

ECOGRAPHY

Review

A practical guide to species trend detection with unstructured data using local frequency scaling (Frescalo)

Romain Goury¹✉, Diana E. Bowler², Colin Harrower², Tamara Münkemüller¹, Jeanne Vallet³, Jon M. Yearsley⁴, Wilfried Thuiller¹ and Oliver L. Pescott²

¹Université Grenoble Alpes, Université Savoie Mont Blanc, CNRS, LECA, Grenoble, France

²Biodiversity Monitoring and Analysis, UK Centre for Ecology and Hydrology, Wallingford, UK

³Muséum national d'histoire naturelle, Conservatoire Botanique National du Bassin Parisien (MNHN/CBNBP), Paris, France

⁴School of Biology and Environmental Science, University College Dublin, Dublin, Ireland

Correspondence: Romain Goury (romain.goury@univ-grenoble-alpes.fr)

Ecography

2026: e08270

doi: [10.1002/ecog.08270](https://doi.org/10.1002/ecog.08270)

Subject Editor:

Michael Krabbe Borregaard

Editor-in-Chief:

Jens-Christian Svenning

Accepted 5 March 2026



Accurately measuring biodiversity change remains a central challenge in ecology. Beyond the general idea of quantifying temporal species frequency changes, several sampling-related biases in data collection remain key methodological challenges to consider. Long-term standardized ecological data are rare, and most available datasets exhibit considerable spatial and temporal variation in sampling effort (i.e. unstructured data). Among the available methods, the local frequency scaling approach (Frescalo) has proven particularly effective at addressing these biases. By applying successive spatial and temporal corrections, Frescalo leverages emergent patterns in species assemblages to correct for variation in survey effort. Compared to other similar approaches, Frescalo is particularly well suited to long-term datasets and those with a high number of species. It is also a versatile method, allowing simultaneous estimation of temporal and spatial changes, or even providing diagnostics for survey design or bias assessment. The method's technical complexity, the level of ecological knowledge required, and the challenges of implementation raise a number of practical issues in its application. In this paper, we present a clear and accessible explanation of the Frescalo methodology, offer a step-by-step roadmap to guide users, and highlight the wide range of applications it supports. To further facilitate its adoption, we also introduce an R package designed to simplify implementation.

Keywords: biodiversity estimation, Frescalo, neighbourhood, sampling effort, species richness, unstructured data

Introduction

In recent years the importance of monitoring biodiversity and assessing long-term species trends has grown significantly (Dornelas et al. 2023), particularly in the context



www.ecography.org

© 2026 The Author(s). Ecography published by John Wiley & Sons Ltd on behalf of Nordic Society Oikos

This is an open access article under the terms of the Creative Commons Attribution License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited.

of the global IPBES assessments and the Kunming–Montreal Global Biodiversity Framework. Although frameworks for quantifying and attributing temporal biodiversity changes have recently emerged (Gonzalez et al. 2023), adapted statistical models have become essential for estimating long-term biodiversity changes. These models are fundamental for providing reliable trend estimates, ideally accompanied by measures of uncertainty (Pescott et al. 2022). Such estimates are essential for making reliable diagnoses of biodiversity status and understanding the underlying causes of observed trends (attribution; Gonzalez et al. 2023, Grace 2024). However, there is a growing need to expand the range of statistical models available to ecologists for estimating trends, especially given the diversity of monitoring data sources.

Monitoring data vary widely in terms of design and coordination, which has significant implications for extracting reliable trend information. Structured monitoring schemes, such as breeding bird surveys, which follow strict sampling protocols and are repeated regularly, provide the most robust data for trend estimation and allow for the use of relatively simple statistical modelling, assuming good coverage of the statistical population. However, these schemes are rare and typically limited to well-studied taxa (e.g. birds and butterflies) or to well-studied areas (e.g. Global North). In contrast, most available ecological data are unstructured or semi-structured (i.e. with associated metadata on survey methods), meaning they have been collected through diverse and often undocumented methods with varying sampling effort across time and space. This heterogeneity introduces variability in species detection and identification (Geldmann et al. 2016), statistical population coverage (Boyd et al. 2023, 2024b), and reporting, posing challenges for trend analysis. Unstructured (and semi-structured) data, such as historical museum collections (e.g. herbaria; Rich 2006), distribution atlases (Stroh et al. 2023), and aggregated species occurrence records databases (e.g. GBIF) offer considerable potential for assessing biodiversity change. These unstructured (and semi-structured) datasets are often in the form of presence-only data, span long temporal periods (sometimes centuries), and cover a broad range of taxa beyond those monitored by structured schemes. Despite their value, these datasets are arguably underutilized for estimating biodiversity trends.

A key challenge in using unstructured data is the variation in recording effort (e.g. the number of visits per site, time spent during each visit) and behavior (e.g. which sites are visited, what is reported). Additionally, unstructured data often lack metadata about survey method, effort, or target species, which makes it difficult to model this variation (Pescott et al. 2015, Kelling et al. 2019). Data availability patterns also reveal systematic taxonomic and geographic biases (Boakes et al. 2010, Troudet et al. 2017). Taken together, these issues mean that the true ecological patterns can be confounded by recorder strategies (Dobson et al. 2020) or related issues, such as historical data curation practices (Pescott et al. 2019). If no attempt is made to control for this, any trend estimates using simple statistical models of

the data would be biased and could lead to poor conservation decisions (Boyd et al. 2023).

To address these challenges, several statistical methods have been developed to account for biases in unstructured data and provide reliable estimates of species trends. Among a different set of methods, including a ‘naive’ linear mixed model approach and the Telfer model (Telfer et al. 2002), Isaac et al. (2014) identified two methods that appeared to have the most potential for detecting trends under different bias scenarios: the occupancy-detection model (MacKenzie et al. 2002) and Frescalo (‘FREquency SCALing LOcal’; Hill 2012). These methods use distinct approaches that vary in their assumptions, making them suitable for different contexts (Pescott 2026). Occupancy-detection models assume that a species’ true occupancy state remains stable over a defined closure period, allowing detectability to be estimated from repeated surveys (MacKenzie et al. 2002). This assumption weakens at larger spatial or temporal scales. Moreover, these models only correct for detectability differences among survey visits and do not account for spatial sampling bias in site selection (MacKenzie et al. 2002). In contrast, Frescalo operates at larger spatio-temporal scales (e.g. multi-year/100 km²; Hill 2012) and uses information on emergent patterns in observed species assemblages, rather than estimating parameters related to observation processes at the level of individual surveys. While occupancy models exploit repeated observations through time, Frescalo uses observations across multiple spatial locations. By attempting to account for both spatial and temporal variation in sampling coverage at larger scales, Frescalo is particularly well-suited to handling biases inherent to unstructured datasets (Geldmann et al. 2016, Binley and Bennett 2023). Frescalo can be seen as building on several preceding approaches designed to remove the biasing effects of variable survey effort, including the Telfer index (Telfer et al. 2002) and the use of ‘benchmark species’ to index effort (Pescott et al. 2019). The Telfer index used an overall adjustment for effort between two survey periods, meaning that it was sensitive to mismatches between the spatial distribution of this changing effort and true changes in distribution. Similarly, early uses of benchmark species (Maes and Swaay 1997) were across entire (statistical) populations, rather than local areas. Frescalo develops both of these approaches by allowing for local variation in confounding between effort and true change (Pescott et al. 2019).

While Frescalo has primarily been applied to plant data (Bijlsma 2013, Blockeel et al. 2014, Pescott et al. 2019, White et al. 2019, Eichenberg et al. 2021, Auffret and Svenning 2022, Suggitt et al. 2023), there is growing interest in extending its use to other taxa, such as moths (Fox et al. 2014), pollinators (Redhead et al. 2018), and multi-taxa analyses (Dyer et al. 2017, Montràs-Janer et al. 2024). This is supported by the fact that Frescalo only uses information within the biodiversity data, and does not rely on associated sampling effort metadata; this means that it has relatively low data requirements. Especially, Frescalo requires presence-only occurrence records, such as Atlas or GBIF-type data,

including historical museum data or similar, although richer data (e.g. presence/absence) could of course be degraded for inclusion if this was deemed appropriate but losing the information of the true absences (see ‘Assumptions’ for detail on these points). Similarly, data that are relatively coarsely resolved in space and/or time are also suitable for use with the method, a feature that can also intrinsically act to reduce bias (Stroh et al. 2023, Boyd et al. 2024a).

The complexity of Frescalo, which involves the use of site neighbourhoods to predict species assemblage properties and a multi-step algorithm grounded in ecological theory, may explain its limited adoption among ecologists relative to other methods. The approach requires knowledge of taxon group ecology, dataset properties, and an understanding of the underlying ecological theory. To make this method more accessible, we present a guide to Hill’s Frescalo method tailored to ecologists who may find the original presentation too mathematical (Hill 2012). This article begins with an intuitive and concise explanation of the method’s key principles, accompanied by conceptual figures to aid understanding. We provide a step-by-step roadmap to help readers apply the method effectively, identify potential pitfalls, and make informed decisions. We have also developed an R package to facilitate the use of Frescalo (<https://github.com/colinharrower/frescalo>) based on the efficient parallelised implementation of White et al. (2019); a direct (i.e. loop-based) R translation of the original Fortran implementation is also provided by Pescott (2025). Mathematical notation throughout the article follows Hill (2012) for easy cross-referencing, and a Glossary is provided for plain language descriptions of model parameters.

