the use of such metrics as optimization targets during model calibration. In the following, we use these terms interchangeably. In this study, the eight different objective functions we use are defined through the eight individual performance metrics introduced below. We therefore use 'objective function' or 'OF' when referring to their role in calibration, and 'metric' when referring to their mathematical formulation or diagnostic properties. To describe them, we introduce the following set of common symbols: Q and P denote streamflow and

precipitation respectively. σ and r are standard deviation and correlation. Subscripts are used to indicate the specifics of all mentioned variables, e.g. Q_s and Q_o show simulated and observed streamflow respectively. Subscript t refers to individual time steps. For correlation, the subscripts r_{pe} and r_{sp} are used for Pearson and Spearman correlation. Overlines are used to indicate temporal means. In sums, the letter n is used to indicate the total length of the time series.

140 The first two metrics are likely the most common objective functions used in hydrology (Melsen et al., 2025). That is, (1) the ~~Nash-Sutcliffe-Efficiency (NSE; Nash and Sutcliffe, 1970)~~ Kling-Gupta-Efficiency (KGE; Gupta et al., 2009) and (2) the ~~Kling-Gupta-Efficiency (KGE; Gupta et al., 2009)~~ Nash-Sutcliffe-Efficiency (NSE; Nash and Sutcliffe, 1970) as defined in Table 2. ~~The next two metrics are just as well known, and~~ Additionally, we use transformations of the first two metrics that aim to improve the representation of low flows, ~~the logarithmic versions of NSE (~~ For KGE, we selected (3) and KGE (a
145 power transformation because of the known issues with logarithmic transformations (Thirel et al., 2024; Santos et al., 2018b). For NSE we selected (4) the logarithmic version because of its long-standing use (Krause et al., 2005); here ϵ is a small positive offset added to streamflow to avoid undefined logarithms at zero flow. In addition, we include four recently developed metrics that intend to improve certain aspects of NSE or KGE or diagnose hydrological model behaviour.

~~Firstly~~ First, this is (5) the non-parametric version of the KGE (KGE-NP; Pool et al., 2018). The major difference ~~to~~ with
150 the standard KGE is that the variability component is no longer based on the ratio of variances, but instead replaced by a flow-duration curve-based term. Additionally, the correlation term uses the Spearman rank correlation instead of the ~~pearson~~ Pearson correlation. Pool et al. (2018) argue that the resulting benefit of their proposed metric is an overall better agreement between observations and simulations, except for high flows. Their reasoning is that more information is contained within the metric because the FDC is more complex than the standard deviation and the Spearman rank correlation leads to improvements
155 for low flows. In the equation of (5) $I(k)$ and $J(k)$ indicate time steps in which the k -th largest flow occurs in the simulated and observed time series respectively.

~~Secondly~~ Second, another adaption of the KGE, ~~the KGE-Split (KGE-Split; Fowler et al., 2018)~~ (6) the KGE-Split (KGE-Split; Fowler et al., 2018) is evaluated. It is calculated like the regular KGE, but instead of calculating one value for the entire time series, a value for each year is calculated, and the mean over these values becomes the actual objective function value. Fowler et al. (2018) developed
160 it to put more emphasis on dry years, since the typically used "least squares" methods give more attention to high flows rather than low flows. The variables are identical to those used in the KGE, with y indicating the evaluated year and Y the total number of years as shown in equation (6).

~~Thirdly, the~~ Third, the (7) signature-based hydrologic efficiency (SHE; Kiraz et al., 2023) normalizes both the bias and the variance term by the mean and variance of the precipitation. Like the KGE-NP, it also uses the Spearman description for the
165 correlation term to improve on the low flow representations. The SHE as used in Kiraz et al. (2023) is defined in equation (7).

~~Fourthly and lastly~~ Fourth, we analyze (8) the Diagnostic Efficiency (DE; Schwemmler et al., 2021) where the bias term is replaced with a mean relative error and the variance term is the integral over all residuals of the relative error based on the flow duration curve. As the name suggests, this metric is designed as a diagnostic tool but will be applied here as an objective function. The general structure, with three parts for bias, variance, and timing, is similar to the KGE, but particularly
170 the variance term proposes an interesting alternative to the KGE. To convert it into a similar objective function as the other

candidates we use: $DE_{OF} = 1 - DE$ $DE_{OF} = 1 - DE$, the only metric for which we mark this calibration/analysis distinction in the notation since its two forms differ in formulation rather than in role alone. In equation (8), $\hat{p}$ denotes the exceedance probability of observed and simulated streamflow and $Q^{\downarrow}(p)$ the streamflow sorted along the flow-duration curve.

All of the eight used metrics can be disaggregated into three components representing bias, variability and correlation. For an easier overview of the differences between the individual metrics we summarized how they represent each component in Table 3.

2.1.3 Study Catchments

This study uses catchments from the CARAVAN database (Kratzert et al., 2023) which combines several large-sample datasets, include many CAMELS (Catchment Attributes and Meteorology for Large-sample Studies) datasets-version of different countries or regions in a standardized way. While the latest update to CARAVAN now includes more than 20000 different catchments (Färber et al., 2025) from over-seven-many existing large-sample hydrology datasets, at the beginning of this study CARAVAN included 2901 catchments. At the time, CARAVAN included catchments in the United States (Addor et al., 2017), Great Britain (Coxon et al., 2020), Brazil (Chagas et al., 2020), Chile (Alvarez-Garretón et al., 2018) and Australia (Fowler et al., 2021) from their respective CAMELS datasets, as well as the LamaH-CE (Central Europe, Klingler et al., 2021) and HYSETS (North America, Arsenault et al., 2020) datasets. CARAVAN provides standardized meteorological forcing, streamflow data, and static catchment attributes (e.g., geophysical, sociological, climatological) with a median data length of 31 years. Our primary goal is to investigate the sensitivity of a large number of different model structures to objective function choice. To keep the analysis manageable (i.e., to "balance depth with breadth" (Gupta et al., 2014)) (Gupta et al., 2014), we defined a subset of ten hydro-climatically differing catchments for our analysis. To determine these catchments, we used the climate classification procedure outlined in (Knoben et al., 2018) Knoben et al. (2018) to quantify the aridity, seasonality and fraction snow in each of the CARAVAN basins. We then used a k-means clustering algorithm to divide the catchments into 10 clusters, and selected the catchment closest to each cluster centroid for use in this work. The location of the selected catchments is shown in Figure 2a and a brief description of some catchment properties is given in Table 4. We also applied a threshold for maximum catchment area of 1000 km^2 to ensure a more direct influence between discharge, signatures and metrics, because for signatures like the Flashiness Index there might be a dampening effect in larger catchments. The full clustering procedure is described in more detail in the Supplementary Material S2.2.

Table 4. Overview of catchments properties for the selected 10 CARAVAN basins. [Details in Supplementary Material S2.2.](#)

Cluster	Abbrev.	Dataset	Area	Aridity <u>Moisture</u>	Seasonality	Snow Frac	Precip	Runoff	Mea
–	–	–	(km^2)	(–)	(–)	(–)	(mm/yr)	(mm/yr)	(–)
1	AUS1	CAMELS-AUS	125.74	-0.39	0.59	0.00	836.5 <u>837</u>	117.3 <u>117</u>	18.2
2	BR6	CAMELS-BR	194.00	-0.28	1.30	0.00	1451.2 <u>1451</u>	490.0 <u>490</u>	22.4
3	AUS6	CAMELS-AUS	124.94	-0.04	1.67	0.00	680.4 <u>680</u>	165.8 <u>166</u>	15.4
4	C12	CAMELS	148.69	0.02	1.66	0.49	773.5 <u>774</u>	301.5 <u>302</u>	2.5
5	C02	CAMELS	285.39	0.03	1.07	0.07	1120.1 <u>1120</u>	361.8 <u>362</u>	10.9
6	GB3	CAMELS-GB	136.53	0.16	1.46	0.00	792.6 <u>793</u>	185.7 <u>186</u>	9.7
7	C03	CAMELS	127.83	0.33	0.85	0.08	1411.5 <u>1412</u>	656.6 <u>657</u>	10.5
8	HYS	HYSETS	328.44	0.34	1.28	0.38	1211.3 <u>1211</u>	660.8 <u>661</u>	3.5
9	GB2	CAMELS-GB	283.60	0.42	1.25	0.00	1209.2 <u>1209</u>	664.9 <u>665</u>	8.0
10	LAM	LamaH-CE	102.29	0.58	0.58	0.32	1777.0 <u>1777</u>	1299.5 <u>1300</u>	0.9

The 10 selected catchments span a wide range of climatic and hydrological settings. They differ substantially in aridity and mean temperature, ranging from humid, cool basins such as LAM and HYS (~~aridity~~moisture index $\approx$ 0.3–0.6, mean temperatures near 0–4 °C) to warm, ~~dry~~ catchments like AUS1 and BR6 (negative ~~aridity indices~~, moisture indices indicating high aridity, mean temperature > 18 °C). Snow fraction varies ~~markedly~~substantially, from snow-free catchments in Australia and Brazil to snow-influenced basins such as C12 and LAM. Precipitation and runoff also contrast strongly, with wetter, high-runoff regions (e.g. LAM, GB2) versus more arid, low-runoff sites (AUS1, ~~AUS6~~BR6), illustrating the climatic and geographic diversity captured by the dataset ~~and clustering approach~~. To avoid known issues with the ERA5-based potential evapotranspiration (PET) estimation from CARAVAN (Clerc-Schwarzenbach et al., 2024), we are using the provided alternative based on Penman-Monteith. We have compared the CARAVAN data for precipitation, potential evapotranspiration and temperature against the native forcing data provided by each catchment’s source dataset (e.g. CAMELS, CAMELS-AUS, CAMELS-BR, CAMELS-GB, HYSETS, or LamaH-CE) and found no systematic bias, confirming the CARAVAN forcing was of sufficient quality for consistent use across all catchments (details in Figures S3–S12 in the Supplementary Material).

2.2 Model Calibration Procedure

All models were run on daily time step and calibrated with the Covariance Matrix Adaptation Evolution Strategy (CMA-ES) algorithm by Hansen et al. (2003) as implemented in the MARRMoT Toolbox. We used a 10-year time period from 2004 to 2014 for calibration and a 10-year time period from 1993 to 2003 for evaluation with a one-year warm-up period each. The evaluation period of the catchment AUS6 had to be shortened to eight years, ~~however, because of large~~ because of data gaps. We conducted ~~3760 (47x8x10)~~18,800 (5x47x8x10) calibration runs to calibrate each of the 47 models on all eight metrics in each of the ten catchments. ~~For all conducted analyses, we decided to analyze the signature representation during the~~ To

consider parameter uncertainty we calibrated each combination on five different seeds and retained the best 500 parameter sets across all seeds.

For all analysis, we focus on the calibration period only, because ~~the impact of the objective function is more direct. Because the signature representation is generally consistent between calibration and validation periods (see Figure S3 in Supplementary~~
220 ~~Material) we believe this to be a feasible approach~~ it provides the closest connection between objective function choice and resulting model behaviour. If objective functions cannot produce accurate signatures during calibration, it is unlikely that they will systematically do so during evaluation. We provide additional plots showing model and signature performance during the evaluation period in the Supporting Information (Figures S13/14 in the Supporting Materials).

2.3 Model Benchmarking

225 ~~As we do not want to contaminate the analysis through poorly performing models, we~~ We establish a baseline performance benchmark to remove poorly performing models from further analysis. Based on earlier arguments and examples (~~Schaeffli and Gupta, 2007;~~ (Seibert, 2001; Schaeffli and Gupta, 2007; Knoben et al., 2020; Knoben, 2024), we use ~~the~~ an ensemble of benchmark models that can be used to establish the minimum performance a model should exceed. The HydroBM package (Knoben, 2024), provides 22 benchmark models that can easily be used. They cover three broad categories: benchmarks derived from streamflow
230 data, such as the daily mean flow ~~as a benchmark, because it timeseries, which~~ accounts for seasonality ~~and provides a more informative baseline than the average flow;~~ benchmarks derived from rainfall-runoff ratios; and very simple empirical models that approximate aggregated catchment behaviour. For each catchment, we compute ~~this benchmark from the performance of the timeseries these benchmark models create by comparing it to the~~ observed discharge data in the calibration period ~~and evaluate it with~~. By calculating all eight performance metrics ~~Models~~ we can identify which benchmark model yields the
235 highest performance for each metric and catchment. This performance will be used as the benchmark to beat. Model runs that fail to outperform ~~the~~ this highest benchmark are excluded from further analysis, ensuring that only sufficiently reliable models are carried forward to the analysis of signature performance. ~~For catchments with a snow fraction larger than 30% we additionally remove all models without a snow module from the analysis should they manage to beat the benchmark.~~

2.4 Signature Selection

240 We selected 15 streamflow signatures to investigate the extent to which objective function choice influences a model's ability to replicate streamflow signatures. Our goal with this selection is to cover different aspects of the flow regime, as well as a range of hydrological processes (McMillan, 2020). The selected signatures are shown in Table 5 and were all calculated using the TOSSH toolbox (Gnann et al., 2021), ~~a more detailed description of each signature can be found in the Supplementary Material S2.5.~~ Most of the considered signatures describe streamflow properties such as the slope of the flow duration curve (FDC Slope, ~~33rd-66th exceedance percentiles~~), the 5th and 95th streamflow ~~percentile percentiles~~ (Q5 and Q95), the high and low flow duration (~~HF Dur & LF Dur~~ HFD & LFD) and frequency (~~HF Freq & LF Freq~~ HFF & LFF) as well as the mean half flow date (MHFD). The slope of the FDC describes the variability of the streamflow regime, indicating how quickly a river responds to rainfall events and how sustained the flows are during dry periods. Representing it correctly therefore indicates

a good representation of in-catchment storage and flashiness behaviour. The 5th and 95th streamflow percentile describe the magnitude of low and high flows. Representing these signatures well indicates a good representation of extreme flow behaviour, while representing the high and low flow duration and frequency well ensures that also ~~the~~ these important characteristics of extreme flow can be met. The MHFD helps to determine if the general timing and seasonality of streamflow is represented well. Additionally, we evaluate the impact of the objective function on the total runoff ratio (Total RR), event runoff ratio (Event RR), baseflow index (BFI), baseflow recession coefficient (BFRC), flashiness index (FI), variability index (VI) and rising limb density (RLD). These are intended to deliver information on the different catchment processes such as the general water balance (Total RR), individual storm responses (Event RR), the baseflow processes (BFI, BFRC) or general water storage behaviour and runoff dynamics (flashiness index, variability index, rising limb density).

