



Interpretable contour level selection for heat maps for gridded data

Tarn Duong¹

Received: 15 September 2025 / Accepted: 8 June 2026

© The Author(s), under exclusive licence to Springer-Verlag GmbH Germany, part of Springer Nature 2026

Abstract

Gridded data formats, where the observed multivariate data are aggregated into grid cells, ensure confidentiality and reduce storage requirements, with the trade-off that access to the underlying point data is lost. Heat maps are a highly pertinent visualisation for gridded data, and heat maps with a small number of well-selected contour levels offer improved interpretability over continuous contour levels. There are many possible contour level choices. Amongst them, density contour levels are highly suitable in many cases. Current methods for computing density contour levels requires access to the observed point data, so they are not applicable to gridded data. To remedy this, we introduce an approximation of density contour levels for gridded data. We then compare our proposed method to existing contour level selection methods, and conclude that our proposal provides improved interpretability for synthetic and experimental gridded data.

Keywords Big data · Highest density region · Jenks contour levels

1 Introduction

Gridded data are an effective format for the storage and dissemination of confidential and high-volume data. Confidential data are aggregated within a grid cell to ensure their privacy, whereas high volume data (aka Big data) are aggregated into grid cells to reduce their storage requirements. A heat map is a common visualisation method for bivariate gridded data where the value of the variable of interest is encoded by a colour scale. A continuous colour scale provides the most complete visualisation of the data (Tobler 1973), though the large number of colours may impede interpretability (Jenks 1963; Dobson 1973). The usual approach to increasing interpretability

✉ Tarn Duong
tarn.duong@gmail.com

¹ Independent Researcher, Paris 75000, France

is to reduce the number of distinct colours by assigning the same colour to an interval of values. The key question is to select an optimal set of contour levels and a colour scale to describe meaningfully the structure of the data surface (Klemelä 2009). For simplicity, we treat the colour scale selection here as a secondary question, since we follow the recommendations in Zeileis et al. (2009, 2020), and focus our attention on the contour level selection.

There are many possibilities for the contour level selection: for example, the ‘natural’ contour levels (Jenks 1967) are widely employed. Despite their widespread acceptance, the spatial relevance of their associated contour regions has not been established rigorously within a statistical framework. In contrast, density contour levels rigorously define spatially relevant regions (Bowman and Foster 1993; Hyndman 1996). A τ -density contour level is associated with the smallest region that contains τ probability mass of the data sample. Whilst the probabilistic interpretation of the density contour regions is crucial for interpreting density estimates, we claim that density contour levels can be extended so that they are also useful for the visualisation of gridded data which are not density estimates. Though we continue to call them ‘density contour levels’ as this is already well-established. The most popular method for calculating these density contour levels requires access to original point data which does not pertain to our case (Hyndman 1996). Thus we propose an approximation of density contour levels which is suitable for gridded data.

Our first motivating data set is a $1 \text{ km} \times 1 \text{ km}$ gridded data set of socioeconomic indicators from the 2019 household census of about 30 million households in France (Insee 2019). The socioeconomic indicators involve income, poverty and housing etc. Due to their sensitive nature, this grid resolution is a trade-off between the dissemination of detailed socio-economic indicators for informed policy decision making, and the guarantee of privacy for individual households. We focus on the population density in the ROI (region of interest) $[1.97^\circ\text{E}, 2.79^\circ\text{E}] \times [48.63^\circ\text{N}, 49.06^\circ\text{N}]$ which consists of 2078 grid cells and includes the Paris city centre and its inner ring of suburbs, as displayed in Fig. 1, in the EPSG 3035 projected coordinate system. The continuous colour scale, ranging from 40,000 persons per km^2 (red) to 20,000 persons per km^2 (orange) to 10,000 persons per km^2 (yellow), provides a visually appealing heat map of the population distribution. We are able to identify qualitatively the densely populated regions (hotspots), though it is challenging to compute quantitatively these hotspots.