Frescalo in a nutshell

The frequency scaling using local occupancy approach of Hill (2012) can provide an unbiased estimate of temporal trends when there has been enough sampling to at least estimate species’ local relative frequencies fairly accurately. Spatial variation in species’ ecologies is addressed by dividing the study area into ‘neighbourhoods’, whereby each site in the analysis is assigned a number of other similar sites nearby that provide an ecologically coherent context within which to understand a target site’s species assemblage. Within this context Frescalo consists of two main steps: the first step is to correct for variation in sampling effort across neighbourhoods for the *overall* time periods being considered (‘Spatial correction’, Fig. 1). This step involves aligning all neighbourhood species frequency curves, which correspond to the distribution of species ranked by their site frequencies (e.g. Fig. 2B). A method of standardising species frequency curves across neighbourhoods is used to ensure their comparability for subsequent steps in the algorithm. The second step is to correct for time-period specific variations in recording effort within and across sites. This temporal variation in effort is accounted for using an index of local recording completeness: the proportion of a suite of locally common species, referred to as ‘benchmark species’, that have been recorded (‘Temporal correction’, Fig. 1). While other methods have also used benchmarks

for similar purposes (Pescott et al. 2019), Frescalo allows the identity and number of these taxa to vary regionally (i.e. by neighbourhood). This ensures that each area has its own species frequency curve and benchmarks, accounting for regional variation in ecology and sampling effort. Section ‘Frequency scaling using local occupancy: a deeper dive’ provides more detail on each of these steps.

Frequency scaling using local occupancy: a deeper dive

Spatial correction

Frescalo models the data in discrete space at two different scales: sites (these may be grid cells or any set of non-overlapping polygons of equal or roughly similar size) and neighbourhoods, which correspond to a number of sites, typically in the order of tens, grouped together (these sites need not all be contiguous, but neighbours will generally need to be geographically close to and ecologically similar to the target site for the assumption that they provide information on a target site’s species assemblage to be reasonable). To predict the species frequency curve for each site Frescalo aggregates the data for the ecological ‘neighbourhood’ for each site, using data for all time points included in the analysis, based on a predefined set of neighbourhood site weights (section ‘Neighbourhoods weights and frequency-weighted mean frequencies’). The resulting species frequency curves for each site are therefore constructed in the context of their associated neighbourhoods (section ‘Standardising local species frequencies’). (Note that temporal information is not used at this stage.)

Neighbourhoods weights and frequency-weighted mean frequencies

Neighbourhoods are a key part of local frequency scaling. As already noted, a neighbourhood is a cluster of sites that are ‘similar’ to a target site. The simplest definition of ‘similar’ is to use geographical distance, where only sites within a set distance from the target site are included in its neighbourhood. Ecological similarity can also be used, for example floristic similarity (Hill 2012), climatic similarity (Auffret and Svenning 2022), and/or land cover similarity (Eichenberg et al. 2021, Stroh et al. 2023). Each site is assigned its own unique neighbourhood, although neighbourhoods for different target sites may overlap and typically do. The definition of site neighbourhoods is very important as it underpins the generation of the large-scale species frequency curves that constitute one of the key propositions of local frequency scaling. The more ecologically defensible the neighbourhoods, the more confident we may be that their species frequency curves will provide useful information about the target sites that they surround.

The first step in the method, which we provide a worked example for below and in Fig. 2, is to calculate the frequency of each species j in the neighbourhood of each site i (f_{ij} , see Glossary). This means that species that are more frequent and/or in neighbourhood locations with larger weights, i.e.

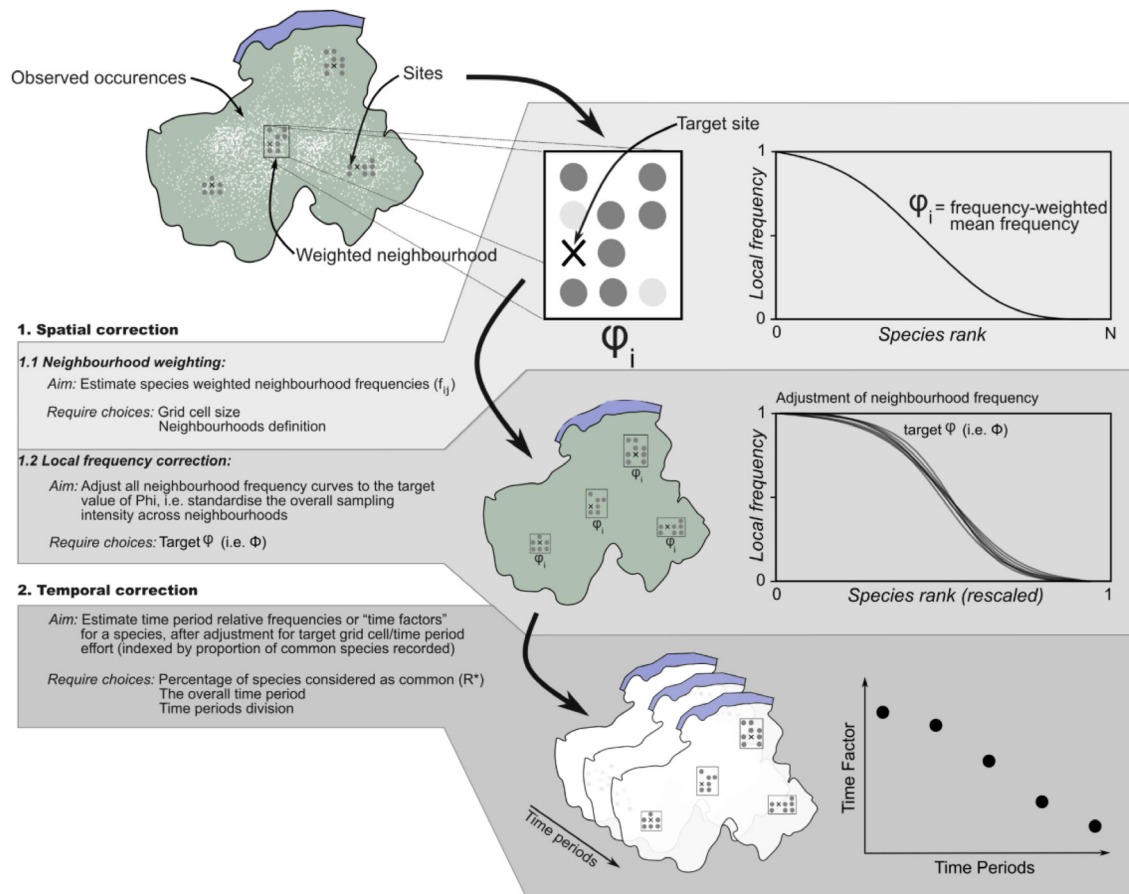


Figure 1. An overview of frequency scaling using local occupancy (Frescalo). The approach is based on two consecutive steps: 1) Spatial correction. This consists in first defining a neighbourhood per site and then standardising the local species frequency curves across neighbourhoods to make them comparable. 2) Temporal correction. This consists of equalising the sum of the observed species occurrences with the sum of the standardised (i.e. effort-corrected) neighbourhood frequencies across sites to obtain an average temporal deviation factor per species and time period (Hill's 'time factor') across the study area. Taken together these time factors represent a species' average temporal trend, conditional on modelling assumptions.

those that are closer and/or more ecologically similar to the target site, are emphasised relative to those species that are rare and/or in more distant or less similar sites. Neighbourhoods in which all sites are equally weighted (i.e. effectively an unweighted neighbourhood) can also be used, and this can be seen simply as a special case of the weighted situation; Fig. 2 uses such an unweighted neighbourhood for simplicity (but see the Supporting information for a weighted example). In what follows we refer to species' weighted frequencies in neighbourhoods simply as 'frequencies' (f_{ij}) to avoid confusion with the statistic later used to summarise a neighbourhood's species frequency curve, the frequency-weighted mean species frequency (below). Whether weighted or unweighted, the amount of neighbourhood 'space' occupied by any species is simply a proportion or frequency, and it is this point that is key for understanding the Frescalo method. We therefore relegate the additional complication of defining species' frequencies relative to weighted neighbourhoods to the Supporting information.