2.5 Analysis of Objective Function Influence

2.5.1 Signature Error Metric

As most signature values have individual ranges and can ~~significantly differ between individual~~ differ substantially between catchments, we analyze normalized signature errors for ~~a~~ a better comparison. ~~The metric is calculated as follows~~ We first define the error at the level of individual retained runs, then aggregate it (additional details are provided in the Supplementary Material S2.4)-(6).

For a single retained run (one catchment, one model, one objective function, one parameter set), the signature error is the deviation of the simulated signature from the observation,

$$\underline{EM_{jk}e_{ijk r}} = \frac{\underline{med_i(S_{ijk} - y_j)}}{|\underline{max(med_i(S_{ijk} - y_j))}|} \underline{S_{ijk r} - y_j}, \quad (1)$$

~~where S_{ijk} where $S_{ijk r}$ is the simulated signature value depending on for model i , catchment j , and objective function k , and retained run r , and y_j refers to is the observed signature value depending on the location, and $med_i(-)$ denotes the median taken over the ensemble index i for catchment j .~~

To place signatures of different magnitudes on a comparable axis, errors are normalized on a per-catchment basis. The normalization denominator D_j is the largest absolute typical error across objective functions for that catchment,

$$\underline{D_j} = \max_k(|\underline{\bar{e}_{jk}}|), \quad (2)$$

where $\bar{e}_{jk}$ is the typical (median) error of the retained ensemble, obtained by first taking the median over the retained runs of each model and then the median across models,

$$\underline{\bar{e}_{jk}} = \underline{\text{med}_i \left(\text{med}_r(e_{ijk r}) \right)}. \quad (3)$$

The sequential median ensures that each model contributes equally to the typical error, independent of how many of its retained runs pass the benchmark.

The range for this metric is limited to $[-1, 1]$, with positive values indicating run-level normalized error is then

$$em_{ijk_r} = \frac{e_{ijk_r}}{D_j} = \frac{S_{ijk_r} - y_j}{\max_k (|\text{med}_i(\text{med}_r(S_{ijk_r} - y_j))|)}. \quad (4)$$

280 Positive values indicate overestimation and negative values indicating underestimation of the observed signature value. These values are calculated on a catchment-by-catchment basis so that if one OF is worst. The normalization is performed catchment by catchment, so that an objective function that is the most biased in all catchments, all error metric values would indicate ± 1 . A value of 0 indicates would reach ± 1 at the median level. Because D_j is set by the typical (median) error while individual runs may deviate further, run-level values can exceed $[-1, 1]$.

285 Aggregating over the retained ensemble gives a catchment- and objective-specific error metric. Since D_j is independent of model and run, the median passes through the normalization,

$$EM_{jk} = \text{med}_i \left(\text{med}_r(em_{ijk_r}) \right) = \frac{\bar{e}_{jk}}{D_j}, \quad (5)$$

which by construction lies in $[-1, 1]$, with 0 indicating that the observed value was is matched exactly.

This initial error metric can be aggregated further by using the median value over all catchments. This gives an error metric specific to a signature. Finally, aggregating across catchments yields a single value per signature and objective function, which can be compared to find is used to identify the objective function that best represents each signature.

$$EM_{k,absk,abs} = \text{abs med}_j \left(\text{med}_j |EM_{jk}| \right) = \text{abs med}_j \frac{\text{med}_i(S_{ijk} - y_j)}{\max(\text{med}_i(S_{ijk} - y_j))} \quad (6)$$

2.5.2 Statistical Testing of Objective Function Influence

295 To better understand which signatures are most affected by the objective function choice, we conducted a paired sample t-test: a pairwise comparison for two objective functions at a time by testing their median value of signature values across catchments. The null hypothesis for the t-test used here is H_0 : The mean values of the distributions are equal to each other. With 8 objective functions this leads to $\sum_{k=1}^n k = \frac{n(n-1)}{2} = \frac{8*7}{2} = 28$ data points $\binom{8}{2} = \frac{8*7}{2} = 28$ pairwise comparisons for each signature. The signature value entering each comparison is the per catchment median error EM_{jk} (Eq. 5), in which model and parameter-set variability have already been collapsed through the sequential median (Sect. 2.5.1); the test thus pairs the two objective functions catchment by catchment across the ten catchments, and the retained top-500 runs are not treated as independent samples.

This analysis highlights which signatures are most strongly influenced by the choice of objective function. We report the resulting p -values for each pairwise comparison and consider differences statistically significant at $\alpha = 0.05$. Not all distributions are expected to differ significantly, as some objective functions are likely to produce similar results. Based on 305 these outcomes, we restrict subsequent analyses to signatures that are commonly significantly impacted by the choice of objective function, which we define as 3040% of the p -values being lower than 0.05.

2.5.3 Random Forest for Attribute Importance

For signatures that show significant sensitivity to the choice of objective function (based on the previous subsection), we apply a We apply random forest (RF) regression to assess the relative importance of catchment, objective function (OF), and model choice. The RF was trained with the vectorized signature error values. We consider two complementary configurations. In the first, the response is the median simulated signature value across the retained ensemble per (model, catchment, OF) combination; in the second, the response is the corresponding median signature error against the observation (taken from 2.5.1) as the response, and the corresponding section 2.5.1). In both cases the predictors are the categorical factor indices for OF (k), model (i), OF (k), and catchment (j) as predictors.

We used the TreeBagger implementation in MATLAB with 300-500 trees and regression mode, extracting out-of-bag permuted predictor importance scores. Negative importance values, which indicate no predictive contribution, were set to zero. For each signature and configuration, importance values were normalized to sum to one, providing a direct measure of the relative influence of the three factors on the signature values. The resulting matrix reports the mean relative importance of catchment, OF, and model choice for the representation of the analyzed signatures.

3 Results

The following sections will cover the outcomes of the benchmarking procedure (Section 3.1). Based on this, the general distribution of signature representation (Section 3.2) for (i) the detailed example of Total Runoff Ratio and (ii) all signatures (aggregated over the models) are shown. Afterwards, we will use the error metric (Equation 5 from Section 2.5.1) to highlight the skills and shortcomings of each objective function in Section 3.3. Further analysis is then used to test the statistical significance of objective function influence (Section 3.4) and assess the relative predictive importance of objective function choice (Section 3.5). Lastly, we use the aggregated error metric (Equation 6 from Section 2.5.1) to find the best performing objective functions regarding signature representation.

3.1 Benchmarking the Models

Figure 3 shows the performance of all 47 calibrated models with their 500 best parameter sets (violins) compared to the calculated benchmarks highest calculated benchmark (red dash). The benchmark scores are calculated as any of the model scores, by comparing a timeseries of benchmark "simulations" against the observations in each basin. The chosen benchmark is the seasonal cycle, given by a repeating sequence of daily mean flows (see Section 2.3). The goal of the benchmarking procedure is to retain only plausible models for the remainder of the analysis, as described in section 2.3.

The benchmark scores vary considerably between catchments and also between objective functions, as do the number of models that are able to beat the benchmark (shown as the number above varies considerably across catchments and objective functions (number shown below each violin). In some basins, the interannual mean (i.e. the benchmark) is a good predictor of the daily flow, and in these basins the benchmark scores tend to be high. Basins three catchments (C12 and LAM are catchments

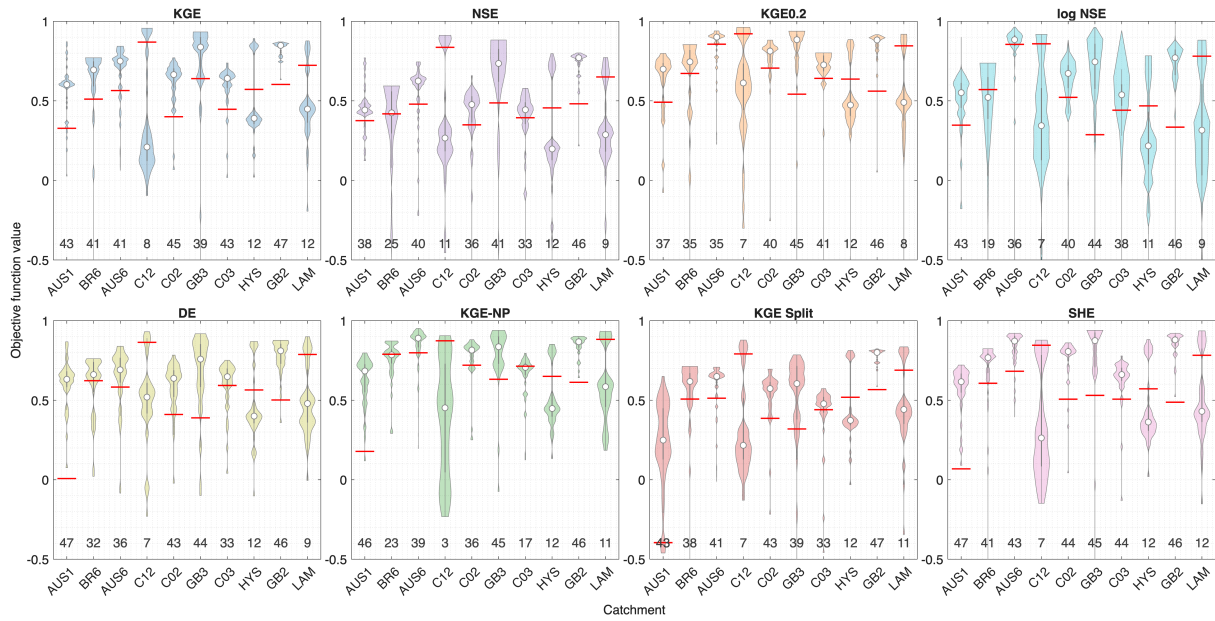


Figure 3. Benchmark Performance (driest catchment on left to wettest catchment on right). The dots represent distribution shows the objective function values across the top 500 parameter sets (across all calibrated seeds) for all 47 individual-models while and the red line show shows the highest benchmark score of the objective function. The numbers above/below each violin indicate the number of models that outperform the benchmark with at least one of their parameter sets. The white circles indicate the median of the data in the violin.

where snow plays a large role, and in these catchments this leads to a strongly seasonal flow regime that is relatively stable between years, and thus well captured by the benchmark. In these basins few models beat the benchmark, though this is not entirely unexpected because only eight of the MARRMoT models have (HYS, LAM) the number of models is comparatively low. The explanation is a strong snow influence, which produces seasonality that only the 12 of the 47 MARRMoT models that include a snow module capable of representing these local processes. Therefore there is a large spread in the number of models that can outperform the benchmark. The C12 and LAM basins, as well as HYS in most cases, contrast with the remaining catchments where typically a larger number of models outperform the benchmark can accurately reproduce. In other catchments, most models beat the benchmark, and the models that fail may be missing important process controls (e.g., groundwater representation) though this is more difficult to diagnose than the snow cases. In a handful of cases all models can beat the benchmark (AUS1 and e.g. GB2 for the NSE-KGE objective function), and for catchment C12 (high snow tendency, high seasonality), there are two objective functions (log NSE, SHE) for which no model was able to outperform the interannual mean benchmark. Models that fail to beat the benchmark might be missing key controls such as snow accumulation and melt, glacier or deep groundwater storage, so calibration cannot recover the correct timing or magnitude of seasonal peaks, allowing the simple daily-mean predictor to do better. For all following analysis, only those models that every combination of catchment and OF, there are at least 3 models exceeding the benchmark ensemble. As may be expected, more complex models (more

parameters and stores), tend to outperform the benchmark are-used more often and were retained more frequently (result not shown).

3.2 How accurate are the signature representations when different OF-objective functions are used?

To begin the analysis of how objective functions affect hydrological adequacy, we compare the simulated signature values with the observed signature values for all eight objective functions. Figure 4 shows the distribution of absolute errors (y-axis) over all catchments (x-axis) for each metric (subplots) for one of the 15 signatures, the Total Runoff Ratio. The plots for all other
360 signatures can be found in the Supplementary Material S2.7 (Figure S5 through S18.9 (Figure S16 through S29)).