Our second motivating data set is non-geospatial gridded data. It consists of the year, the month and the mean monthly GISTEMP temperature anomaly, for all years from 1880 to 2024. It is a reconstructed data set, drawing from many high-volume point data sources, and the aggregation into grid cells is carried out to ensure comparable statistical reliability for all grid cells, so that it can serve as a global reference data source (Vose et al. 2021; Lenssen et al. 2024; Gistemp Team 2025). We focus on the temperature anomalies in the ROI $[111.28^\circ\text{E}, 129.43^\circ\text{E}] \times [35.79^\circ\text{S}, 12.68^\circ\text{S}]$ which covers the state of Western Australia. It comprises 1740 grid cells for each year–month combination from 1880-01 to 2024-12. The heat map plots the year on the horizontal axis and the month on the vertical axis. The long term trends in Fig. 2 is from cool anomalies (blue) in the first half of the year in the 1880–1900s to hot

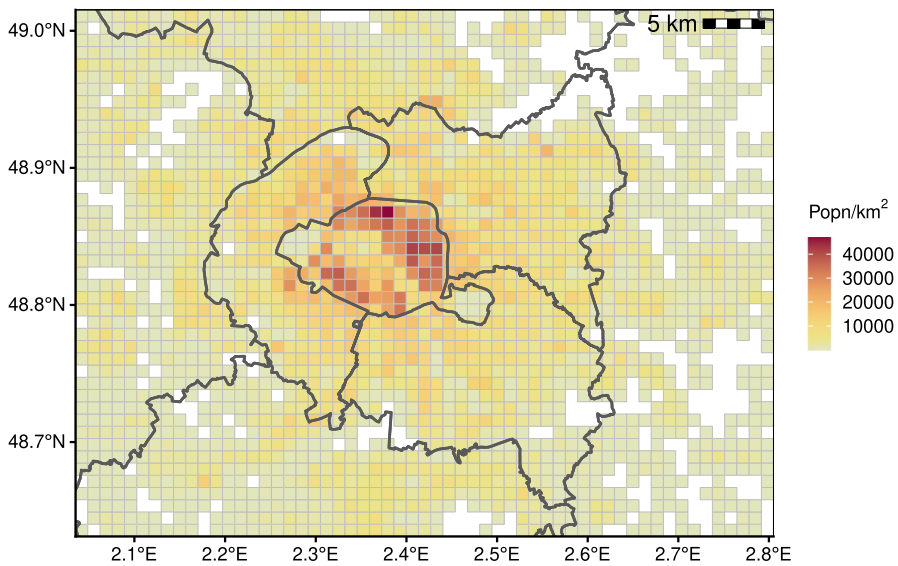


Fig. 1 Heat map for geospatial gridded data. Population of the Paris city region on $1 \text{ km} \times 1 \text{ km}$ grid, with continuous sequential ‘Heat’ colour scale

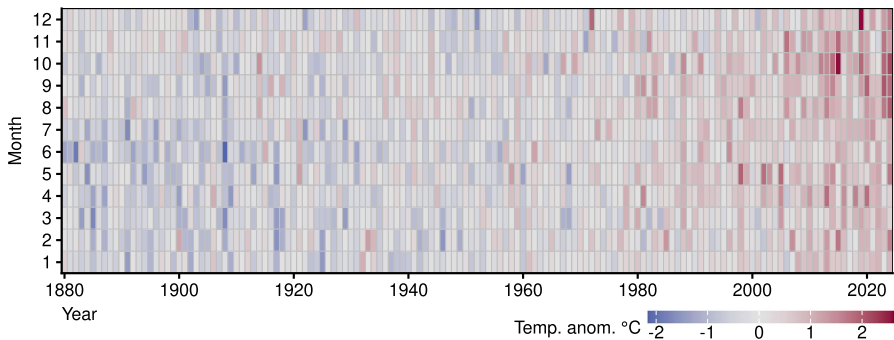


Fig. 2 Heat map for non-geospatial gridded data. Year-monthly surface temperature anomaly time series in Western Australia, with continuous diverging ‘Red-Blue’ colour scale

anomalies (red) in the second half of the year since the 2000s, though the visual impression of these trends is not overly strong due to the continuous colour scale.