Once all species' frequencies have been calculated for a neighbourhood, a type of mean frequency can be defined: this is the (self-)weighted mean frequency, where the weights are the species frequencies themselves. Rather than treating all species equally, as would the simple arithmetic average of all species' frequencies in a neighbourhood, the frequency-weighted mean frequency treats site/species occurrences equally, meaning that the average is pulled towards higher frequencies representing commoner species. A key insight of Hill (2012) was that this frequency-weighted mean frequency could be rearranged into the ratio of the average species richness of the neighbourhood to the reciprocal of Simpson's index (i.e. Hill number 2 or N_2). Hill numbers can be interpreted as the effective number of species, and are metrics that only depend on relative abundances, not absolute numbers (Jost 2006). This provides a useful way of allowing the 'structure' of species assemblages to be compared on a common scale. In other words, the N_2 metric is influenced by the shape of the species rank-frequency curves rather than

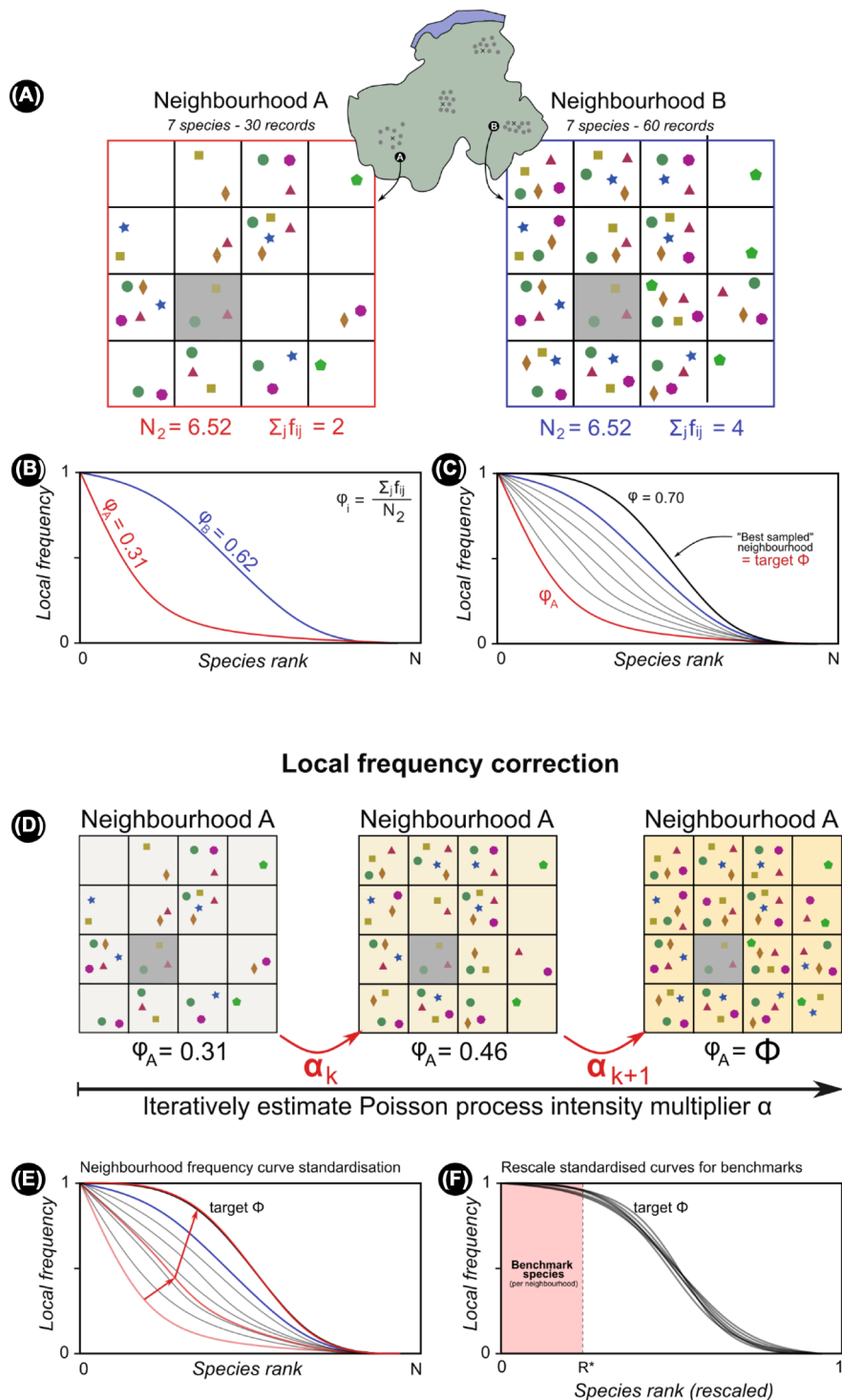


Figure 2. Spatial correction. The neighbourhood calculations are illustrated by showing two different unweighted neighbourhoods represented in red and blue, divided into different sites delimited by black grid lines (A). The target site is represented by a grey shaded square. Each small coloured symbol is a different species. For each neighbourhood A and B, the patterns of local species frequencies are plotted as rank-frequency curves, with their frequency-weighted mean local frequencies (ϕ_i) given in text (B); (C) expands this to a larger number of imaginary neighbourhoods, and depicts the local frequency curve for the well-sampled neighbourhood that acts as the adjustment target Φ in black. The local frequency standardisation is illustrated by the iterative estimation of an intensity multiplier α via its log-link for a given neighbourhood (D), (E), corresponding to the steps required to make neighbourhoods comparable (F). Note that while the x-axis maximum is given as 1 in (F), scaled ranks beyond 1 are possible when the number of observed species exceeds the number predicted.

by the frequencies' absolute magnitude; see also Section 'Standardising local species frequencies' and Pescott (2026).

To illustrate this, consider two neighbourhoods: A, which is relatively under-sampled (30 records), and B, which is well-sampled (60 records, Fig. 2A). In this example, neighbourhood A is also equivalent (in expectation) to randomly deleting half of the records in neighbourhood B. The inverse Simpson index (i.e. N_2) is the same between neighbourhoods A and B, which means that the species in each neighbourhood have the same relative frequency distribution. However, the sum of the species' observed frequencies (equivalent to the average species richness of the neighbourhood) of A is half of B, but as N_2 is unchanged the ratio ϕ_A is half of ϕ_B (Fig. 2A, B). Therefore, based on the ϕ_i values, we can say that the sampling intensity in neighbourhood A is half of the sampling intensity in neighbourhood B. However, we note that this interpretation only works if N_2 is the same in the neighbourhoods compared. The question now is: how can we ensure that observations from different large-scale neighbourhoods sampled at different intensities are comparable in this way? The method used to achieve this is described in the next section.

Standardising local species frequencies

The standardisation of local species frequencies across neighbourhoods is required to make them comparable, i.e. to align all neighbourhood species frequency curves towards that indicated by a specific value of ϕ (Fig. 2F). This target ϕ value (hereafter referred to as Φ as per Hill 2012) can be informally thought of as the ϕ value of the 'best sampled' neighbourhood (Fig. 2C), although in fact any value could be used. As Hill (2012) states, 'the precise value of Φ is not critical, but it should correspond to a thorough search of the neighbourhood'. In other words, Φ is generally set so that it is close to the largest observed value of ϕ across neighbourhoods.

Another way of understanding this process is to understand that α_k (Fig. 2D) can be seen as a multiplier that 'fattens' the pattern of each species frequency within a multivariate Poisson point process (see a visual animation of the process here: <https://github.com/sacrevet/Frequency-rescaling-demo>). If the value of α is high, the neighbourhood correction is strong and species weighted frequencies are all scaled up. However, this step does not artificially create 'new' species, but rather can be understood as increasing existing species' relative frequencies proportionally. N_2 is unchanged given that it is invariant to multiplication applied equally across species' frequencies, and the α multiplier only affects the numerator of the ratio formula for ϕ ; α then effectively acts as an index of how much 'fattening' of species' neighbourhood frequencies was required to achieve the specified frequency-weighted mean frequency Φ . This could also be interpreted as harmonising the expected frequency of species across all neighbourhoods (Fig. 2F). Returning to our example in Fig. 2, if we set Φ to 0.70 and find that ϕ_A is 0.31, we want ϕ_A to be adjusted towards 0.70. Therefore, through an iterative algorithm, each species frequency (f_{ij}) for neighbourhood A can

be interpreted as having been multiplied by some value α via its log-link, which is updated at each algorithm step k (note that this iterative process is not the only way of solving for α , but it is the approach implemented by Hill (2012)). Once all (transformed) species frequencies have been multiplied by α_k , ϕ_A is recalculated according to the formula described in words above, i.e. $\sum f_{ij}/N_2$. Since N_2 is constant, and where α_k is greater than one for under-sampled neighbourhoods, the value of ϕ_A will increase at each step until reaching Φ (although, in reality, depending on the successive approximation algorithm used, ϕ_i may also overshoot and be subsequently reduced towards Φ).

This spatial correction is applied over all time periods in the analysis taken together, i.e. the species frequency curve thus adjusted is time-independent and is assumed to represent a hypothetical 'true' all-time species frequency curve for a neighbourhood. This does not directly correct for temporal sampling bias, but it provides a neat way of subsequently decomposing this curve into two elements: one relating to the 'true' ecological state of a species within a neighbourhood at a given point in time and one relating to sampling effort.