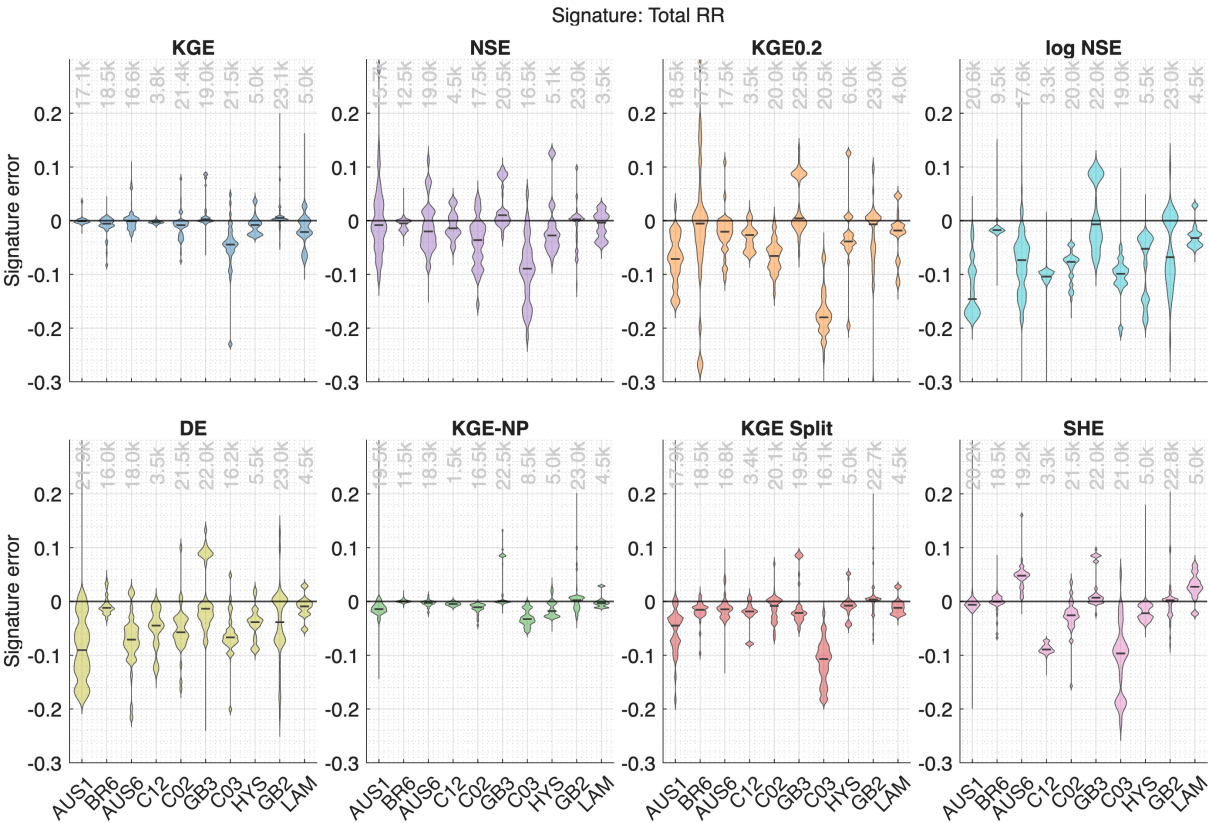


Figure 4. Absolute signature error (simulated - observed) on Total RR values for all objective functions, calculated as the difference between simulated. The shown distributions include all models and observed signature value their 500 best parameter sets that passed the benchmark. For The numbers above the combination violins indicate the number of catchment C12 and objective functions log-NSE and SHE, no runs (model outperformed the benchmark x parameter set) shown within each violin. A number of 23,5k would indicate all models and parameter sets were retained.

In Figure 4, there are three things of note. First, the location of the violin plots shows whether the signature values are underestimated (negative values), overestimated (positive values) or relatively unbiased (centered around zero). Second, the

size of the violins indicates whether the signature predictions are well-constrained (~~small-short~~ violin) or not (~~large-long~~ violin). Third, the catchments are ordered from driest (left) to wettest (right) indicating if signature error correlates with climate.

365 ~~Interestingly, there-There~~ is no clear pattern in Runoff Ratio errors related to the catchments' aridity values. Instead, there appear to be two categories of results: objective functions that return mostly unbiased models in all catchments (KGE and KGE-NP), and objective functions where the signature representation of the resulting models varies strongly per catchment (e.g. NSE: relatively unbiased in some basins, large variability in signature errors in others). The KGE-NP ~~contains-constrains~~ the variability within the models much better and is therefore assessed as the best available objective function for the general

370 water balance as represented through the Total RR signature.

~~Figure 4 shows that the-For the remaining six objective functions, the~~ Runoff Ratio tends to be underestimated ~~rather than overestimated by the given selection of objective functions and positive errors are relatively rare~~. This means that generally the models tend to put too much water into storage and/or evaporation or other sink terms, which may become relevant when modelling climate change impacts.

375 To show the results for more signatures in ~~a comprehensive way, it is necessary to aggregate the results in some form. This is why we are using the median-an aggregated way,~~ Figure 5 shows the median simulated signature value over all models ~~(exceeding the benchmark in performance)in-Figure 5above the benchmark (underlying data can be seen in Fig S16-S29).~~ The median was chosen to achieve the most representable ~~representation-summary~~ of the model ~~spreadensemble~~, which is not susceptible to outliers. This gives an overall impression on the capability of the different models to represent the

380 hydrological signatures of interest. It also allows a first impression of the impact different objective functions have on signature representation. In Figure 5, the coloured dots indicate the median ~~modeled-modelled~~ signature value for each objective function, while the black line shows the signature values calculated from observations. The grey violins indicate the variability in the modelled signature value if all individual model results with their 500 best parameter sets are considered (every model that beats the benchmark for any objective function). The plot therefore allows insights on how well a signature value can generally

385 be represented through the different tested models and gives a first impression on the influence the objective function has on these representations. The x-axis denotes the 10 selected catchments and the y-axis gives the signature value in its native unit (see Table 5).

Initially, we can use this plot to compare the violin ranges to the observed values of signatures. For the vast majority of cases, the observed value is always within the boundaries of the model ~~spreadensemble~~, which indicates that the tested conceptual

390 models are able to capture the hydrologic variability across diverse catchments. The spread of the colored dots around the observed values indicates different degrees of accuracy in the signature representation depending on the calibration metric used. For some signatures, the range of signature values is notably smaller (e.g. Total RR) than for others (e.g. Rising Limb Density). There are signatures (e.g. Q95, BFI) where the objective functions differ strongly from each other and signatures where different objective functions lead to similar modelled values (e.g. MHFD, ~~HF-Dur~~HFD). To move beyond these very

395 broad conclusions, the next section will use the previously introduced error metric (Equation 5 from Section 2.5.1) to assess the relative skill of each objective function in more detail.

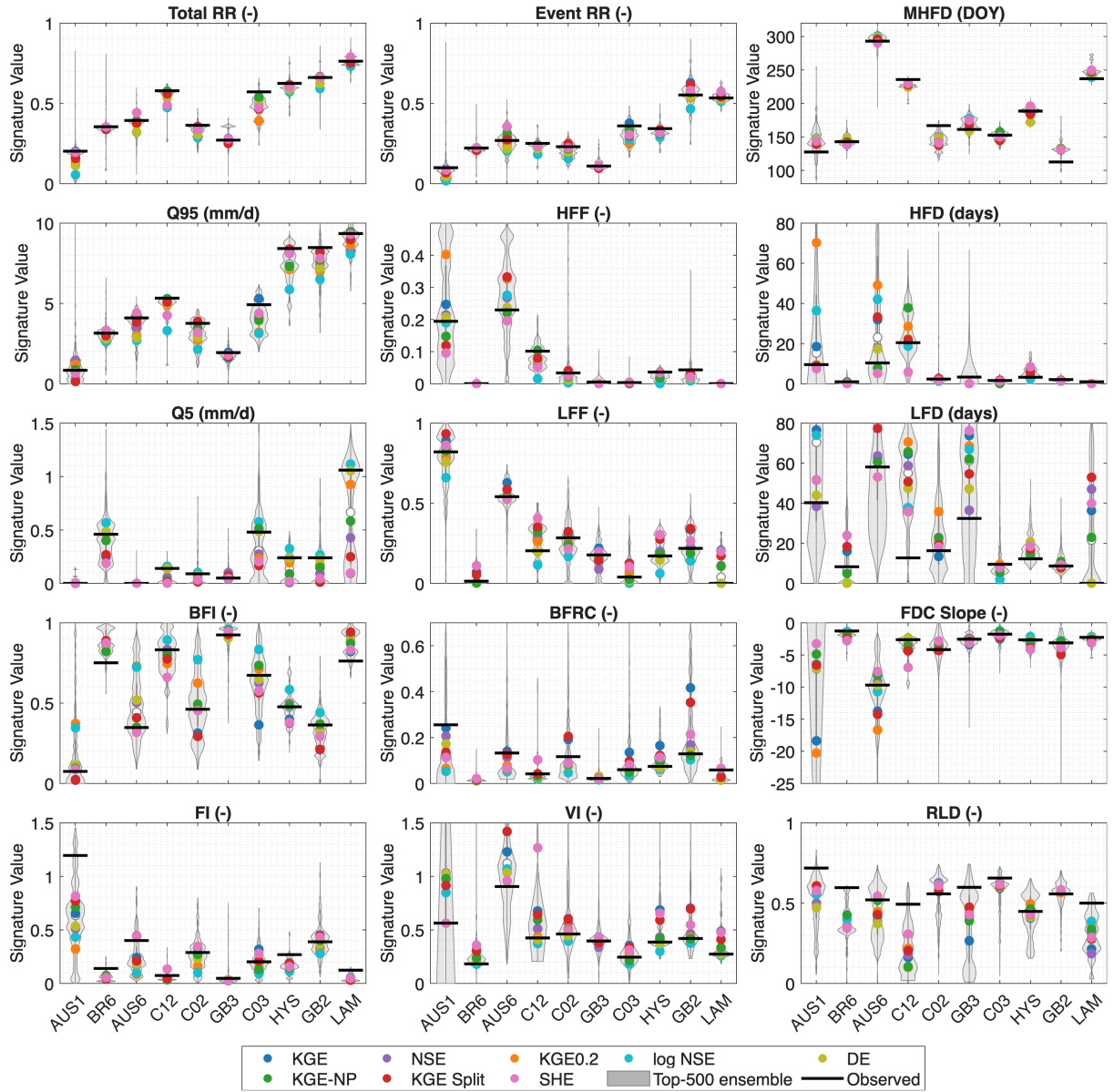


Figure 5. Comparison of signature representation capabilities across the different objective functions. The coloured points represent the median signature (across all [retained](#) models calibrated with a specific OF). The black dashes indicate the observed signature values. The grey violins represent the variability of the modelled signature values [of across the individual retained models with their top 500 parameter sets](#). In catchment AUS1, the observed value for the FDC slope is NaN, because the observed flow is 0 mm/d for more than 66% of the time steps in calibration period.

3.3 Which **OFs** objective functions show strengths or weaknesses for a specific hydrological signature?

In this section we use Equation 5 to normalize the signature errors using the observed signature value.

Figure 6 shows the normalized error in signature representation for all objective functions. Generally, we can see that the performance varies largely depending both on the signature (each violin in a plot) and the catchment (each dot in a violin/catchments/models (within the violins)).

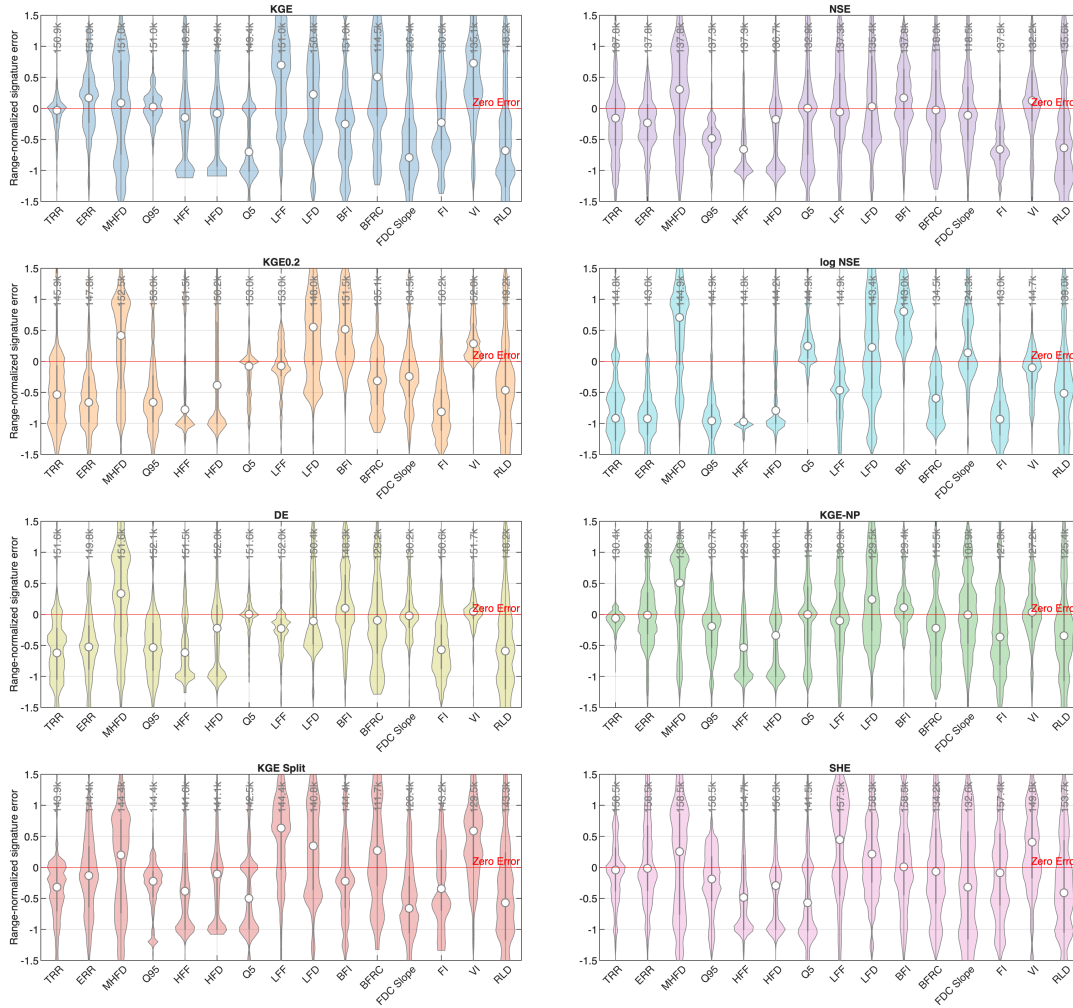


Figure 6. Error Metric as given in Equation 5 for all Objective Functions. The error metric is calculated for each objective function (subplot) per signature (x-axis) and catchment (dots data in violin). A value of 0 indicates no error in signature representation and a value of ± 1 indicates the largest under- or overestimation error. A value of ± 1 would indicate a performance equal the most extreme over/underestimation for the median across models and runs. As we are showing the top 500 runs for all benchmark-exceeding models explicitly, the values can exceed these boundaries. A number of 235.0k would indicate all models and parameter sets were retained.

The KGE in Figure 6 has good performance for the Total Runoff Ratio and the high flow percentile Q95 (violin ~~is clustered closely~~has a smaller spread around the zero error line). This does not only apply to the median performance (white dot in the violin) but also has a high consistency among catchments (~~spread of violin and parameter sets~~(violin spread)). However, particularly for the low flow percentile Q5, LFF, FDC slope, ~~BFI~~, and Variability Index the representation obtained from calibrating to KGE relative to the other objective functions is often worse.