This manuscript is organised as follows. In Sect. 2, we define rigorously density contour levels, and introduce our proposed approximation of density contour levels for gridded data. In Sect. 3, we verify the efficacy of this approximation for synthetic and experimental gridded data, as well as the increased interpretability of density contour levels over existing contour level estimation methods. We end with some concluding remarks.

2 Density contours

We begin by examining the special case for a probability density function. A τ -density contour R_τ is defined as the smallest region that contains the upper τ probability of the density function f

$$R_\tau = \{\mathbf{x} : f(\mathbf{x}) \geq f_\tau\} \text{ and } \mathbb{P}(\mathbf{X} \in R_\tau) = \tau \quad (1)$$

where the random variable $\mathbf{X}$ is drawn from f (Bowman and Foster 1993; Hyndman 1996). This definition of R_τ in Eq. (1) involves the implicitly defined τ -density contour level f_τ . Hyndman (1996) asserts an alternative explicit definition of f_τ as the $(1 - \tau)$ -quantile of $Y = f(\mathbf{X})$, i.e.

$$f_\tau = F_Y^{-1}(1 - \tau) \quad (2)$$

where F_Y is the cumulative distribution function of Y . This quantile is well-defined as Y is a univariate random variable. Then, $\mathbb{P}(\mathbf{X} \in R_\tau) = \int_{\mathbb{R}^d} f(\mathbf{x}) \mathbf{1}\{f(\mathbf{x}) \geq f_\tau\} d\mathbf{x} = \mathbb{E}[\mathbf{1}\{f(\mathbf{X}) \geq f_\tau\}] = \mathbb{P}(f(\mathbf{X}) \geq f_\tau) = \mathbb{P}(Y \geq f_\tau) = 1 - F_Y(f_\tau) = 1 - F_Y(F_Y^{-1}(1 - \tau)) = \tau$, which verifies the probability associated with R_τ , i.e. the f_τ in Eqs. (1) and (2) are the same. The proof that R_τ has the minimal hypervolume amongst all sets of probability mass τ is more involved and the interested reader is urged to consult Hyndman (1996). We call R_τ a ‘density contour region’, following Polonik (1995), Chacón and Duong (2018), whilst Hyndman (1996) prefers ‘highest density region’.

A set of m density contour levels $0 < \tau_1 < \dots < \tau_m < 1$, and their corresponding m density contour regions $R_{\tau_1}, \dots, R_{\tau_m}$ facilitate an intuitive interpretation of density visualisations. Common choices include the quartile contours $\tau = 0.25, 0.5, 0.75$ or the odd decile contours $\tau = 0.1, 0.3, 0.5, 0.7, 0.9$, or the decile contours $\tau = 0.1, \dots, 0.9$, depending on the intricacy of structure of the density f . The upper limit of a suitable number of contour levels is most likely about eight since this approaches the limit of human abilities for visual differentiation (Jenks 1963; Delicado and Vieu 2015).

In Fig. 3, we illustrate the additional interpretability of density contour levels over a continuous colour scale for a multimodal Gaussian mixture density $3/11N((-1, 1), 1/8[1, 0; 0, 1]) + 4/11N((0, 0), 1/8[1, 9/10; 9/10, 1]) + 3/11N((1, -1), 1/8[1, 0; 0, 1])$. In Fig. 3a is the heat map of the density function with a continuous colour scale, ranging from red (high density) to orange (mid) to yellow (low). We are able to ascertain some qualitative characteristics, such as the multimodality and the relative heights/locations of the modes, they are not evident. In Fig. 3b, the solid black curves are the decile contour regions $f_{0.9} = 0.066, f_{0.7} = 0.18, f_{0.5} = 0.29, f_{0.3} = 0.38, f_{0.1} = 0.51$. With this set of contour levels, the structure of the density function is more clearly ascertained.

The density contour level f_τ in Eq. (2) remains difficult to compute since it requires a closed form of both the density f and the cumulative distribution of $f(\mathbf{X})$. The mostly commonly deployed method is a quantile-based approximation from a random sample $\mathbf{X}_1, \dots, \mathbf{X}_n$ drawn from the common density f (Hyndman 1996).