Temporal correction

The temporal correction estimates a metric (hereafter 'time factor') for each species across time periods after the spatial corrections have been applied to all neighbourhoods; taken together, these time factors constitute a species' temporal trend. Once all neighbourhoods have been rescaled to have the frequency-weighted local mean frequency of the best sampled neighbourhood Φ , the adjusted frequency of each species (i.e. f'_{ij} see Glossary) can be extracted. These values are independent of time, and can be considered to refer to an idealised and standardised species frequency curve from which period-specific deviations can be estimated.

The standardisation of local species frequency curves also serves another purpose within Frescalo: it provides the analyst with a simple and objective way of selecting benchmark species. Recall that benchmarks are locally common species, the recording of which is assumed to be a useful index of effort (Latour and van Swaay 1992, Pescott et al. 2019). Mathematically, the benchmark species are defined for each neighbourhood as some specified proportion of the standardised species frequency curve. These neighbourhood-level benchmark estimations therefore explicitly consider variations in species assemblages dependent on spatial context. While benchmarks could in theory be chosen outside of the Frescalo algorithm, the ordered lists of species rank in neighbourhoods, R_{ij} , arising from the neighbourhood frequency curves, can be made comparable across neighbourhoods by dividing by the expected number of species in that neighbourhood $\sum f'_{ij}$. This means that R'_{ij} always runs from 0 up to (and sometimes just above) 1 regardless of neighbourhood richness. This extra normalisation step means that we can define benchmarks simply as the species with $R'_{ij} < R'$ using a single value of R' across neighbourhoods. (Note that the rank-normalised adjusted species frequencies are not used

directly elsewhere in the Frescalo algorithm.) The effort index arising from this benchmark definition is the proportion of benchmarks found in a given site in a given time period (S_{ij}), which corresponds to the number of the benchmark species found in site i at time t divided by the total number of benchmark species found in this site over all time periods. The selected benchmark species are assumed to not show substantial temporal trends on average. It is important to note that this method is applied at a large spatial scale. Consequently, a large number of species is recorded within each assemblage, allowing benchmark species to be reliably captured and represented while accounting for variation across the landscape.

Once this index of sampling effort is in place, we have a way of estimating how much of the deviation between our time-independent standardised species frequencies and our observed data can be accounted for by sampling effort: any remaining deviation is assumed to represent true ecological pattern. This is done by estimating the model-based species frequencies by combining neighbourhood level sampling intensity (f'_{ij}) and site level sampling effort (S_{ij}); these are then summed across sites for each species and time period ($\Sigma_i Q_{ijt}$). The time factor (x_{jt}) is therefore the value that adjusts the sum of modelled estimates of a species frequency to match the empirical observations (i.e. the sum of the detected occurrences by species and time period across sites, $\Sigma_i P_{ijt}$). This represents a temporal deviation for a species from its modelled frequency. Because our benchmark species have been chosen to be (on average) temporally stable, deviations of x_{jt} away from 1 quantify how much each species' frequency has changed relative to the benchmarks.

Assumptions

The recorder effort problem

The frequent lack of information about the effort that has gone into collecting biodiversity data has often hampered research. This problem was identified early on in the modern literature as the 'recorder effort problem' (Prendergast et al. 1993). A first assumption of Frescalo is that *the probability of finding a species at a site can be estimated from its frequency in the neighbourhood*, conditional on a thorough search of the area having been conducted. This assumption distinguishes Frescalo from some other modelling approaches such as occupancy-detection models (Mackenzie et al. 2002) by modelling species abundance at a coarser scale, but it is closely related to the closure hypothesis of occupancy detection models (Pescott 2026). Occupancy models typically adjust for imperfect detection using repeat visits (and often by adjusting for covariates thought to index effort expended, e.g. list length; van Strien et al. 2013) at small spatial scales, whereas Frescalo uses information (almost always arising from numerous separate visits) across neighbourhoods at large scales to estimate the true local species frequency curve (Pescott et al. 2019). This could be framed as the difference between modelling the actual data-generating process and modelling emergent patterns in aggregated data (Frank 2009).

Neighbourhood frequencies and species discovery

A key aspect of the Frescalo algorithm is transitioning from site-specific occurrences to estimated neighbourhood frequencies, relying on three key assumptions. First, *the target site must be similar to those in its neighbourhood*, requiring well-defined variables to describe site similarity among a neighbourhood (Section 'Neighbourhoods weights and frequency-weighted mean frequencies'). Additionally, neighbourhoods are assumed to remain constant over time, without accounting for temporal changes (e.g. land use modifications). The second assumption is that *species discovery can be modelled by a Poisson process*. This requires that the average rank of a species in a local frequency curve over time is not biased. That is, ranks should ideally index the truth, regardless of changes in sampling intensity. According to Hill (2012), '[t]he chance of a given species being discovered is the outcome of a two-stage stochastic process. The first stage concerns the type and duration of the visit, while the second concerns the frequency with which a given species is encountered during a visit of a certain type. When these processes are combined, each species will have a standard probability of being recorded on a visit. Under most assumptions about the nature of this two-stage process, the discovery of less common species will be a rare event', i.e. species counts in the neighbourhood behave like a Poisson process. To rephrase this assumption, we could say that the discovery process is stationary across space and time. Common species remain proportionally common even if you tweak the benchmark proportion. Indeed, the sensitivity analysis of Hill (2012) shows that varying R' over a wide range only has negligible effects on the estimated trends. Finally, *in a well-sampled neighbourhood, there is a characteristic weighted mean species frequency* which is the same for all neighbourhoods. This assumption is taken into account in the choice of the parameter Φ (section 'Standardising local species frequencies'), which specifies that a well-sampled neighbourhood corresponds to a certain weighted mean species frequency value (i.e. Φ). According to Hill (2012), 'in most applications, the exact choice of Φ is not critical' (Hill 2012 for a sensitivity analysis). Note that the ϕ_i values across neighbourhoods can be standardised without explicitly modelling areas of different richnesses. Although ϕ_i depends on the average species richness of the neighbourhood (i.e. Σf_{ij}) and the 'effective species number' (N_2), their ratio is independent of local species richness and evenness, and instead indexes sampling intensity in a standardised way (Fig. 2A). This implies that, at a given level of sampling effort (i.e. Φ), the adjusted species frequencies in each neighbourhood (f'_{ij}) are expected to be comparable (Fig. 2F).

Trend estimation

The time factor (section 'Temporal correction') represents the relative frequency of each species within its rescaled neighbourhood. Since this is a relative estimation, *benchmark species should ideally remain stable over time on average*. If all benchmark species increase at the same rate as the target species, the time factor would yield a flat trend, obscuring actual changes.

Therefore, it is essential that reference species exhibit a (flat) constant trend on average over time, otherwise the observed trend of the target species may become unreliable. Thus, the larger the number of species considered, the greater the likelihood that a subset of common species will show stable trends over time. This is particularly true of very common species, which generally structure habitats and are therefore uncorrelated with human activity. Furthermore, Frescalo is based on changes in species presence rather than abundance. While population changes are likely to be widespread, changes in species presence will be slower for abundant species, especially at a regional scale.

Frescalo: a road map

Deciding on the spatio-temporal units of the analysis

To run Frescalo, the dataset must be collapsed to presence-only data at the scale of the analysis, where each row represents a species observation linked to a geographic location and a defined time period. Applying the model therefore requires two key decisions. First, a spatial unit (which we refer to as sites) must be chosen, assigning each observation to a specific site. These units can vary, including grid-based resolutions (e.g. hectads, tetrads) or administrative boundaries (e.g. municipalities, departments). Alternatively, custom spatial units could be defined based on biogeographical knowledge for specific taxa or environments. In atlas-type datasets, the spatial unit often corresponds to the sampling grid. Second, appropriate time periods must be selected (both overall and divisions for trend calculations). Typical options range from years to broader intervals. Spatial and temporal decisions are interconnected, as spatial units should be selected to ensure sufficient data coverage across all time periods. These choices should be made in collaboration with experts familiar with the dataset. Additionally, approaches to quantifying the potential 'risk-of-bias', such as those in Boyd et al. (2022), can help assess dataset structure and inform the optimal selection of spatial and temporal units. In particular, unstructured and semi-structured data inherently involve uneven sampling across spatial units and years, leading to irregular spatial and temporal distributions of records. Datasets with high temporal resolution often lack consistent sampling effort over time. Older records, for instance, may be limited by digitization gaps or accuracy issues, reducing their usability and the number of available observations. Additionally, external factors such as funding fluctuations or targeted surveys for specific taxa or habitats can create temporary, systematic shifts in sampling effort. Beyond temporal inconsistencies, spatial distribution can also be affected by these variations, resulting in uneven sampling across both space and time – patterns that are highly dataset-dependent. While Frescalo is designed to correct for such biases, prior knowledge of dataset-specific sampling discrepancies is crucial for refining these corrections. For example, if a dataset exhibits strong temporal unevenness due to the presence of older records, extending the length of the initial time periods can help aggregate

more data and reduce estimation uncertainty. Defining time periods incorrectly can introduce significant uncertainties in estimating the time factor, ultimately biasing the overall temporal trend. Thoughtful selection of spatial and temporal units is therefore essential to ensure robust trend analysis.