The NSE performs mediocre on almost all of the evaluated signatures: the range within catchments (~~dots in and runs~~(spread of each violin)) is ~~large and there is no clear pattern in signature estimation~~often large. This indicates that the suitability of the NSE metric can be very ~~catchment-dependent~~High Flow Frequency, Flashiness Index ~~catchment-~~and ~~model-dependent~~HFF, FI, RLD and Q95 are usually underestimated by models calibrated on the NSE. Compared to the KGE, the performance on Q5, ~~FDC slope~~, and ~~BFI~~ is ~~better in specific cases, but subject to large variability~~and FDC slope is typically better.

For ~~log KGE~~KGE0.2 and log NSE, we see similar patterns in signature error. Both (~~log KGE~~KGE0.2 and log NSE) underestimate the runoff generation and high flows as well as the reactivity of the catchment (Flashiness Index). The log NSE is often the worst objective function among the eight investigated, indicated by ~~values close to ± 1~~ high or low values. Both metrics show improvements compared to ~~regular KGE and NSE~~the KGE for FDC Slope and Q5, but not necessarily for low flow duration (~~LF Dur~~LFD) and frequency (~~LF Freq~~LFF). It is important to notice that the FDC Slope can also be met well if the runoff is systematically biased, which is the case with underestimation here. In many of the evaluated 15 signatures, the ~~log KGE~~KGE0.2 shows a better performance than the log NSE (violin closer to 0), ~~but this cannot be generalized across locations based on this assessment~~.

The diagnostic efficiency generally performs similarly to the ~~log KGE~~It shows better performance for Q5 and KGE0.2. It mainly improves representation of the BFI but has comparable issues for Total RR and Q95. Compared to the KGE, the errors are constrained better across catchments (smaller violins). In conclusion, the diagnostic efficiency proves to be a ~~veritable~~viable alternative for low-flow calibration, while offering benefits on additional signatures such as BFI and VI).

The non-parametric version of the KGE performs ~~slightly~~ worse than the KGE on high flows, but improves other signatures strongly, such as the BFI ~~and FDC Slope~~FDC Slope, Q5 and LFF. The KGE-NP is among the objective functions that perform best overall(~~smaller violins~~). The expectations (general improvement due to more information being included) for KGE-NP are thus largely met, with the main benefit being better incorporation of flow variability for a larger fraction of flow percentiles.

~~For KGE Split, we see similar improvements as for KGE-NP but they are not as strong as for the non-parametric KGE as the violin spread remains large~~KGE Split performs similar to the regular KGE, though shows a tendency towards larger errors. Like the NSE, it has no signature that is represented both more accurately and consistently than alternative objective functions. Compared to the regular KGE, there are slight improvements on selected signatures (FDC Slope, BFI, Variability Index), ~~but none are clear enough to recommend this OF for a specific aspect~~.

Lastly, the SHE is one of the objective functions with the largest range in signature representation. The general patterns are similar to the KGE (as expected due to their similar formulation), except for selected signatures like FDC ~~slope~~Slope. This indicates that a change from value-based correlation to rank-based can be influential for signature representation.

The violin plots also show that some signatures are affected very little by the choice of objective function such as the Rising Limb Density and the Mean Half Flow Date and the Rising Limb Density, this is indicated both by a large spread of the violin plots and, which is indicated by a similar distributions among objective functions. This suggests that either the models or the catchment mainly influence the results. As the signature value cannot be simulated well from any of the models, this shows the limits of objective function influence.

~~This Section~~ The Supplementary Material (Figures S30 and S31) includes plots on the analysis with 25 and 100 parameter sets in comparison to the 500 used here. Since the results remain comparable, we assume that parameter uncertainty plays only a minor role here. Additionally, the SM also provides an example of the error metric collapsed to one value per catchment as it is being used in the assessment of overall skill (Figure S32), which emphasizes the differences more, but does not account for model- and run-level uncertainty.

3.4 Does the objective function choice significantly affect signature representation?

This previous section showed the skills and shortcomings of the individual objective functions regarding signature representation. However, it is not easy to assess whether the observed influences that were described can be considered significant for the process representation. To investigate this, ~~the next section uses statistical testing~~.

3.5 Does the OF choice significantly affect signature representation?

we apply statistical testing in this section. Figure 7 shows the distributions of p-values, when testing the significance of signature values calculated with different OFs against each other. ~~The red and black lines mark the commonly used significance levels of 0.05 and 0.1. Every signature that is below these values can be considered to have a statistically significant difference between the paired samples of signature value distributions (in other words, the two objective functions in the pair lead to statistically different values for the signature across all included models). For these signatures, the choice of objective function thus plays a large role.~~

All signatures show a large range in p-values, which implies that there are both objective functions which have similar and different representation for each signature. Some OFs typically behave similar, e.g. the distributions of log KGE KGE0.2 and log NSE are often almost identical comparable for low flow signatures. Figure 7 shows that the choice of objective function is ~~relevant only for the following tested~~ strongly affects the following signatures: the Total RR, Event RR, Q5, Q95, LF Freq LFF, BFI, BRFC BFRC, and the Flashiness Index. For the remaining 7 signatures (Low and High Flow Duration, HF Freq HFF, Mean Half Flow Date, FDC Slope, Variability Index, Rising Limb Density) a change of objective function does not lead to a clear shift in signature representation for most of the combinations of objective functions. Table S4 S3 in the Supplementary Material S2-8 11 shows the frequency of objective ~~functions~~ function pairs reaching p-values of below 0.10 or 0.05 respectively.

By comparing Figure 7 to Figure 5 we can conclude that there can be ~~a lot of~~ considerable spread in the representation of the signatures that do not show a significant difference. This is because the spread for signatures like LF Low Flow Duration, Variability Index, or Rising Limb Density is often not consistent, meaning that when comparing the distribution of signature values of two objective functions with each other, there is variation as to which one has higher or lower values. Therefore, the

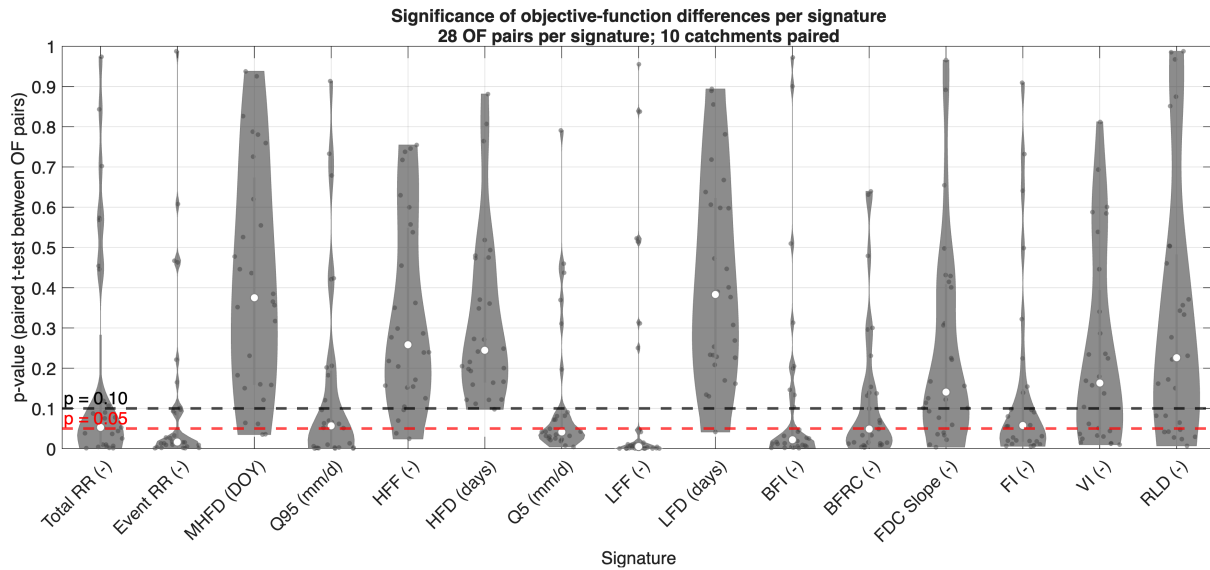


Figure 7. Significance Test between signature representations of the different objective functions. If the median value of a violin plot (white dot) lies below the red line ($p=0.05$) and black ($p=0.1$) lines indicate commonly used significance thresholds. Every signature that is below these values can be considered to have a statistically significant difference between the significance-levels paired samples of 0.05 or 0.1 signature value distributions (in other words, the choice of two objective function plays a significant role in the representation of this pair lead to statistically different values for the signature across all included models).

significance often shows no significant impact because the signal is not clear. The significance test indicates that the variation is not driven by the objective function and we speculate that in these cases other influences such as catchment characteristics or models drive the spread seen in the signature representation.

3.5 What controls simulated signature values?

By running a random forest analysis regarding the importance of the three varying components of this modelling experiment (Figure 8 shows the random forest feature importance for catchment, model, objective function), we can further evaluate the importance of the objective function choice by comparing it to the other two varying factors. Figure 8 shows the results of a random forest feature importance analysis, and OF (parameter-set variability is considered as a median value across the 500 parameter ensemble).

We can see that all signatures are dominated by the influence of Reading the value and error panels together clarifies the role of each factor. Catchment dominates the absolute signature values which is expected, since the basins were deliberately chosen to be hydroclimatically diverse. However, its importance for the signature error is markedly lower and roughly equal across the eight signatures. Objective function has comparatively limited control over the value, yet its importance rises for the catchment, implying that the strongest driver of signature representation is location change. This is reasonable as climatic

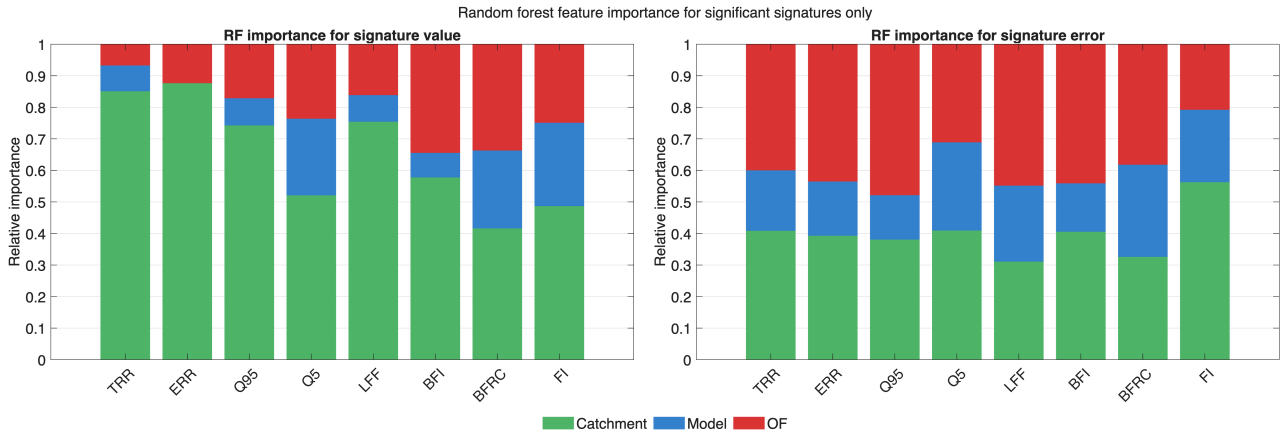


Figure 8. Random forest feature importance analysis of all varying components in this modelling experiment Forest Feature Importance based on predicting the simulated signature value (catchments, models, objective functions left) and signature error (right) using the median over the top 500 performing runs.

and landscape attributes are typically viewed as the most important drivers of hydrologic behaviour. The choice of objective function, however, has relevant predictive importance for a large number of the signatures. Variation in objective function is typically even more impactful than a change of the hydrological model. Particularly for the error, most strongly for Total RR, Event RR and Q95, and typically exceeds that of model choice. Model structure contributes throughout and gains relevance for low-flow and baseflow signatures (BFI and BFRC) and, BFRC, Q5 the importance of the objective function combined with the model choice even), where together with OF choice it challenges the catchment as the most important predictor. Model choice seems to play a more important role in representing Q5, BFRC and the Flashiness Index leading predictor. The pattern is consistent across signatures: the catchment largely controls the order of magnitude of the simulated signatures, but the choice of objective function and model structure play a large role in how closely those magnitudes match the observations.

3.6 Overall Skill of Signature Representation

Our results suggest that all objective functions can represent some signatures well while having shortcomings in others. This implies that, for our selection of objective functions when selecting an objective function, there is not a single one-choice that will ensure realistic values across the spectrum of streamflow signatures we investigated. Yet, we are able to can identify objective functions that generally have good process representation on most signatures. We will therefore investigate across several signatures by investigating the accumulated errors of all significantly influenced signatures as identified from shown in Figure 7. For this, we use equation 6 to get one median error value per signature and objective function. Figure 9 shows this accumulated error for the eight tested objective functions.