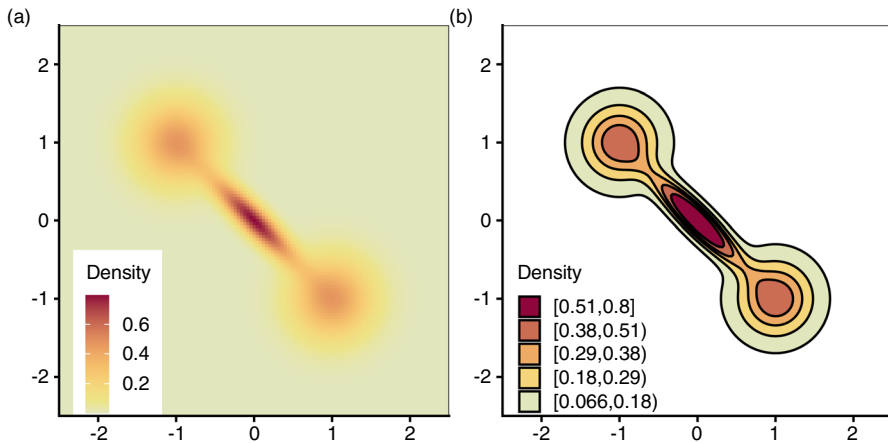


Fig. 3 Gaussian mixture density $4/11N((-1, 1), 1/8[1, 0; 0, 1]) + 3/11N((0, 0), 1/8[1, 9/10; 9/10, 1]) + 4/11N((1, -1), 1/8[1, 0; 0, 1])$. **(a)** Heat map with continuous 'Heat' colour scale, red (high) to orange (mid) to yellow (low). **(b)** Heat map with decile contour levels

From this sample data, we construct an estimator $\hat{f}$ of the target density f . Within a parametric framework, maximum likelihood estimators are popular, and within a non-parametric framework, histograms and kernel density estimators are widely used. The estimation methodology does not need to be specified here as the following apply for any consistent density estimator. The τ -density contour level f_τ is consistently estimated by plug-in estimators of Eqs. (1–2), i.e., $\tilde{f}_\tau$ is the sample $(1 - \tau)$ -quantile of $\hat{f}(\mathbf{X}_1), \dots, \hat{f}(\mathbf{X}_n)$. This approach is outlined in Algorithm 1. The inputs are the random sample of size n and the probability τ . The output is the approximate τ -density contour level $\tilde{f}_\tau$. In Step 1, we draw a random data sample $\mathbf{X}_1, \dots, \mathbf{X}_n$ of size n from the common density f . In Step 2, we compute the density estimate $\hat{f}$ from the data sample. In Step 3, we compute the transformed data sample by evaluating the density estimate at the data sample values. In Step 4, the density contour level is approximated as the $(1 - \tau)$ -quantile of these transformed data values.

Input: random sample $\mathbf{X}_1, \dots, \mathbf{X}_n$, n sample size, τ probability

Output: $\tilde{f}_\tau$ approximate density contour level

- 1: Compute density estimate $\hat{f}$ from $\mathbf{X}_1, \dots, \mathbf{X}_n$
 - 2: Evaluate density estimate at data sample values $\hat{f}(\mathbf{X}_1), \dots, \hat{f}(\mathbf{X}_n)$
 - 3: Approximate density contour level $\tilde{f}_\tau := (1 - \tau)$ -quantile of $\hat{f}(\mathbf{X}_1), \dots, \hat{f}(\mathbf{X}_n)$
-

Algorithm 1 Quantile-based approximation of density contour level

The calculations in Algorithm 1 are intuitive, efficient, and valid for any data dimension d and any consistent density estimator $\hat{f}$, which has contributed to its wide acceptance. However, we cannot employ this algorithm for gridded data since it requires direct access to the original point data sample. So we search for alternatives to Steps 2–3 which avoid explicit computations with the point data $\mathbf{X}_1, \dots, \mathbf{X}_n$.