Defining neighbourhoods

The purpose of defining neighbourhoods is to define the rank abundance curves (i.e. the rank-frequency curves) at the species pool level. Each site has its own neighbourhood of other sites for which the rank abundance curve is derived. As we emphasised in section 'Spatial correction', neighbourhoods are at the heart of the Frescalo method, so it is necessary to characterise them as well as possible. Neighbourhood selection involves two main steps, such as 1) selecting the closest sites to the target, and then 2) filtering these sites by those that are most ecologically similar to the target. To first select the closest sites to the target, geographical distance is a key variable to consider, as environmental variables are often structured by autocorrelation. This may allow us to capture a set of variables that we often don't have. Once we have selected a set of sites close to the target, we filter these sites to retain the top N most similar sites (e.g. Hill uses 200 → 100 as default, i.e. 50%). This second step is based on the values of covariates that are considered to influence the similarity of species assemblages. The typical covariates used to define neighbourhoods are, for example, edaphic similarity, temperature, geology, land cover, and elevation. However, the choice of covariates depends on the study. For example, geology is only of interest if there is a contrast between calcareous and acidic rocks throughout the study area. Therefore, a wide range of variables including plant similarity, and climatic or topographic variables, could also be used. We note that neighbourhoods are not defined according to covariates that are thought to only affect sampling intensity, such as accessibility. In Fig. 2, we present a simplified scenario in which the number of species observed (i.e. recording intensity) varies within the neighbourhood of a given site to show how the rank-frequency curves are affected. In reality, sites are weighted by their similarity to the focus based on the selection of covariates (i.e. in Fig. 2, the weights are binary: cells included in a neighbourhood have a weight of 1, while other cells have a weight of 0. In typical applications, however, the weights are permitted to be continuous between 0 and 1). Each location can have any possible weight value between 0 and 1, reflecting how much it contributes to the neighbourhood for each location. This feature is intended to optimise the relevance of the derived rank abundance curve for each site, so that it can later be used to calculate survey effort and adjust for undersampling. However, no sensitivity analysis provides us with information on the optimal neighbourhood size (i.e. number of sites to consider), nor the optimal variables to consider when defining a neighbourhood, which is highly dependent on the analysis (e.g. grain, richness, disturbance, and recording effort; see Eichenberg et al. 2021 or Hill 2012, Auffret and Svenning 2022 for different neighbourhood definitions). While Hill (2012) showed that the

results are robust to neighbourhood size and weighting exponents, sensitivity analysis checks (e.g. $K = 50, 100, 150$) are required for each specific data.

Choose parameter values that reflect a well-recorded neighbourhood

To correct the local frequency of each neighbourhood (Fig. 2E), the target value Φ corresponding to a well-recorded neighbourhood must be chosen a priori. This value is defined as 0.74 by default (Hill 2012); Φ must lie between $1/n$ and 1 (with n being the number of species in the neighbourhood), and Hill's default 0.74 sits near the upper end of typical plant-data values. However, the definition of the Φ value can strongly influence the results, especially when focusing on under-recorded taxa (e.g. bryophytes). In such cases, using the default value of Φ will not sufficiently correct the predicted number of species in each neighbourhood, resulting in an underestimation of the species standardised probabilities (Fig. 3B), which will then be reflected in the estimate of the time factor, which will also be underestimated, thus biasing the trends. Increasing the Φ value in such cases will correct the underestimation of the predicted number of species (Hill 2012, Auffret and Svenning 2022). However, if the focus is on species that are not too rare, small variations of Φ around the default value (i.e. 0.74) will not affect the trend estimates (sensitivity analysis in Hill 2012). Another possibility is to

take the same approach as Hill (2012), calculating the corresponding ϕ value for each species (Section 'Neighbourhoods weights and frequency-weighted mean frequencies') and then taking the 98.5th percentile to determine Φ . However, as noted by Hill (2012), very little is known about the behaviour of Φ and further work is required.

In addition, to estimate the proportion of common species (S_{ij} , Fig. 3), it is necessary to define the threshold for benchmark species (R). By default, this threshold is set at 0.27 (Hill 2012), meaning the top 27% of species, based on their normalised rank R'_{ij} , are considered benchmarks. Benchmark species are determined within each neighbourhood, rather than being fixed across all neighbourhoods. Sensitivity analyses (Hill 2012, Auffret and Svenning 2022) have shown that while R influences the absolute values of the time factor, it may often have relatively little impact on the overall temporal trend patterns. However, when working at finer spatial scales or when the most ubiquitous species are excluded, a lower R value may be necessary (Auffret and Svenning 2022). Decreasing R ensures that only the most locally abundant species are included as benchmarks (Fig. 2F), but at the cost of reducing the number of benchmark species available. While our goal is to select stable species, lowering R also decreases the precision of time factor estimation. On the other hand, increasing R and admitting more benchmarks can undermine the modelling assumptions if the added species display

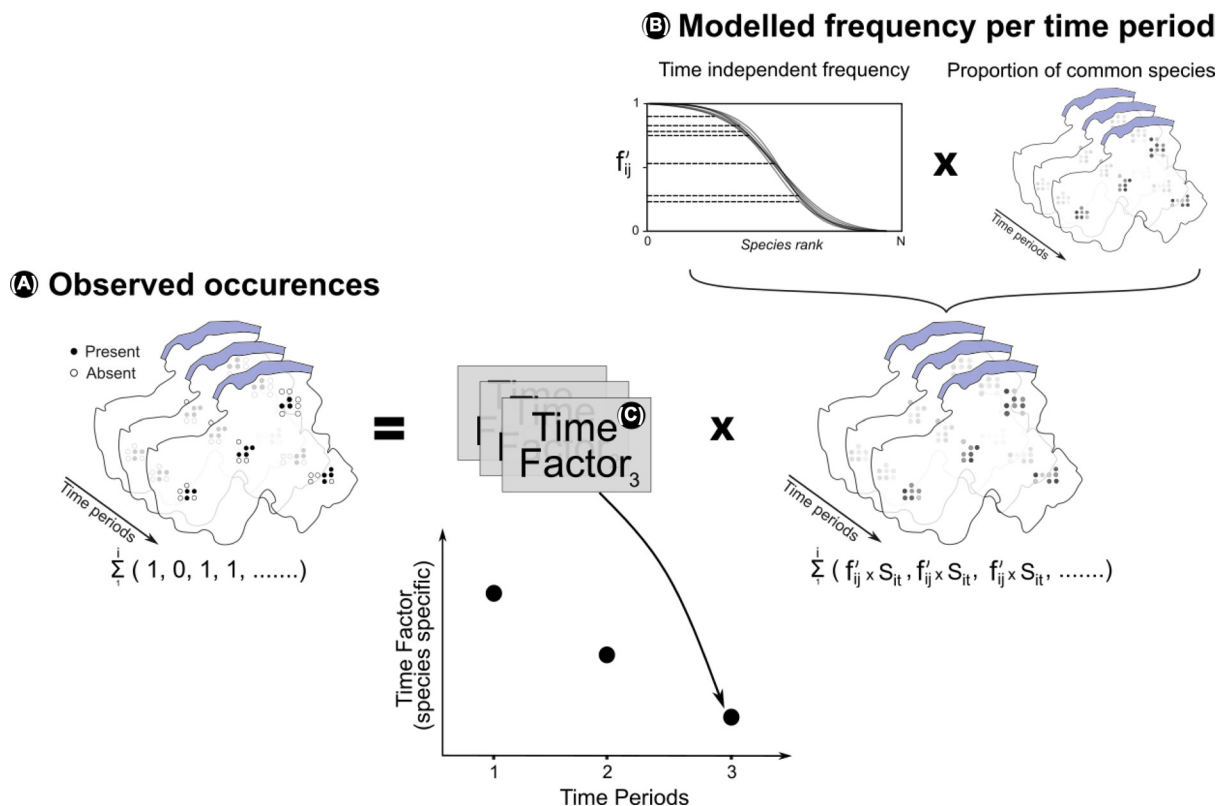


Figure 3. The temporal correction is an adjustment of the observed occurrences (A) by the modelled frequency (B) per time period. Observed occurrences (A) correspond to the sum of the presences per species in the raw dataset. The modelled frequency per species (B) is obtained from the corrected frequency (Fig. 2) and the proportion of benchmark species recorded (Fig. 2) per site/time period combination. The time factor (C) corresponds to the temporal values that adjust the sum of the observed occurrences to the sum of the modelled frequencies.

more systematic bias in how they have been recorded over time: this is a particular risk if analysts are assessing groups of species that are not typically recorded together in the field (cf. Coomber et al. 2021). Ultimately the analyst must be reasonably confident that an optimal balance has been found between ensuring benchmark stability and maintaining model plausibility.

Potential pitfalls

Here, we highlight the different key points that can cause estimation errors when using Frescalo. In most semi- or unstructured datasets, the data are gathered from different sources, which can lead to a number of additional biases.