This plot indicates a few interesting results. The modeled overall error in signature representation varies strongly depending on the objective function chosen. Considering the entire set of signatures with significant metric influence, the KGE-NP,

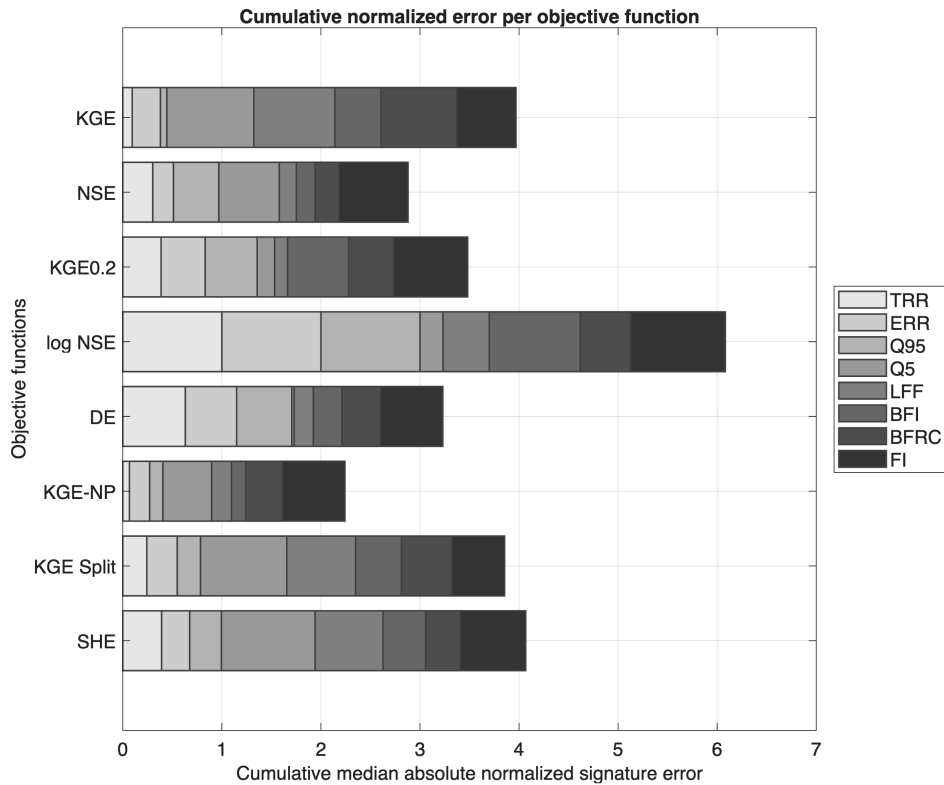


Figure 9. Cumulative error in significantly influenced signatures per tested objective function. [This was calculated using the median across the top 500 parameter sets per model.](#)

[SHENSE](#), and DE show the best overall performance. Conversely, this study identifies the log NSE as [the-worst-objective functions-overall-having the overall largest errors](#). This, however, does not immediately imply that certain OFs should not be used as each has its strengths and weaknesses. As shown before, the KGE is the best metric for high flow conditions (small errors for Q95) and works well for Total RR, while it has considerable weaknesses for low flow representation (e.g. BFI, BFRC and Q5). And while the log NSE performs [well-decently](#) on low flow conditions (Q5) it is considerably worse for [all-most](#) other significant signatures. Similarly, despite the KGE-NP showing the best overall performance, it still comes with weaknesses in Q5 and FI representation.

[As a robustness check, we repeated the central signature-error and objective-function ranking analysis for the evaluation period. The evaluation-period results broadly confirm the calibration-period patterns. As expected, signature errors generally increase during evaluation, but the relative strengths and weaknesses of the objective functions remain largely similar. In particular, the ranking of objective functions for the OF-sensitive signatures is comparable between calibration and evaluation, suggesting that the main conclusions are not an artefact of calibration-period performance alone. Detailed evaluation-period results are provided in the Supporting Information \(Fig. S33 in the Supporting Materials\).](#)

We structured our discussion into three parts. First, we synthesize the main takeaways for signature representation across models, catchments, and objective functions. We condense our findings into some practical recommendations. Second, we discuss the correlation between different signatures and how they relate to the identified impact of objective function choice. Third, we outline the limitations of our study and potential future work.

520 4.1 Main Takeaways and Practical Recommendations

Across 47 model structures and 10 hydro-climatically diverse basins, OF choice exerts a systematic but selective influence on simulated streamflow signatures. A paired significance test indicated that 8 of 15 signatures vary significantly with OF choice (Fig. 7), whereas timing descriptors such as MHFD and RLD are largely OF-insensitive. This may in part be due to the selection of OFs, as objective functions that specifically target time offset (e.g Mean Absolute Peak Time Error; MAPTE) are less frequently used (Ehret and Zehe, 2011) than the OFs we did select ~~and were thus not tested in our study~~. Generally, no single OF excels at all tested signatures. Considering the cumulative error across OF-sensitive signatures (Fig. 9), KGE-NP yields the lowest overall error, with ~~SHE, DE, and NSE~~ NSE and DE performing competitively in some basins but less consistently across signatures. ~~And while log-transformed OFs~~ The DE and the selected transformed OFs (log NSE, KGE0.2) reproduce low flow metrics better (Q5, ~~LF-Freq~~), ~~they~~ LFF, ~~but particularly the latter~~ come with considerable tradeoffs in the representation of other signatures such as Total RR, baseflow related signatures, or HF signatures.

Our ~~analyses~~ analysis hence contributes to the available information of specific strengths and weaknesses of individual OFs. KGE, for example, most reliably reproduces Q95 (Mizukami et al., 2019), and KGE-NP is one of the best OFs for total runoff ratio and ~~baseflow characteristics (Figure 5 and 6)~~ for the baseflow index (BFI), where it leads, while sharing the best representation of baseflow recession (BFRC) with DE. Log NSE ~~in comparison~~, in comparison, showed few strengths for representing hydrological signatures in our experiments, ~~but log KGE did well~~ performing well only on Q5 ~~and the variability index (VI)~~. However, ~~the DE shows those same strengths, while having~~ DE reproduces Q5 comparably while carrying a smaller overall error, which is why we recommend DE for that signature and DE or KGE0.2 for low-flow frequency rather than the transformed NSE. The KGE-Split, NSE, and SHE show a balanced error field with individual differences in cumulative errors (Fig. 9). ~~While information like this helps to guide the choice of OF for purpose-based model calibration, it remains important to consider weaknesses of the individual metrics that have not been part of our experiments. Log KGE, for example, has been shown to lead to numerical issues and should be avoided according to Santos et al. (2018a), and the KGE has been criticized for leading to counterbalancing errors as described in (Cinkus et al., 2023). Thus, it remains important to consider all strengths and weaknesses of a specific OF and how it may influence the acquired modelling results; within this group, NSE is the most reliable for Event RR, where it ranks alongside KGE-NP, and KGE-Split is otherwise unremarkable but reproduces the Flashiness Index best of all tested OFs.~~

We also compared the influence of OF choice on signature representation with other influencing factors such as catchment characteristics or model structure choice. Our random-forest analysis showed that catchment characteristics are the dominant

control of signature representation, with OF choice typically more influential than model choice ~~and therefore placing second~~ (Figure 8). Model choice does, however, gain relevance for signatures connected to low flow or baseflow representation, while
550 OF choice seemed especially relevant for baseflow related signatures. ~~Together,~~

When reducing the impact of the conducted catchment selection by training the random forest on signature errors instead of signature values, the importance of the objective functions and model selection increased considerably for most signatures. This indicates the importance of appropriate objective function and model choice across hydro-climatic differences, which may be an overlooked aspect in large sample studies, potentially indicating a path for more detailed evaluation using signatures and their deviations. The broader analysis also contextualizes for which characteristics objective function choice can be influential, while others are predefined by catchment (MHFD) or a combination of catchment and model (RLD), which is shown in figure S34 in the SM. In summary, our results support the argument for a purpose-based model calibration, that considers ~~multiple specific~~ aspects of the flow regime, ~~and multi-objective calibration setups,~~ rather than defaulting to a familiar single metric (Mai, 2023; Jackson et al., 2019).
555

For the signatures that were found to be significantly impacted by OF choice, we used our results to develop guidelines for which OF will perform well in reproducing a given signature as shown in Table 6. We can conclude that ~~every OF leads most OF lead~~ to the best representation for at least one signature, highlighting ~~that no OF is best across all signatures~~ the need for appropriate selection. However, if the goal is a balanced representation of the water-balance and variability across hydrologic regimes, KGE-NP is a strong option. If low flows are central (ecological flows, drought assessment), preference should be
560 given to DE or ~~log-KGE~~ KGE0.2. If peak flows are the priority (flood risk, infrastructure design), KGE generally performs best for Q95 and related high-flow signatures. Reflecting back on our review of metric choice justification (Fig. 1), we strongly recommend to document the trade-offs one is willing to accept explicitly (e.g., low-flow degradation under KGE) and, where possible, to verify unaffected signatures (e.g., MHFD, RLD) since OF changes have limited influence there.

~~When multiple aspects of the flow regime matter simultaneously, multi-objective calibration has been shown to provide considerable benefits (Hernandez-Suarez et al., 2018; Nemri and Kinnard, 2020). Multi-objective formulations can expose the trade-offs explicitly via a Pareto front and allow practitioners to choose solutions that best satisfy competing goals (Yapo et al., 1998; Efstratiadis et al., 2000). If a single composite metric must be used, a transparent weighting that reflects decision priorities can be helpful. Vis et al. (2015), for instance, argue for directly incorporating multiple ecological flow characteristics in the OF to improve the relevance of the model for management applications. Similarly, signature values have directly been used in calibration, either as additional objectives or as constraints, to target process realism (Gupta et al., 2008; Pool et al., 2017). This can potentially improve baseflow behaviour, low-flow frequency, or water-balance partitioning when those aspects are central. However, incorporating signatures does not guarantee a better hindcast/forecast skill in all settings, and can even degrade time-series fit if misapplied (Araya et al., 2023).~~
570

4.2 Impact of Signature Correlations

580 Signatures are not independent from each other. When investigating correlations among observed signatures (Fig. 10) we noticed three clusters with similar patterns. First, a *runoff/high-flow* cluster (Total RR, Event RR, and Q95) shows strong

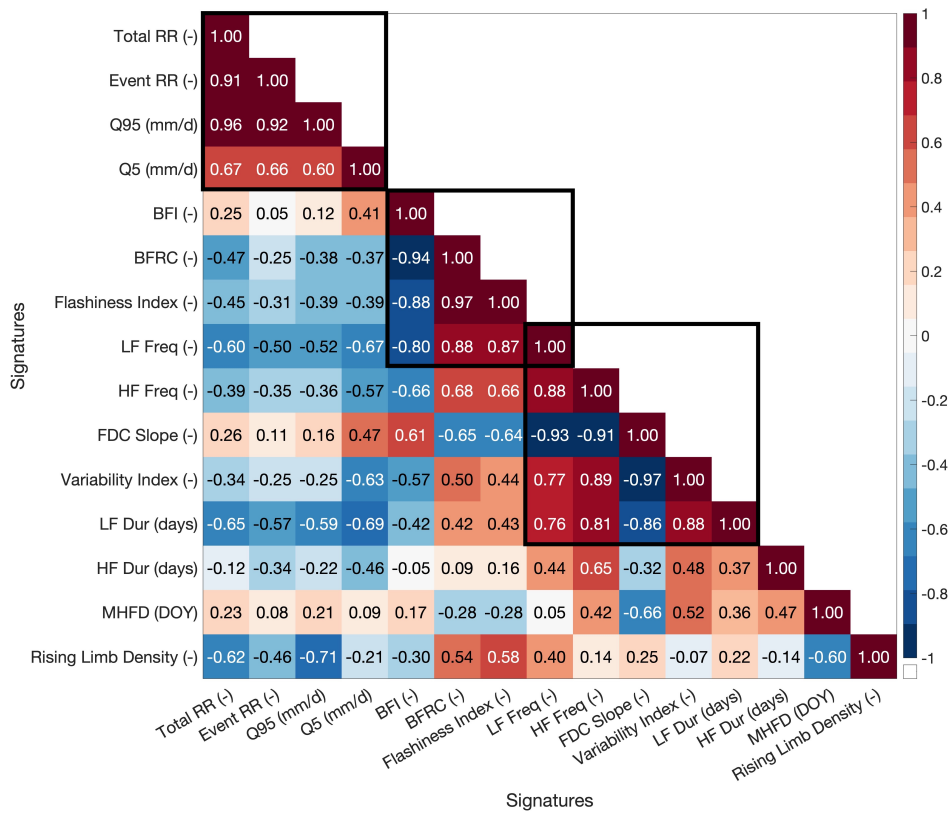


Figure 10. Correlation test between signatures based on Observations. Please note, that the signatures have been reorganized in this plot to highlight the emerging patterns. Supplement S2.9-14 Fig. S19-S35 provides a similar plot for the [comparing the observed correlations between-to](#) the modelled [signaturesones](#).

positive association, linking long-term water partitioning to the reaction to events. This aligns with aridity being a primary control on mean runoff (Berghuijs et al., 2017). Second, a *baseflow/storage* cluster, including BFI, BFRC and FI, exhibits tight coupling (higher BFI $\leftrightarrow$ lower BFRC and FI) and a strong relationship to low-flow frequency, reflecting how sustained baseflow suppresses low-flow events. Similar to the first cluster, this group was also sensitive to objective-function (OF) choice. Third, a *variability/threshold* cluster (FDC slope, **HF****HFF**/**LF**-frequency**LFF**, low-flow duration, VI) captures distributional shape and threshold exceedance. These signatures, derived from the hydrograph, partly encode duplicate information and were mostly OF-insensitive. Several other signatures showed weak or inconsistent ties to these clusters, such as High-Flow Duration, Mean Half-Flow Date, Rising Limb Density.

These clusters help explain why improving one signature typically moves its cluster-mates, yet they do not imply that different OFs are needed for each cluster. For example, KGE-NP performed best for many signatures in the first two clusters (e.g., Total/Event RR, BFI).

We identified some overlap between signatures that are strongly/weakly affected by OF choice and those that were identified as more/less predictable from hydrologic drivers by Addor et al. (2018): some signatures (e.g., FDC slope, durations) tend to be OF-insensitive with low spatial ~~predictability~~predictability, whereas water-balance and high-flow signatures were more OF-sensitive and spatially predictable. A notable exception is Mean Half-Flow Date, which was well predicted in Addor et al. (2018) but showed little OF sensitivity in our study suggesting a stronger climatic control.

If we consider the structural differences among OFs some of these patterns become clearer. For *bias/mean* terms, KGE's mean component directly constrains volumes, explaining its systematic advantage for Total RR and Q95. OFs that deviate from this approach (NSE, DE, log ~~KGE~~/NSE; Table 3) typically perform worse for these metrics (Gupta et al., 2009; Mizukami et al., 2019). Across objective functions, Total RR tends to be underestimated (Fig. 4), reflecting systematic water-balance biases related to excessive storage or evapotranspiration. This pattern underscores the importance of bias terms in objective functions, which directly constrain mean flow and promote a more realistic overall water balance (Mizukami et al., 2019). This underestimation is strongest for low-flow-focused objectives(~~log-metrics~~), as also noted by Vis et al. (2015). ~~The~~Our results do not support the claim that KGE resolves NSE's ~~low-flow~~'s low-flow insensitivity or performs well simultaneously for high and low flows (Althoff and Rodrigues, 2021); ~~but note~~; note, however, that Althoff and Rodrigues (2021) investigates more basins (albeit in a specific biome), while we test a larger number of models.