Taxonomic variability

Taxonomy plays a crucial role when analysing long-term data. The taxonomic conception of field botanists and global referees can vary widely, which can cause several issues. For instance, whether a plant is identified as a subspecies may depend on the botanist's expertise, the reference flora consulted, or the observation period. This inconsistency can have significant consequences, as a subspecies may suddenly be divided into multiple species, directly affecting species trend analyses (cf. Jansen and Dengler 2010, Pescott et al. 2018). However, this bias concerns not only Frescalo, but also all models that estimate the presence of species. One way to address this challenge is to conduct analyses at the species or even higher taxonomic levels, such as genera. This approach should be applied with caution, particularly for genera in which species are phylogenetically very closely related yet differ markedly in their ecological attributes, but it is especially well suited to taxa with complex taxonomy (e.g. apomictic groups). Engaging in clear discussions with experts on this matter is essential. Another important concern involves rare or attractive species. These species are often surveyed more extensively. In contrast, common species may be under-recorded because recorders become accustomed to them. This can result in 'average' sampling intensity corrections being less accurate: targeted surveys for rare species mean that survey efforts for common species may not reflect efforts for rare ones. To tackle this issue, it can be helpful to analyse the correlation between spatial distribution (the number of grid cells occupied) and the total number of occurrences per species.

Sampling bias and data correction

Sampling effort often fluctuates over time when using semi-structured or unstructured data. While Frescalo is designed to correct for this bias, certain precautions can help optimize the accuracy of the algorithm. First, it is essential to check the phenology of surveys across years. In some years, sampling may have been more intensive and spread across multiple seasons, which can significantly impact trend estimates for species detectable only during specific seasons, such as vernal species. If some of the earlier data collection was focused on spring species and the more recent data collection was not,

then the apparent decline observed would be a mere artefact of changing sampling effort. To prevent this, it may be crucial to monitor the distribution of occurrences throughout the year and ensure it remains consistent between the time periods analysed. If certain seasons are disproportionately represented within some time periods, it may be necessary to exclude the affected species from the analysis or to combine the different seasons to avoid such bias. Another important issue arises from nested datasets. Semi-structured or unstructured datasets often aggregate smaller datasets from different regions. Some of these may focus exclusively on specific taxa, such as orchids, ferns, or macrophytes, which can skew the dataset (Stroh et al. 2023). This violates the assumption that species are sampled in proportion to their true frequencies. Sudden changes in species occurrence could merely reflect shifts in survey methodology. Since Frescalo pools data across years, it will struggle to handle such inconsistencies when they are strongly correlated with the chosen time periods in the analysis. Understanding the objective of each survey (i.e. each sub-dataset, if disaggregation is possible) is therefore essential. One solution is to remove imbalanced datasets or exclude taxa that are over-sampled during certain periods (e.g. Blockeel et al. 2014, chapter 5), or to stratify the analysis by survey type if metadata allow. However, many other biases associated with unstructured data must be taken into account. For example, strong spatial biases may arise when observers select sites that are easily accessible, such as those near roads (Tiago et al. 2017), or sites that contain rare species (Johnston et al. 2023). Large body size or colourfulness may also significantly impact sampling pressure (Callaghan et al. 2021, Stoudt et al. 2022). Finally, the lack of metadata from unstructured schemes (Bowler et al. 2022) can restrict the estimation of species trends a posteriori.

Handling absences

When a species is not recorded during a specific time period across all neighbourhoods, Frescalo's time factor solver pulls x_{jt} towards zero, predicting complete absence. This approach differs somewhat from other modelling frameworks, like occupancy-detection models, which can account for the possibility that a species was present but missed during surveys (conditional on information richness in the data, as always). The assumption of definite absence can potentially introduce bias in trend analyses. However, if the time period is well-chosen (spanning several years) and there are no occurrences in the entire surrounding large-scale area, it may be reasonable to assume that the species was truly absent. Nonetheless, conducting a sensitivity analysis is a useful way to assess how these time periods might affect trend predictions.

Opportunities to use Frescalo to answer ecological questions about biodiversity changes

Detecting long-term changes in species' distributions

The Frescalo method provides a correction that generates time series for each species, tracking changes in species frequency

over time relative to common species. These time series are assumed to reveal true ecological patterns, as they reflect relative changes in frequency once effort adjustments have been made. The time factor for each species can also be compared across time periods. A time factor of 1 indicates that a species' frequency matches that of the average benchmark species within a given time period. Time factors also enable comparisons between species: for instance, if species A has a time factor of 1.0 and species B has 0.5, species B's relative frequency is half that of the benchmarks (and half that of species A). These properties make the time factor useful for analysing temporal species trends and identifying species that are either increasing or decreasing in frequency. However, each time factor estimation is surrounded by uncertainty, which must be accounted for to accurately assess species abundance changes (Pescott et al. 2022). Some models, such as multilevel models (Gelman and Hill 2007), explicitly handle uncertainty around each observation, allowing for lower error estimation of temporal dynamics. This approach helps to assess both relative and absolute uncertainty, leading to better categorisation of species according to their temporal trend. Alternative approaches for identifying long-term changes using Frescalo have been proposed. These include bootstrapping and a posteriori classification (Pescott et al. 2022), as well as the estimation of species-specific occurrence changes over time (Eichenberg et al. 2021).

Describing the spatial patterns of distribution change

Beyond temporal analysis, the Frescalo method can also be applied to examine spatial patterns of species. Following the approach suggested by Bijlsma (2013) and applied by Eichenberg et al. (2021), Frescalo outputs can be used to map species occurrence probabilities at the site level across different time periods. These maps help identify areas where significant changes in species frequency have occurred. Additionally, they can serve as valuable tools for communicating species distributions to stakeholders, highlighting knowledge gaps, and informing conservation efforts.

Another interesting point is that the probabilistic species distributions generated by Frescalo are somewhat similar to those that could be obtained from species distribution models (SDM; Guisan and Thuiller 2005, Thuiller 2024) applied to the different time periods. In one case, the probabilistic estimates are derived given the benchmark species and neighbourhoods, while in the other, they are given by the environmental variables measured and extrapolated at the sites i for a given time period. It will be interesting to investigate whether they give comparable estimates. However, Frescalo corrects for temporal and spatial bias, which would also need to be integrated in the SDM applications (Chauvier et al. 2021). Something interesting could then be to produce species probability distributions from the SDMs, for the different time periods, correcting for simple spatial bias (distance to road, cities), and see if the SDMs still manage to capture the species niche and reproduce distribution shifts over time. If sampling bias does not have much effect on how the species occupies

its environmental niche geographically, then it should not be too problematic for the SDM. The challenge for such a comparison may be finding the appropriate and relevant environmental variables over time.

Untapped opportunities: investigating the spatio-temporal biases of biodiversity monitoring schemes

Frescalo can also be a valuable tool for enhancing biodiversity monitoring schemes. As previously mentioned, the method is particularly effective in correcting unstructured or semi-structured datasets, such as atlas-type datasets. However, these datasets often carry inherent historical biases – whether practical (e.g. challenging areas to sample), financial (e.g. insufficient funding for sampling), or historical (e.g. intensive sampling by private recorders in specific areas) – all of which contribute to variability in sampling effort. Consequently, stakeholders, such as field naturalists, need tools to identify areas that are under- or over-sampled in order to refine future sampling plans. While modifying ongoing sampling strategies may be challenging due to limited funding and time constraints, Frescalo can provide valuable insights. As discussed in section 'Standardising local species frequencies', the Poisson point process used in Frescalo fits each neighbourhood's rank-frequency curve to the 'best sampled' neighbourhood, with an α value reflecting the intensity of this process. A higher α indicates that a neighbourhood is farther from the 'best sampled' neighbourhood (defined by Φ), meaning it was under-sampled. Therefore, the α values for each neighbourhood throughout the study period serve as a proxy for sampling intensity, with high α values indicating under-sampled neighbourhoods and low α values indicating well-sampled ones. Once these α values are mapped, stakeholders can visually assess discrepancies in sampling intensity across sites and adjust their sampling strategies accordingly to ensure more uniform coverage across the area.

Conclusion

Hill's Frescalo method fills a critical gap by enabling robust estimation of species-frequency trends from unstructured occurrence data, correcting both spatial and temporal sampling biases through a two-step neighbourhood-based algorithm. Yet its uptake has arguably been limited by perceived mathematical complexity and a lack of accessible implementations. In this paper, we have:

1. **Demystified the algorithm**, clearly explaining 1) spatial standardisation via frequency scaling to a common benchmark Φ , and 2) temporal correction and trend estimation using benchmark-derived effort indices S_{it} and time factors x_{jt} , respectively.
2. **Provided a practical roadmap**, highlighting key decisions – site and period definitions, neighbourhood construction, parameter choices (Φ , R') – and drawing attention to common pitfalls around taxonomy, seasonality, nested datasets, and absences.

3. **Pointed to an open-source R package**, which streamlines and speeds up algorithm execution and visualization, lowering the barrier for ecologists to apply Frescalo to their own datasets.

Beyond trend estimation, Frescalo's outputs – spatial multipliers α_i , per-period effort indices s_{it} , and species time factors x_{jt} – also offer powerful diagnostics for:

- **survey design**, by mapping under- and over-sampled regions;
- **bias assessment**, by comparing Frescalo maps with SDM or occupancy-model outputs; and
- **spatial change detection**, by producing maps of local frequency shifts.

Looking forward, we see several promising extensions:

- **Bayesian or hierarchical occupancy hybrids**, e.g. to jointly estimate α_i , x_{jt} and their uncertainties;
- **integration with SDMs**, e.g. using α_i as bias-correction grids to test concordance with environmental-driven predictions; and
- **real-time monitoring**, by updating α_i and x_{jt} with 'live streaming' citizen-science data.

By clarifying the methodology, offering concrete guidance, and providing ready-to-use software, we aim to spark broader adoption of Frescalo. This will enable more accurate, transparent, and reproducible assessments of biodiversity change in an era of increasingly large but frequently unstructured biological datasets.

Glossary

Parameter	Definition	Biological meaning
S_i	Number of searches made in a site	-
λ_{ij}	Discovery rate of a species j in a site i through a Poisson-process model	-
f_{ij}	Observed frequency of species j in the (weighted) neighbourhood of site i	Observed frequency of species in the neighbourhood of site i
N_2	Correspond to the inverse of Simpson index (Hill 1973)	Relative abundances of common species (Roswell et al. 2021)
ϕ_i	Weighted mean of the observed frequencies of species j in neighbourhood i (i.e. f_{ij}), for all time, for a standardized level of sampling (i.e. N_2)	Expected species frequency within a neighbourhood, also mentioned as a measure of the sampling intensity. In other words, this is the ratio of the mean species richness to the 'effective number of common species' (i.e. N_2)
Φ	Correspond to the ϕ_i for the well-sampled neighbourhood	Value is used as the target ϕ (i.e. Φ) for all neighbourhoods if no ϕ value is set as default
f'_{ji}	Expected frequency of species j in neighbourhood i , for all time, for a standardized level of sampling	Frequency of the species in the neighbourhood of site i after correction (simulating a thorough search). A proxy for the 'true' discoverability- or effort-standardised neighbourhood species rank-frequency curve (Pescott 2026)
α_i	Scaling factor applied to the intensity of the discovery-related Poisson process (Fig. 2); iteratively estimated to move ϕ_i towards Φ . This is the conversion of all species' f_{ij} to f'_{ji} (so by definition ϕ_i to Φ) for a site's neighbourhood	Proxy of a sampling-intensity multiplier, which means that if the α is high, the neighbourhood correction is strong and species weighted frequencies are increased
R^*	Correspond to the proportion of the species defined as benchmarks (i.e. species above a defined rank) specific to a given neighbourhood	Corresponds to the most common species within a neighbourhood
S_{it}	Proportion of locally frequent species (R^*), used to index sampling effort in particular sites and time periods	Proxy of local sampling effort
P_{ijt}	True underlying probability that a species j is recorded in a site i at a given time t	Empirically, represented by the observed occurrence data (i.e. presence/absence of a species j in sites i at time t), which are assumed to incorporate both effort-related and true ecological signals
$\Sigma_i P_{ijt}$	Sum of the observed occurrences which is equivalent to the sum of true underlying occurrence probabilities across all sites i at a given time period t	Sum of the observed occurrence data across all sites i for a species j at a given time period t
Q_{ijt}	Modelled estimate of a species frequency after adjusting neighbourhood level sampling intensity (i.e. f'_{ji}) and site level sampling effort (i.e. S_{it})	Corrected modelled species frequency by the sampling effort for a given time period
$\Sigma_i Q_{ijt}$	Sum of the modelled estimate of a species frequency after adjusting neighbourhood level sampling intensity and site level sampling effort across all sites i at a given time period t	Sum of the corrected modelled frequency across all sites i for a species j at a given time period t
x_{jt} = time factor	Corresponds to the multiplier which requires adjusting the $\Sigma_i Q_{ijt}$ to equal the $\Sigma_i P_{ijt}$. This is assumed to represent a true ecological pattern	Corresponds to the temporal deviation at a given time of a species compared to the most common species and its modelled frequency (i.e. Q_{ijt}). A value of 1 means that the species is as much recorded as the average of the benchmark species

Funding – This research is a product of the IMPACTS group funded by the Centre for the Synthesis and Analysis of Biodiversity (CESAB) of the Foundation for Research on Biodiversity (FRB) and the Ministry of Ecological Transition. This work was also funded through the European Union's Horizon Europe under grant agreement no. 101134954 (Obsession) and by Biodiversa+, the European Biodiversity Partnership under the 2021–2022 BiodivProtect joint call for research proposals, co-funded by the European Commission (no. 101052342) and with the funding organizations ANR. This research was partially supported by NERC, through the UKCEH National Capability for UK Challenges Programme NE/Y006208/1.

Conflict of interest – The authors declare no conflict of interest.

Author contributions

Romain Goury: Conceptualization (equal); Methodology (equal); Visualization (equal); Writing – original draft (lead). **Diana E. Bowler:** Conceptualization (equal); Supervision (equal); Writing – review and editing (equal). **Colin Harrower:** Software (equal); Writing – review and editing (equal). **Tamara Münkemüller:** Funding acquisition (equal); Supervision (equal); Writing – review and editing (equal). **Jeanne Vallet:** Methodology (equal); Writing – review and editing (equal). **Jon M. Yearsley:** Software (equal); Writing – review and editing (equal). **Wilfried Thuiller:** Funding acquisition (equal); Supervision (equal); Writing – review and editing (equal). **Oliver L. Pescott:** Conceptualization (equal); Software (equal); Supervision (equal); Writing – review and editing (equal).

Transparent peer review

The peer review history for this article is available at <https://www.webofscience.com/api/gateway/wos/peer-review/eco.g.08270>.

Data availability statement

Data sharing is not applicable to this article as no new data were created or analyzed in this study.

Supporting information

The Supporting information associated with this article is available with the online version.

References

- Auffret, A. G. and Svenning, J.-C. 2022. Climate warming has compounded plant responses to habitat conversion in northern Europe. – *Nat. Commun.* 13: Article 1.
- Bijlsma, R. J. 2013. The estimation of species richness of Dutch bryophytes between 1900 and 2011. Documentation of VBA-procedures based on the Frescalo program. – <https://research.wur.nl/en/publications/the-estimation-of-species-richness-of-dutch-bryophytes-between-19>.
- Binley, A. D. and Bennett, J. R. 2023. The data double standard. – *Methods Ecol. Evol.* 14: 1389–1397.
- Blockeel, B. and Hill, M. O. 2014. Atlas of British and Irish bryophytes. – British Bryological Society, www.britishbryologicalsociety.org.uk/publications/atlas-of-british-and-irish-bryophytes.
- Boakes, E. H., McGowan, P. J. K., Fuller, R. A., Chang-qing, D., Clark, N. E., O'Connor, K. and Mace, G. M. 2010. Distorted views of biodiversity: spatial and temporal bias in species occurrence data. – *PLoS Biol.* 8: e1000385.
- Bowler, D. E., Bhandari, N., Repke, L., Beuthner, C., Callaghan, C. T., Eichenberg, D., Henle, K., Klenke, R., Richter, A., Jansen, F., Bruehlheide, H. and Bonn, A. 2022. Decision-making of citizen scientists when recording species observations. – *Sci. Rep.* 12: 11069.
- Boyd, R. J., Powney, G. D., Burns, F., Danet, A., Duchenne, F., Grainger, M. J., Jarvis, S. G., Martin, G., Nilsen, E. B., Porcher, E., Stewart, G. B., Wilson, O. J. and Pescott, O. L. 2022. ROBITT: a tool for assessing the risk-of-bias in studies of temporal trends in ecology. – *Methods Ecol. Evol.* 13: 1497–1507.
- Boyd, R. J., Powney, G. D. and Pescott, O. L. 2023. We need to talk about nonprobability samples. – *Trends Ecol. Evol.* 38: 521–531.
- Boyd, R. J., Bowler, D. E., Isaac, N. J. B. and Pescott, O. L. 2024a. On the trade-off between accuracy and spatial resolution when estimating species occupancy from geographically biased samples. – *Ecol. Modell.* 493: 110739.
- Boyd, R. J., Stewart, G. B. and Pescott, O. L. 2024b. Descriptive inference using large, unrepresentative nonprobability samples: an introduction for ecologists. – *Ecology* 105: e4214.
- Callaghan, C. T., Poore, A. G. B., Hofmann, M., Roberts, C. J. and Pereira, H. M. 2021. Large-bodied birds are over-represented in unstructured citizen science data. – *Sci. Rep.* 11: 19073.
- Chauvier, Y., Thuiller, W., Brun, P., Laverne, S., Descombes, P., Karger, D. N., Renaud, J. and Zimmermann, N. E. 2021. Influence of climate, soil, and land cover on plant species distribution in the European Alps. – *Ecol. Monogr.* 91: e01433.
- Coomer, F. G., Smith, B. R., August, T. A., Harrower, C. A., Powney, G. D. and Mathews, F. 2021. Using biological records to infer long-term occupancy trends of mammals in the UK. – *Biol. Conserv.* 264: 109362.
- Dobson, A. D. M. et al. 2020. Making messy data work for conservation. – *One Earth* 2: 455–465.
- Dornelas, M., Chase, J. M., Gotelli, N. J., Magurran, A. E., McGill, B. J., Antão, L. H., Blowes, S. A., Daskalova, G. N., Leung, B., Martins, I. S., Moyes, F., Myers-Smith, I. H., Thomas, C. D. and Vellend, M. 2023. Looking back on biodiversity change: lessons for the road ahead. – *Philos. Trans. R. Soc. B* 378: 20220199.
- Dyer, R. J., Gillings, S., Pywell, R. F., Fox, R., Roy, D. B. and Oliver, T. H. 2017. Developing a biodiversity-based indicator for large-scale environmental assessment: a case study of proposed shale gas extraction sites in Britain. – *J. Appl. Ecol.* 54: 872–882.
- Eichenberg, D., Bowler, D. E., Bonn, A., Bruehlheide, H., Grescho, V., Harter, D., Jandt, U., May, R., Winter, M. and Jansen, F. 2021. Widespread decline in central European plant diversity across six decades. – *Global Change Biol.* 27: 1097–1110.
- Fox, R., Oliver, T. H., Harrower, C., Parsons, M. S., Thomas, C. D. and Roy, D. B. 2014. Long-term changes to the frequency of occurrence of British moths are consistent with opposing and synergistic effects of climate and land-use changes. – *J. Appl. Ecol.* 51: 949–957.
- Frank, S. A. 2009. The common patterns of nature. – *J. Evol. Biol.* 22: 1563–1585.
- Geldmann, J., Heilmann-Clausen, J., Holm, T. E., Levinsky, I., Markussen, B., Olsen, K., Rahbek, C. and Tøttrup, A. P. 2016. What determines spatial bias in citizen science? Exploring four