For *variability* terms, replacing variance ratios with FDC-based components (as done for the KGE-NP) or integrating relative errors along the FDC (as done for the DE) re-targets calibration toward distributional shape and mid/low flows, improving BFI, Q5, and Event RR in our results (Fig. 6). ~~While log-metrics and DE enhance low-flow behaviour, log transformations carry numerical pitfalls (Santos et al., 2018a); 1/Q variants can retain sensitivity with fewer issues (Pushpalatha et al., 2012).~~

For *correlation* terms, moving from Pearson to Spearman (KGE-NP, SHE, ~~DE~~) reduces sensitivity to extremes and stabilizes correlation. A direct comparison of KGE and SHE, where the dependence component is the only structural change, suggests ~~gains improved representation~~ for FDC slope, VI, and baseflow metrics, with similar or slightly worse performance for high-flow signatures (Kiraz et al., 2023). Annual ~~reweighting via Split-KGE~~ weighting using the KGE-Split yielded mixed, context-dependent benefits and no systematic improvement of low-flow characteristics over baseline KGE in our experiments (Fig. 9; Fowler et al., 2018).

Overall, each ~~OF~~objective function emphasizes the hydrograph component it encodes(~~bias, variability, dependence~~), ~~none dominates across all~~ and no single OF performs best across all signatures. Within this picture, the variability term emerges as a particularly influential lever: changes to it propagated to more signatures than changes to the correlation term, an effect most pronounced for low- and mid-flow signature, while bias terms remain essential for water-balance and event signatures. This ~~aligns component-specific behaviour is consistent~~ with evidence that ~~flood characteristics are more sensitive to metric choice than droughts (Melsen et al., 2019) and with comparative studies of KGE vs. NSE(Mizukami et al., 2019)~~signature performance depends strongly on which aspect of the hydrograph a metric targets (Mizukami et al., 2019; Melsen et al., 2019), and reinforces that OF selection should be guided by the signatures of interest rather than by a search for a universally optimal metric.

Streamflow transformations act as a further handle on signature representation, operating not on a single component but on the flow values entering all three. By compressing the dynamic range, the power and log transformations (KGE0.2, log NSE) shift calibration emphasis toward low and mid flows, which is reflected in their improved Q5 and LFF representation at the expense of high-flow and water-balance signatures (Fig. 6). The choice and strength of transformation is therefore itself a design decision shaping which signatures a calibration favours, complementing the component-level differences discussed above (Thirel et al., 2024).

4.3 Limitations and Further Research

Our study aimed to fill a specific research gap in the current understanding of the impact of objective function choice on signature representation; using a larger number of models for this type of analysis (see across diverse hydrological regimes (in comparison to Table 1). To enable a thorough investigation of the results we necessarily had to impose certain limits in other parts of the experimental design (i.e., having to “balance depth with breadth”, Gupta et al. (2014)).

First, compared to most studies in Table 1 we use a limited number of catchments. This keeps computational cost manageable and allows more detailed analysis into individual modelling results than would otherwise be possible. We selected these basins to ensure a hydro-climatic spread, as is common practice in these scenarios (see for example the 12 MOPEX basins that have been used for a large number of studies Duan et al. (e.g. 2006); van Werkhoven et al. (e.g. 2008); Carrillo et al. (e.g. 2011) (e.g. Duan et al., 2006; van Werkhoven et al., 2008; Carrillo et al., 2011). Compared to existing studies, our selected number of objective functions and signatures is fairly typical, while our number of models is clearly higher. The presented study found relevant differences in signature representation depending on the model structure, which might indicate a limitation in generalizing signature influence from single-model studies. Conversely, our broader model sample implies that conclusions drawn from the many-basin, single-model studies in Table 1 may themselves be difficult to transfer to other model structures, given that we find model choice to be a non-negligible, signature-dependent control.

Second, our experiments necessarily exclude parameter uncertainty: every combination of model, catchment and objective function is calibrated only once. To briefly investigate the impact of parameter uncertainty, we re-calibrated one model (GR4J) with multiple random seeds (Fig. S20 in the Supplementary Material S2.10, catchment C03). Optima were stable across seeds for KGE, NSE, and DE, whereas log-KGE, KGE-NP, and Split-KGE returned different parameter sets with similar OF scores. Importantly however, the best-performing parameter sets had very similar parameters and signature representation was nearly identical across these sets, showing that in this specific case parameter uncertainty is not a large concern. However, more work on this is needed. we focus on single-objective calibration because of its prevalence in the current literature. When multiple aspects of the flow regime matter simultaneously, multi-objective calibration has been shown to provide considerable benefits (Hernandez-Suarez et al., 2018; Nemri and Kinnard, 2020). Multi-objective formulations can expose the trade-offs explicitly via a Pareto front and allow practitioners to choose solutions that best satisfy competing goals (Yapo et al., 1998; Efstratiadis and Koutsoyiannis, 2003). If a single composite metric must be used, a transparent weighting that reflects decision priorities can be helpful. Vis et al. (2015), for instance, argue for directly incorporating multiple ecological-flow characteristics in the OF to improve the relevance of the model for management applications. Similarly, signature values have directly been used in calibration, either as additional

objectives or as constraints, to target process realism (Gupta et al., 2008; Yilmaz et al., 2008; Wagener and Montanari, 2011; Euser et al., 2011), and our component-level insights into individual metric strengths and weaknesses can help identify which signatures or complementary metrics to add when a particular aspect of the flow regime is the target. This can potentially improve baseflow behaviour, low-flow frequency, or water-balance partitioning when those aspects are central. However, incorporating signatures does not guarantee a better hindcast/forecast skill in all settings, and can even degrade time-series fit if misapplied (Araya et al., 2023)

Third, we note the following three methodological limitations in our study design: (1) The application of the median to calculate representative values within the error metric is an important choice, as this does not capture information on the spread of the signature representation. We accounted for that by visually inspecting the underlying violins, but did not explicitly quantify the spread within the applied metrics. (2) The random-forest importance scores should be read as a relative diagnostic within this specific design rather than a universal attribution of variance. The three predictors differ in cardinality (47 models, 10 catchments, 8 objective functions), which can affect permutation-based importance, and benchmark filtering removes poorly performing combinations beforehand. The scores therefore reflect experiment-specific influence rather than general rankings of hydrological controls. (3) Finally, we did not perform any sensitivity analysis before calibrating the models, and simply calibrated all model parameters. Across all calibration runs, approximately 10% of parameters reached one of their bounds. We found no indication that these boundary hits were systematically concentrated in a specific catchment, model, or objective function. Still, boundary hits can indicate parameter insensitivity, missing processes, data issues, or overly restrictive parameter ranges. They are therefore difficult to interpret in large multi-model calibration experiments and should be considered as an additional diagnostic of calibration reliability in future benchmarking studies.

Going forward, we see four priorities to build on the individual strengths of existing work (Table 1). First, broaden external validity with a larger, stratified basin sample (including ephemeral and groundwater-dominated systems) and potentially a wider model palette spanning physically-based and ~~ML models, using hierarchical/mixed-effects analyses to partition variance among catchment, OF choice, and model structure~~ machine-learning models. Second, test cluster-aware multi-objective calibration that deliberately pairs complementary OFs—e.g., a water-balance/high-flow metric (KGE or KGE-NP) with a low-flow/storage metric (DE or ~~log-KGE~~ $KGE_{0.2}$), optionally adding ~~an FDC-shape term—and evaluate gains via Pareto hypervolume and the change in worst-signature error within each cluster~~ signature-based components. Third, make component-level experiments less ambiguous by swapping KGE-type components one-by-one in a controlled ~~scaffold (mexwan: runoff-ratio vs. P -normalized mean; variability: stdev-ratio vs. FDC-shape; dependence: Pearson vs. Spearman vs. short-lag correlation)~~ setup, and test weighted KGE-type variants; ~~an accompanying analytical test case should trace how components propagate into variability metrics (e.g., VI).~~ Finally, propagate forcing/observation and signature-method uncertainty (e.g., baseflow separation) and explore ~~information-rich~~ information-rich objectives that go beyond the bias-variability-correlation trio, such as distributional divergences along the flow-duration curve, ~~mutual information-style dependence measures, or terms tied to process diagnostics (recession-slope distributions, intermittency, precipitation-discharge hysteresis), to~~ link calibration more directly to hydrologic processes (Kirchner, 2006). Designs like KGE-NP already go in this direction and offer a principled path to retain low-flow sensitivity without relying solely on log transforms.

5 Conclusions

We calibrated 47 conceptual model structures across 10 hydro-climatically diverse basins using 8 objective functions (OFs) and evaluated how well the top-500 parameter sets for each model simulate 15 streamflow signatures. Three results stand out. First, OF choice exerts a *selective* influence: 8 of 15 tested signatures are OF-sensitive, while others (e.g. mean half-flow date and rising-limb density, and often event durations) change only *little marginally* with OF choice. Second, no single OF performs best across all signatures; aggregated across OF-sensitive signatures, KGE-NP yields the smallest overall error across signatures. Third, *the direction of change is differences are* interpretable from OF design: metrics emphasizing mean volumes favour water balance and high flows; FDC-based and rank-based terms improve storage/low-flow behaviour.

These patterns translate into simple steps for guiding OF choice. For accurate water balance and high-flow magnitude, KGE (or KGE-NP if distributional shape matters) is effective. When low flows and storage-related signatures are central, DE or *log-KGE (noting log-transform caveats) can KGE0.2 should* be preferred. *When multiple signature clusters matter simultaneously, multi-objective calibration (e.g. pair KGE-NP with DE or a high-flow metric) can help to manage trade-offs explicitly.*

Catchment differences dominate overall variability, with OF choice typically While catchment differences dominate the signature values produced, it is the analysis of signature errors that matters for our purpose, since it is unsurprising that geography largely sets observed signature values. When attributing the error in reproducing each signature, the choice of objective function and model structure rise in importance, typically with OF more influential than model structure *in our setup*. Model structure nonetheless exerts a non-negligible, signature-dependent control on this error, particularly for low-flow and baseflow signatures, implying that findings from single-model studies cannot be assumed to generalise easily. Our findings are still bounded by the selected basins, conceptual structures, *signature methods, and deterministic calibration signatures, and calibration procedure but our use of 47 different models reduces that ambiguity*. Priorities for future work include *cluster-aware more systematic testing and/or* multi-objective formulations, *as well as* broader regime and *structure coverage (snow/glacier and groundwater-explicit models), and low-flow objective functions that retain sensitivity without log pitfalls.* model structure coverage.

In short, knowing the Multiple studies have indicated that OF choice is important for representing hydrologic signatures. Although process-oriented calibration may ultimately benefit most from multi-objective formulations, current practice still often relies on single-metric objective functions. For guidance on process-based modelling efforts, our study provides insights on the impact and limitations of objective function choice in a diverse selection of catchments, conceptual models and selected signatures. This is both useful as a validation approach for existing studies and guidance for further endeavours. Understanding the trade-offs in metric selection offers a practical path to models that reproduce the signatures that matter for the question at hand.

Code and data availability

. All code used in this analysis will be provided in a GitHub repository. The used data can be downloaded from the CARAVAN repository
730 on Zenodo.

Author contributions

. PW: methodology, data curation, formal analysis, investigation, visualization, writing - original draft, writing - review and editing; DS: supervision, conceptualization, methodology, investigation, writing - original draft, writing - review and editing; WJMK: supervision, investigation, writing - review and editing; NS: writing - review and editing, supervision

735 Competing interests

. No competing interests are being declared.

Acknowledgements

. The authors are grateful to the Center for Information Services and High Performance Computing [Zentrum für Informationsdienste und Hochleistungsrechnen (ZIH)] at TUD Dresden University of Technology, for providing its facilities for high throughput calculations.
740 PW and WK were supported by the Cooperative Institute for Research to Operations in Hydrology (CIROH) with funding under award NA22NWS4320003 from the NOAA Cooperative Institute Program. The statements, findings, conclusions, and recommendations are those of the author(s) and do not necessarily reflect the opinions of NOAA. During the preparation of this ~~work~~manuscript, some of the authors used ChatGPT ~~in order to improve language and readability~~of the original draft, with caution. After using these tools, the. The authors reviewed and edited ~~the content as needed~~all AI-assisted text and take full responsibility for the content~~of the publication~~.

- Addor, N., Newman, A. J., Mizukami, N., and Clark, M. P.: The CAMELS data set: catchment attributes and meteorology for large-sample studies, *Hydrology and Earth System Sciences*, 21, 5293–5313, <https://doi.org/10.5194/hess-21-5293-2017>, 2017.
- Addor, N., Nearing, G., Prieto, C., Newman, A. J., Vine, N. L., and Clark, M. P.: A Ranking of Hydrological Signatures Based on Their Predictability in Space, *Water Resources Research*, 54, 8792–8812, <https://doi.org/10.1029/2018WR022606>, 2018.
- 750 Althoff, D. and Rodrigues, L. N.: Goodness-of-fit criteria for hydrological models: Model calibration and performance assessment, *Journal of Hydrology*, 600, 126 674, <https://doi.org/10.1016/j.jhydrol.2021.126674>, 2021.
- Alvarez-Garreton, C., Mendoza, P. A., Boisier, J. P., Addor, N., Galleguillos, M., Zambrano-Bigiarini, M., Lara, A., Puelma, C., Cortes, G., Garreaud, R., McPhee, J., and Ayala, A.: The CAMELS-CL dataset: catchment attributes and meteorology for large sample studies – Chile dataset, *Hydrology and Earth System Sciences*, 22, 5817–5846, <https://doi.org/10.5194/hess-22-5817-2018>, 2018.
- 755 Araya, D., Mendoza, P. A., Muñoz-Castro, E., and McPhee, J.: Towards robust seasonal streamflow forecasts in mountainous catchments: impact of calibration metric selection in hydrological modeling, *Hydrology and Earth System Sciences*, 27, 4385–4408, <https://doi.org/10.5194/hess-27-4385-2023>, 2023.
- Arsenault, R., Brissette, F., Martel, J.-L., Troin, M., Lévesque, G., Davidson-Chaput, J., Gonzalez, M. C., Ameli, A., and Poulin, A.: A comprehensive, multisource database for hydrometeorological modeling of 14,425 North American watersheds, *Scientific Data*, 7, 243, <https://doi.org/10.1038/s41597-020-00583-2>, 2020.
- 760 Bennett, N. D., Croke, B. F., Guariso, G., Guillaume, J. H., Hamilton, S. H., Jakeman, A. J., Marsili-Libelli, S., Newham, L. T., Norton, J. P., Perrin, C., Pierce, S. A., Robson, B., Seppelt, R., Voinov, A. A., Fath, B. D., and Andreassian, V.: Characterising performance of environmental models, *Environmental Modelling & Software*, 40, 1–20, <https://doi.org/10.1016/j.envsoft.2012.09.011>, 2013.
- Berghuijs, W. R., Larsen, J. R., Van Emmerik, T. H. M., and Woods, R. A.: A Global Assessment of Runoff Sensitivity to Changes in Precipitation, Potential Evaporation, and Other Factors, *Water Resources Research*, 53, 8475–8486, <https://doi.org/10.1002/2017WR021593>, 2017.
- 765 Carrillo, G., Troch, P. A., Sivapalan, M., Wagener, T., Harman, C., and Sawicz, K.: Catchment classification: hydrological analysis of catchment behavior through process-based modeling along a climate gradient, *Hydrology and Earth System Sciences*, 15, 3411–3430, <https://doi.org/10.5194/hess-15-3411-2011>, 2011.
- 770 Chagas, V. B. P., Chaffe, P. L. B., Addor, N., Fan, F. M., Fleischmann, A. S., Paiva, R. C. D., and Siqueira, V. A.: CAMELS-BR: hydrometeorological time series and landscape attributes for 897 catchments in Brazil, *Earth System Science Data*, 12, 2075–2096, <https://doi.org/10.5194/essd-12-2075-2020>, 2020.
- Cinkus, G., Mazzilli, N., Jourde, H., Wunsch, A., Liesch, T., Ravbar, N., Chen, Z., and Goldscheider, N.: When best is the enemy of good – critical evaluation of performance criteria in hydrological models, *Hydrology and Earth System Sciences*, 27, 2397–2411, <https://doi.org/10.5194/hess-27-2397-2023>, 2023.
- 775 Clerc-Schwarzenbach, F., Selleri, G., Neri, M., Toth, E., Van Meerveld, I., and Seibert, J.: Large-sample hydrology – a few camels or a whole caravan?, *Hydrology and Earth System Sciences*, 28, 4219–4237, <https://doi.org/10.5194/hess-28-4219-2024>, 2024.
- Coxon, G., Addor, N., Bloomfield, J. P., Freer, J., Fry, M., Hannaford, J., Howden, N. J. K., Lane, R., Lewis, M., Robinson, E. L., Wagener, T., and Woods, R.: CAMELS-GB: hydrometeorological time series and landscape attributes for 671 catchments in Great Britain, *Earth System Science Data*, 12, 2459–2483, <https://doi.org/10.5194/essd-12-2459-2020>, 2020.
- 780

- Duan, Q., Schaake, J., Andréassian, V., Franks, S., Goteti, G., Gupta, H., Gusev, Y., Habets, F., Hall, A., Hay, L., Hogue, T., Huang, M., Leavesley, G., Liang, X., Nasonova, O., Noilhan, J., Oudin, L., Sorooshian, S., Wagener, T., and Wood, E.: Model Parameter Estimation Experiment (MOPEX): An overview of science strategy and major results from the second and third workshops, *Journal of Hydrology*, 320, 3–17, <https://doi.org/10.1016/j.jhydrol.2005.07.031>, 2006.
- 785 Efstratiadis, A. and Koutsoyiannis, D.: One decade of multi-objective calibration approaches in hydrological modelling: a review, *Hydrological Sciences Journal*, 55, 58–78, <https://doi.org/10.1080/02626660903526292>, 2010.
- Ehret, U. and Zehe, E.: Series distance – an intuitive metric to quantify hydrograph similarity in terms of occurrence, amplitude and timing of hydrological events, *Hydrology and Earth System Sciences*, 15, 877–896, <https://doi.org/10.5194/hess-15-877-2011>, 2011.
- Euser, T., Winsemius, H. C., Hrachowitz, M., Fenicia, F., Uhlenbrook, S., and Savenije, H. H. G.: A framework to assess the realism of model structures using hydrological signatures, *Hydrology and Earth System Sciences*, 17, 1893–1912, [https://doi.org/10.5194/hess-17-](https://doi.org/10.5194/hess-17-1893-2013)
790 1893-2013, 2013.
- Fowler, K., Peel, M., Western, A., and Zhang, L.: Improved Rainfall-Runoff Calibration for Drying Climate: Choice of Objective Function, *Water Resources Research*, 54, 3392–3408, <https://doi.org/10.1029/2017WR022466>, 2018.
- Fowler, K. J. A., Acharya, S. C., Addor, N., Chou, C., and Peel, M. C.: CAMELS-AUS: hydrometeorological time series and landscape
795 attributes for 222 catchments in Australia, *Earth System Science Data*, 13, 3847–3867, <https://doi.org/10.5194/essd-13-3847-2021>, 2021.
- Färber, C., Plessow, H., Mischel, S. A., Kratzert, F., Addor, N., Shalev, G., and Looser, U.: GRDC-Caravan: extending Caravan with data from the Global Runoff Data Centre, *Earth System Science Data*, 17, 4613–4625, <https://doi.org/10.5194/essd-17-4613-2025>, 2025.
- Garcia, F., Folton, N., and Oudin, L.: Which objective function to calibrate rainfall-runoff models for low-flow index simulations?, *Hydrological Sciences Journal*, 62, 1149–1166, <https://doi.org/10.1080/02626667.2017.1308511>, 2017.
- 800 Gnann, S. J., Coxon, G., Woods, R. A., Howden, N. J., and McMillan, H. K.: TOSSH: A Toolbox for Streamflow Signatures in Hydrology, *Environmental Modelling and Software*, 138, 104 983, <https://doi.org/10.1016/j.envsoft.2021.104983>, 2021.
- Gupta, H. V., Wagener, T., and Liu, Y.: Reconciling theory with observations: elements of a diagnostic approach to model evaluation, *Hydrological Processes*, 22, 3802–3813, <https://doi.org/10.1002/hyp.6989>, 2008.
- Gupta, H. V., Kling, H., Yilmaz, K. K., and Martinez, G. F.: Decomposition of the mean squared error and NSE performance criteria:
805 Implications for improving hydrological modelling, *Journal of Hydrology*, 377, 80–91, <https://doi.org/10.1016/j.jhydrol.2009.08.003>, 2009.
- Gupta, H. V., Perrin, C., Blöschl, G., Montanari, A., Kumar, R., Clark, M., and Andréassian, V.: Large-sample hydrology: a need to balance depth with breadth, *Hydrology and Earth System Sciences*, 18, 463–477, <https://doi.org/10.5194/hess-18-463-2014>, 2014.
- Hallouin, T., Bruen, M., and O'Loughlin, F. E.: Calibration of hydrological models for ecologically relevant streamflow predictions: a
810 trade-off between fitting well to data and estimating consistent parameter sets?, *Hydrology and Earth System Sciences*, 24, 1031–1054, <https://doi.org/10.5194/hess-24-1031-2020>, 2020.
- Hansen, N., Müller, S. D., and Koumoutsakos, P.: Reducing the Time Complexity of the Derandomized Evolution Strategy with Covariance Matrix Adaptation (CMA-ES), *Evolutionary Computation*, 11, 1–18, <https://doi.org/10.1162/106365603321828970>, 2003.
- Hernandez-Suarez, J. S., Nejadhashemi, A. P., Kropp, I. M., Abouali, M., Zhang, Z., and Deb, K.: Evaluation of the impacts of
815 hydrologic model calibration methods on predictability of ecologically-relevant hydrologic indices, *Journal of Hydrology*, 564, 758–772, <https://doi.org/10.1016/j.jhydrol.2018.07.056>, 2018.

- Jackson, E. K., Roberts, W., Nelsen, B., Williams, G. P., Nelson, E. J., and Ames, D. P.: Introductory overview: Error metrics for hydrologic modelling – A review of common practices and an open source library to facilitate use and adoption, *Environmental Modelling & Software*, 119, 32–48, <https://doi.org/10.1016/j.envsoft.2019.05.001>, 2019.
- 820 Jacobs, B., Tobi, H., and Hengeveld, G. M.: Linking error measures to model questions, *Ecological Modelling*, 487, 110562, <https://doi.org/10.1016/j.ecolmodel.2023.110562>, 2024.
- Jakeman, A., Letcher, R., and Norton, J.: Ten iterative steps in development and evaluation of environmental models, *Environmental Modelling & Software*, 21, 602–614, <https://doi.org/10.1016/j.envsoft.2006.01.004>, 2006.
- Jakeman, A. J., Sawah, S. E., Cuddy, S., Robson, B., McIntyre, N., and Cook, F.: Good Modelling Practice Principles, Tech. rep., 2018.
- 825 Kiraz, M., Coxon, G., and Wagener, T.: A Signature-Based Hydrologic Efficiency Metric for Model Calibration and Evaluation in Gauged and Ungauged Catchments, *Water Resources Research*, 59, <https://doi.org/10.1029/2023WR035321>, 2023.
- Kirchner, J. W.: Getting the right answers for the right reasons: Linking measurements, analyses, and models to advance the science of hydrology, *Water Resources Research*, 42, <https://doi.org/10.1029/2005WR004362>, 2006.
- Klingler, C., Schulz, K., and Herrnegger, M.: LamaH-CE: LARge-SaMple DATA for Hydrology and Environmental Sciences for Central Europe, *Earth System Science Data*, 13, 4529–4565, <https://doi.org/10.5194/essd-13-4529-2021>, 2021.
- 830 Knoben, W. J. M.: Setting expectations for hydrologic model performance with an ensemble of simple benchmarks, *Hydrological Processes*, 38, <https://doi.org/10.1002/hyp.15288>, 2024.
- Knoben, W. J. M., Woods, R. A., and Freer, J. E.: A Quantitative Hydrological Climate Classification Evaluated With Independent Streamflow Data, *Water Resources Research*, 54, 5088–5109, <https://doi.org/10.1029/2018WR022913>, 2018.
- 835 Knoben, W. J. M., Freer, J. E., Fowler, K. J. A., Peel, M. C., and Woods, R. A.: Modular Assessment of Rainfall–Runoff Models Toolbox (MARRMoT) v1.2: an open-source, extendable framework providing implementations of 46 conceptual hydrologic models as continuous state-space formulations, *Geoscientific Model Development*, 12, 2463–2480, <https://doi.org/10.5194/gmd-12-2463-2019>, 2019a.
- Knoben, W. J. M., Freer, J. E., and Woods, R. A.: Technical note: Inherent benchmark or not? Comparing Nash–Sutcliffe and Kling–Gupta efficiency scores, *Hydrology and Earth System Sciences*, 23, 4323–4331, <https://doi.org/10.5194/hess-23-4323-2019>, 2019b.
- 840 Knoben, W. J. M., Freer, J. E., Peel, M. C., Fowler, K. J. A., and Woods, R. A.: A Brief Analysis of Conceptual Model Structure Uncertainty Using 36 Models and 559 Catchments, *Water Resources Research*, 56, <https://doi.org/10.1029/2019wr025975>, 2020.
- Kratzert, F., Nearing, G., Addor, N., Erickson, T., Gauch, M., Gilon, O., Gudmundsson, L., Hassidim, A., Klotz, D., Nevo, S., Shalev, G., and Matias, Y.: Caravan - A global community dataset for large-sample hydrology, *Scientific Data*, 10, <https://doi.org/10.1038/s41597-023-01975-w>, 2023.
- 845 Krause, P., Boyle, D. P., and Bäse, F.: Comparison of different efficiency criteria for hydrological model assessment, *Advances in Geosciences*, 5, 89–97, <https://doi.org/10.5194/adgeo-5-89-2005>, 2005.
- Mai, J.: Ten strategies towards successful calibration of environmental models, *Journal of Hydrology*, 620, 129414, <https://doi.org/10.1016/j.jhydrol.2023.129414>, 2023.
- McMillan, H.: Linking hydrologic signatures to hydrologic processes: A review, *Hydrological Processes*, 34, 1393–1409, <https://doi.org/10.1002/hyp.13632>, 2020.
- 850 Melsen, L. A., Teuling, A. J., Torfs, P. J., Zappa, M., Mizukami, N., Mendoza, P. A., Clark, M. P., and Uijlenhoet, R.: Subjective modeling decisions can significantly impact the simulation of flood and drought events, *Journal of Hydrology*, 568, 1093–1104, <https://doi.org/10.1016/j.jhydrol.2018.11.046>, 2019.

Melsen, L. A., Puy, A., Torfs, P. J. J. F., and Saltelli, A.: The rise of the Nash-Sutcliffe efficiency in hydrology, *Hydrological Sciences Journal*, 70, 1248–1259, <https://doi.org/10.1080/02626667.2025.2475105>, 2025.

Mendoza, P. A., Clark, M. P., Mizukami, N., Gutmann, E. D., Arnold, J. R., Brekke, L. D., and Rajagopalan, B.: How do hydrologic modeling decisions affect the portrayal of climate change impacts?, *Hydrological Processes*, 30, 1071–1095, <https://doi.org/10.1002/hyp.10684>, 2016.

Mizukami, N., Rakovec, O., Newman, A. J., Clark, M. P., Wood, A. W., Gupta, H. V., and Kumar, R.: On the choice of calibration metrics for “high-flow” estimation using hydrologic models, *Hydrology and Earth System Sciences*, 23, 2601–2614, <https://doi.org/10.5194/hess-23-2601-2019>, 2019.