- recording schemes with different proficiency requirements. – *Divers. Distrib.* 22: 1139–1149.
- Gelman, A. and Hill, J. 2007. Data analysis using regression and multilevel/hierarchical models. – Cambridge Univ. Press.
- Gonzalez, A., Chase, J. M. and O'Connor, M. I. 2023. A framework for the detection and attribution of biodiversity change. – *Philos. Trans. R. Soc. B* 378: 20220182.
- Grace, J. B. 2024. An integrative paradigm for building causal knowledge. – *Ecol. Monogr.* 2024: e1628.
- Guisan, A. and Thuiller, W. 2005. Predicting species distribution: offering more than simple habitat models. – *Ecol. Lett.* 8: 993–1009.
- Hill, M. O. 1973. Diversity and evenness: a unifying notation and its consequences. – *Ecology* 54: 427–432.
- Hill, M. O. 2012. Local frequency as a key to interpreting species occurrence data when recording effort is not known. – *Methods Ecol. Evol.* 3: 195–205.
- Isaac, N. J. B., van Strien, A. J., August, T. A., de Zeeuw, M. P. and Roy, D. B. 2014. Statistics for citizen science: extracting signals of change from noisy ecological data. – *Methods Ecol. Evol.* 5: 1052–1060.
- Jansen, F. and Dengler, J. 2010. Plant names in vegetation databases – a neglected source of bias. – *J. Veg. Sci.* 21: 1179–1186.
- Johnston, A., Matechou, E. and Dennis, E. B. 2023. Outstanding challenges and future directions for biodiversity monitoring using citizen science data. – *Methods Ecol. Evol.* 14: 103–116.
- Jost, L. 2006. Entropy and diversity. – *Oikos* 113: 363–375.
- Kelling, S., Johnston, A., Bonn, A., Fink, D., Ruiz-Gutierrez, V., Bonney, R., Fernandez, M., Hochachka, W. M., Julliard, R., Kraemer, R. and Guralnick, R. 2019. Using semistructured surveys to improve citizen science data for monitoring biodiversity. – *BioScience* 69: 170–179.
- Latour, J. and van Swaay, C. 1992. Dagvinders als indicatoren voor de regionale milieukwaliteit. – *Levende Nat.* 93: 19–22.
- MacKenzie, D. I., Nichols, J. D., Lachman, G. B., Droege, S., Andrew Royle, J. and Langtimm, C. A. 2002. Estimating site occupancy rates when detection probabilities are less than one. – *Ecology* 83: 2248–2255.
- Maes, D. and van Swaay, C. A. M. 1997. A new methodology for compiling national Red Lists applied to butterflies (Lepidoptera, Rhopalocera) in Flanders (N-Belgium) and the Netherlands. – *J. Insect Conserv.* 1: 113–124.
- Montràs-Janer, T., Suggitt, A. J., Fox, R., Jönsson, M., Martay, B., Roy, D. B., Walker, K. J. and Auffret, A. G. 2024. Anthropogenic climate and land-use change drive short- and long-term biodiversity shifts across taxa. – *Nat. Ecol. Evol.* 8: 739–751.
- Pescott, O. L. 2025. An R translation of Hill's Fortran code for Frescalo [ver. 1.0.0]. – <https://doi.org/10.5281/zenodo.15305437>, <https://doi.org/10.1111/bij.12581>.
- Pescott, O. L. 2026. Unifying occupancy-detection and local frequency scaling (Frescalo) models. – *Ecol. Modell.* 511: 111367.
- Pescott, O. L., Humphrey, T. and Walker, K. 2018. A short guide to using British and Irish plant occurrence data for research. – <https://doi.org/10.13140/RG.2.2.33746.86720>.
- Pescott, O. L., Walker, K. J., Pocock, M. J. O., Jitlal, M., Outhwaite, C. L., Cheffings, C. M., Harris, F. and Roy, D. B. 2015. Ecological monitoring with citizen science: the design and implementation of schemes for recording plants in Britain and Ireland. – *Biol. J. Linn. Soc.* 115: 505–521.
- Pescott, O. L., Humphrey, T. A., Stroh, P. A. and Walker, K. J. 2019. Temporal changes in distributions and the species atlas: how can British and Irish plant data shoulder the inferential burden? – *Br. Ir. Bot.* 1: 250–282.
- Pescott, O. L., Stroh, P. A., Humphrey, T. A. and Walker, K. J. 2022. Simple methods for improving the communication of uncertainty in species' temporal trends. – *Ecol. Indic.* 141: 109117.
- Prendergast, J. R., Wood, S. N., Lawton, J. H. and Eversham, B. C. 1993. Correcting for variation in recording effort in analyses of diversity hotspots. – *Biodivers. Lett.* 1: 39–53.
- Redhead, J. W., Woodcock, B. A., Pocock, M. J. O., Pywell, R. F., Vanbergen, A. J. and Oliver, T. H. 2018. Potential landscape-scale pollinator networks across Great Britain: structure, stability and influence of agricultural land cover. – *Ecol. Lett.* 21: 1821–1832.
- Rich, T. C. G. 2006. Floristic changes in vascular plants in the British Isles: geographical and temporal variation in botanical activity 1836–1988. – *Bot. J. Linn. Soc.* 152: 303–330.
- Roswell, M., Dushoff, J. and Winfree, R. 2021. A conceptual guide to measuring species diversity. – *Oikos* 130: 321–338.
- Stoudt, S., Goldstein, B. R. and de Valpine, P. 2022. Identifying engaging bird species and traits with community science observations. – *Proc. Natl Acad. Sci. USA* 119: e2110156119.
- Stroh, P. A., Walker, K. J., Humphrey, T. A., Pescott, O. L. and Burkmar, R. J. 2023. Plant atlas 2020: mapping changes in the distribution of the British and Irish flora. – Princeton Univ. Press.
- Suggitt, A. J., Wheatley, C. J., Aucott, P., Beale, C. M., Fox, R., Hill, J. K., Isaac, N. J. B., Martay, B., Southall, H., Thomas, C. D., Walker, K. J. and Auffret, A. G. 2023. Linking climate warming and land conversion to species' range changes across Great Britain. – *Nat. Commun.* 14: 6759.
- Telfer, M. G., Preston, C. D. and Rothery, P. 2002. A general method for measuring relative change in range size from biological atlas data. – *Biol. Conserv.* 107: 99–109.
- Thuiller, W. 2024. Ecological niche modelling. – *Curr. Biol.* 34: R225–R229.
- Tiago, P., Ceia-Hasse, A., Marques, T. A., Capinha, C. and Pereira, H. M. 2017. Spatial distribution of citizen science casuistic observations for different taxonomic groups. – *Sci. Rep.* 7: 12832.
- Troutet, J., Grandcolas, P., Blin, A., Vignes-Lebbe, R. and Legendre, F. 2017. Taxonomic bias in biodiversity data and societal preferences. – *Sci. Rep.* 7: 9132.
- van Strien, A. J., van Swaay, C. A. M. and Termaat, T. 2013. Opportunistic citizen science data of animal species produce reliable estimates of distribution trends if analysed with occupancy models. – *J. Appl. Ecol.* 50: 1450–1458.
- White, H. J., Gaul, W., Sadykova, D., León-Sánchez, L., Caplat, P., Emmerson, M. C. and Yearsley, J. M. 2019. Land cover drives large scale productivity–diversity relationships in Irish vascular plants. – *PeerJ* 7: e7035.