Nash, J. and Sutcliffe, J.: River flow forecasting through conceptual models part I — A discussion of principles, *Journal of Hydrology*, 10, 282–290, [https://doi.org/10.1016/0022-1694\(70\)90255-6](https://doi.org/10.1016/0022-1694(70)90255-6), 1970.

Nemri, S. and Kinnard, C.: Comparing calibration strategies of a conceptual snow hydrology model and their impact on model performance and parameter identifiability, *Journal of Hydrology*, 582, 124 474, <https://doi.org/10.1016/j.jhydrol.2019.124474>, 2020.

Pool, S., Vis, M. J. P., Knight, R. R., and Seibert, J.: Streamflow characteristics from modeled runoff time series – importance of calibration criteria selection, *Hydrology and Earth System Sciences*, 21, 5443–5457, <https://doi.org/10.5194/hess-21-5443-2017>, 2017.

Pool, S., Vis, M., and Seibert, J.: Evaluating model performance: towards a non-parametric variant of the Kling-Gupta efficiency, *Hydrological Sciences Journal*, 63, 1941–1953, <https://doi.org/10.1080/02626667.2018.1552002>, 2018.

Pushpalatha, R., Perrin, C., Moine, N. L., and Andréassian, V.: A review of efficiency criteria suitable for evaluating low-flow simulations, *Journal of Hydrology*, 420–421, 171–182, <https://doi.org/10.1016/j.jhydrol.2011.11.055>, 2012.

Santos, L., Thirel, G., and Perrin, C.: Technical note: Pitfalls in using log-transformed flows within the KGE criterion, *Hydrology and Earth System Sciences*, 22, 4583–4591, <https://doi.org/10.5194/hess-22-4583-2018>, 2018a.

Santos, L., Thirel, G., and Perrin, C.: Technical note: Pitfalls in using log-transformed flows within the KGE criterion, *Hydrology and Earth System Sciences*, 22, 4583–4591, <https://doi.org/10.5194/hess-22-4583-2018>, 2018b.

Schaeffli, B. and Gupta, H. V.: Do Nash values have value?, *Hydrological Processes*, 21, 2075–2080, <https://doi.org/10.1002/hyp.6825>, 2007.

Schwemmler, R., Demand, D., and Weiler, M.: Technical note: Diagnostic efficiency – specific evaluation of model performance, *Hydrology and Earth System Sciences*, 25, 2187–2198, <https://doi.org/10.5194/hess-25-2187-2021>, 2021.

Seibert, J.: On the need for benchmarks in hydrological modelling, *Hydrological Processes*, 15, 1063–1064, <https://doi.org/10.1002/hyp.446>, 2001.

Seiller, G., Roy, R., and Anctil, F.: Influence of three common calibration metrics on the diagnosis of climate change impacts on water resources, *Journal of Hydrology*, 547, 280–295, <https://doi.org/10.1016/j.jhydrol.2017.02.004>, 2017.

Thirel, G., Santos, L., Delaigue, O., and Perrin, C.: On the use of streamflow transformations for hydrological model calibration, *Hydrology and Earth System Sciences*, 28, 4837–4860, <https://doi.org/10.5194/hess-28-4837-2024>, 2024.

Trotter, L. and Knoben, W. J. M.: MARRMoT v2.1, <https://doi.org/10.5281/zenodo.6484372>, 2022.

Trotter, L., Knoben, W. J. M., Fowler, K. J. A., Saft, M., and Peel, M. C.: Modular Assessment of Rainfall–Runoff Models Toolbox (MARRMoT) v2.1: an object-oriented implementation of 47 established hydrological models for improved speed and readability, *Geoscientific Model Development*, 15, 6359–6369, <https://doi.org/10.5194/gmd-15-6359-2022>, 2022.

van Werkhoven, K., Wagener, T., Reed, P., and Tang, Y.: Characterization of watershed model behavior across a hydroclimatic gradient, *Water Resources Research*, 44, <https://doi.org/10.1029/2007WR006271>, 2008.

- Vis, M., Knight, R., Pool, S., Wolfe, W., and Seibert, J.: Model Calibration Criteria for Estimating Ecological Flow Characteristics, *Water*, 7, 2358–2381, <https://doi.org/10.3390/w7052358>, 2015.
- Wagener, T. and Montanari, A.: Convergence of approaches toward reducing uncertainty in predictions in ungauged basins, *Water Resources Research*, 47, 2010WR009469, <https://doi.org/10.1029/2010WR009469>, 2011.
- 895 Wavren, R. v., Groot, S., Scholten, H., Geer, F. v., Wösten, J., Koeze, R., and Noort, J.: Good Modelling Practice Handbook, Tech. rep., 1999.
- Yapo, P. O., Gupta, H. V., and Sorooshian, S.: Multi-objective global optimization for hydrologic models, *Journal of Hydrology*, 204, 83–97, [https://doi.org/10.1016/S0022-1694\(97\)00107-8](https://doi.org/10.1016/S0022-1694(97)00107-8), 1998.
- Yilmaz, K. K., Gupta, H. V., and Wagener, T.: A process-based diagnostic approach to model evaluation: Application to the NWS distributed
900 hydrologic model, *Water Resources Research*, 44, 2007WR006716, <https://doi.org/10.1029/2007WR006716>, 2008.

Table 2. Equations of the eight performance metrics used as objective functions in this study. Numbers (1–8) correspond to metric references in the text.

No.	Metric	Equation
(1)	Nash-Kling-Sutcliffe Efficiency (NSE) <u>Gupta Efficiency (KGE)</u>	$NSE = 1 - \frac{\sum_{t=1}^n (Q_{o,t} - Q_{s,t})^2}{\sum_{t=1}^n (Q_{o,t} - \bar{Q}_o)^2} \quad KGE = 1 - \sqrt{\left(\frac{\bar{Q}_s}{\bar{Q}_o} - 1\right)^2 + \left(\frac{\sigma_s}{\sigma_o} - 1\right)^2 + (r_{pe} - 1)^2}$
(2)	Kling-Nash-Gupta Efficiency (KGE) <u>Sutcliffe Efficiency</u> (NSE)	$KGE = 1 - \sqrt{\left(\frac{\bar{Q}_s}{\bar{Q}_o} - 1\right)^2 + \left(\frac{\sigma_s}{\sigma_o} - 1\right)^2 + (r_{pe} - 1)^2} \quad NSE = 1 - \frac{\sum_{t=1}^n (Q_{o,t} - Q_{s,t})^2}{\sum_{t=1}^n (Q_{o,t} - \bar{Q}_o)^2}$
(3)	log-NSE <u>Power-transformed KGE (KGE0.2)</u>	$\log NSE = 1 - \frac{\sum_{t=1}^n (\log(Q_{o,t}) - \log(Q_{s,t}))^2}{\sum_{t=1}^n (\log(Q_{o,t}) - \log(\bar{Q}_o))^2} \quad KGE_{0.2} = 1 - \sqrt{\left(\frac{\bar{Q}_s^{0.2}}{\bar{Q}_o^{0.2}} - 1\right)^2 + \left(\frac{\sigma_{0.2,s}}{\sigma_{0.2,o}} - 1\right)^2 + (r_{pe,0} - 1)^2}$
(4)	log-KGE <u>logarithmic NSE (log NSE)</u>	$\log KGE = 1 - \sqrt{\left(\frac{\log(\bar{Q}_s)}{\log(\bar{Q}_o)} - 1\right)^2 + \left(\frac{\sigma_{\log,s}}{\sigma_{\log,o}} - 1\right)^2 + (r_{p,\log} - 1)^2}$ $\tilde{Q}_{o,t} = \log(Q_{o,t} + \epsilon), \quad \tilde{Q}_{s,t} = \log(Q_{s,t} + \epsilon), \quad \epsilon > 0$ $\log NSE = 1 - \frac{\sum_{t=1}^n (\tilde{Q}_{o,t} - \tilde{Q}_{s,t})^2}{\sum_{t=1}^n (\tilde{Q}_{o,t} - \bar{\tilde{Q}}_o)^2}$
(5)	Non-parametric (KGE-NP)	KGE $KGE_{NP} = 1 - \sqrt{\left(\frac{\bar{Q}_s}{\bar{Q}_o} - 1\right)^2 + \left(1 - \frac{1}{2} \left \frac{Q_s(I(k))}{n\bar{Q}_s} - \frac{Q_o(J(k))}{n\bar{Q}_o} \right \right)^2 + (r_{sp} - 1)^2}$
(6)	KGE-Split <u>KGE-Split</u>	$KGE_Split = \frac{1}{Y} \sum_{y=1}^Y \left(1 - \sqrt{\left(\frac{\bar{Q}_{s,y}}{\bar{Q}_{o,y}} - 1\right)^2 + \left(\frac{\sigma_{s,y}}{\sigma_{o,y}} - 1\right)^2 + (r_{pe,y} - 1)^2} \right)$ $KGE_Split = \frac{1}{Y} \sum_{y=1}^Y \left(1 - \sqrt{\left(\frac{\bar{Q}_{s,y}}{\bar{Q}_{o,y}} - 1\right)^2 + \left(\frac{\sigma_{s,y}}{\sigma_{o,y}} - 1\right)^2 + (r_{pe,y} - 1)^2} \right)$
(7)	Signature-based Efficiency (SHE)	Hydrologic $SHE_g = 1 - \sqrt{\left(\frac{\bar{Q}_s/\bar{P}_o}{\bar{Q}_o/\bar{P}_o} - 1\right)^2 + \left(\frac{\sigma_s/\sigma_p}{\sigma_o/\sigma_p} - 1\right)^2 + (r_{sp} - 1)^2}$ $SHE = 1 - \sqrt{\left(\frac{\bar{Q}_s/\bar{P}_o}{\bar{Q}_o/\bar{P}_o} - 1\right)^2 + \left(\frac{\sigma_s/\sigma_p}{\sigma_o/\sigma_p} - 1\right)^2 + (r_{sp} - 1)^2}$
(8)	Diagnostic Efficiency (DE)	$DE_{OF} = \sqrt{\left(\frac{1}{n} \sum_{t=1}^n \frac{Q_{s,t} - Q_{o,t}}{Q_{o,t}}\right)^2 + \left(\int_0^1 \left \frac{Q_s^\downarrow(p) - Q_o^\downarrow(p)}{Q_o^\downarrow(p)} - \frac{\bar{Q}_s - \bar{Q}_o}{\bar{Q}_o} \right dp\right)^2 + (r_{pe} - 1)^2}$

Table 3. Overview of the three components of the eight metrics

Metric	Bias	Variability	Correlation
KGE	$\frac{\overline{Q_s}}{\overline{Q_o}}$	$\frac{\sigma_s}{\sigma_o}$	Pearson
NSE	$\frac{\overline{Q_s - Q_o}}{\sigma_o}$	$\frac{\sigma_s}{\sigma_o}$	$2 \cdot \text{Pearson} \cdot \frac{\sigma_s}{\sigma_o}$
log KGE <u>KGE0.2</u>	as KGE	as KGE	as KGE
log NSE	as NSE	as NSE	as NSE
KGE-NP	as KGE	FDC-based	Spearman
KGE-Split	as KGE	as KGE	as KGE
SHE	as KGE	as KGE	Spearman
DE	Bias FDC-based	Residuals FDC-based	Spearman <u>Pearson</u>

Table 5. Overview of selected signatures.

Signature	Abbreviation	Category	Units
Slope of Flow Duration Curve <u>(33-66)</u>	FDC Slope	Streamflow <u>Flow Variability</u>	-
5 th Streamflow Percentile	Q5	Streamflow <u>Low Flow Magnitude</u>	mm/day
95 th Streamflow Percentile	Q95	Streamflow <u>High Flow Magnitude</u>	mm/day
High <u>Low</u> Flow Frequency	HF <u>Freq</u> LFF	Streamflow <u>Low Flow Occurrence</u>	-
Low <u>High</u> Flow Frequency	LF <u>Freq</u> HFF	Streamflow <u>High Flow Occurrence</u>	-
High <u>Low</u> Flow Duration	HF <u>Dur</u> LFD	Streamflow <u>Low Flow Duration</u>	days/event
Low <u>High</u> Flow Duration	LF <u>Dur</u> HFD	Streamflow <u>High Flow Duration</u>	days/event
Mean Half Flow Date	MHFD	Streamflow <u>Flow Timing</u>	day of year
Total Runoff Ratio	Total RR / <u>TRR</u>	Water Balance	-
Event Runoff Ratio	Event RR / <u>ERR</u>	Partitioning/Connectivity	-
Baseflow Index	BFI	Baseflow	-
Baseflow Recession Coefficient	BFRC	Baseflow	1/day
Flashiness Index	FI	Water <u>Storage</u> Runoff <u>Dynamics</u>	-
Variability Index	VI	Water Storage	-
Rising Limb Density	RLD	Channel <u>Processes</u> Runoff <u>Dynamics</u>	rises/day

Table 6. Recommendations for objective functions that can represent the 15 considered signatures well according to our study results. Where a signature was not significantly affected by OF choice, we report “not significant”.

Signature	Best OF
Total Runoff Ratio	KGE-NP, <u>KGE</u>
Event Runoff Ratio	SHE , KGE-NP, <u>NSE</u>
Mean Half Flow Date	Not significant
FDC Slope	Not significant
5th Flow Percentile (Q5)	log KGE , DE
Low-Flow Duration	Not significant
Low-Flow Frequency	DE, log KGE <u>KGE0.2</u>
Baseflow Index (BFI) KGE-NP, DE 95th Flow Percentile (Q95)	KGE, SHE
High-Flow Duration	Not significant
High-Flow Frequency	Not significant
Baseflow <u>Index (BFI)</u>	<u>KGE-NP</u>
<u>Baseflow</u> Recession Coefficient (BFRC)	SHE, NSE, DE , <u>KGE-NP</u>
Rising Limb Density (RLD)	Not significant
Flashiness Index (FI)	Split-KGE, SHE <u>KGE-Split</u>
Variability Index (VI)	DE (where relevant) <u>Not significant</u>