For our proposed alternative in Algorithm 2, the inputs are the density estimate $\hat{f}$, a set of M grid points $G = \{\mathbf{g}_1, \dots, \mathbf{g}_M\}$, and the probability τ . In Step 1, we

compute the hypervolume δ of a grid cell from G . In Steps 2–4, we compute the M probability elements for the grid G : the j th probability element p_j is obtained by multiplying the hypervolume δ and the density estimate value $\hat{f}(\mathbf{g}_j)$ at the grid point $\mathbf{g}_j$. In Steps 5–6, we compute the order statistics of the M probability elements, and their partial sums. In Step 7, we search for the smallest cell index j^* such that the partial sum is greater than or equal to τ , i.e. $P_{j^*} = p_{(j^*)} + \dots + p_{(M)}$ exceeds τ by the least amount. By construction, the union $\hat{R}_\tau = \mathbf{g}_{(j^*)} \cup \dots \cup \mathbf{g}_{(M)}$ is the smallest region, comprised of the grid cells, with probability mass $\geq \tau$, i.e. $\hat{R}_\tau$ is an approximation of R_τ . In Step 8, the density contour level is approximated by $\hat{f}(\mathbf{g}_{j^*})$, which corresponds to $\hat{f}_\tau$ from Algorithm 1.

Input: $\hat{f}$ density estimate, G estimation grid, τ probability

Output: $\hat{f}_\tau$ approximate density contour level

- 1: Compute hypervolume of grid cell δ
 - 2: **for** $j := 1$ to M **do**
 - 3: Evaluate density estimate at grid point $\hat{f}(\mathbf{g}_j)$
 - 4: Compute probability element $p_j := \delta \hat{f}(\mathbf{g}_j)$
 - 5: **end for**
 - 6: Compute order statistics $p_{(1)} \leq \dots \leq p_{(M)}$ of probability elements
 - 7: Compute partial sums $P_j := \sum_{\ell=j}^M p_{(\ell)}$ of order statistics
 - 8: Find smallest index $j^* \in \{1, \dots, M\}$ such that $P_{j^*} \geq \tau$
 - 9: Approximate density contour level $\hat{f}_\tau := \hat{f}(\mathbf{g}_{j^*})$
-

Algorithm 2 Grid-based approximation of density contour level

For a non-uniform grid, the calculation of the grid cell hypervolume in Step 1 is moved inside the for loop so it is calculated for each grid cell. Due to the invariance of density contour levels, Algorithm 2 can be adapted to any non-negative function g with a finite integral over the grid G . With $\hat{g}$, a grid-based estimate of g , we calculate S which is the sum of all probability elements of $\hat{g}$ on the grid G . If we normalise the grid-based values as $\hat{f} = \hat{g}/S$, then $\hat{f}$ is a suitable density estimate input in Algorithm 2. Since the order statistics of $p_1/S, \dots, p_M/S$ remain the same as those for $p_1, \dots, p_M$, then the approximate contour level for g is $\hat{g}_\tau = S\hat{f}_\tau$. This is the situation for the Paris population data. These contour levels are a sequence of numerical values, so a sequential colour scale with the monotonically increasing luminance is the most suitable for their visualisation, since this colour scale ‘is easily decoded by humans and matched to high-low differences in the underlying data’ (Zeileis et al. 2020).

Furthermore, for a function g whose range includes negative and positive values, then we partition $g = g^+ - g^-$, where $g^-(\mathbf{x}) = -\mathbf{1}\{g(\mathbf{x}) \leq 0\}g(\mathbf{x})$ and $g^+(\mathbf{x}) = \mathbf{1}\{g(\mathbf{x}) > 0\}g(\mathbf{x})$, where $\mathbf{1}\{\cdot\}$ is the indicator function. Since g^-, g^+ are non-negative functions, then we can apply Algorithm 2 to each of them, resulting in contour levels $\hat{g}_\tau^-, \hat{g}_\tau^+$ for each of the negative and positive values of g , yielding the contour regions $\{\mathbf{x} : \hat{g}(\mathbf{x}) \leq \hat{g}_\tau^-\}$ and $\{\mathbf{x} : \hat{g}(\mathbf{x}) \geq \hat{g}_\tau^+\}$, respectively. Whilst these contour regions do not have a probabilistic interpretation, $\hat{g}_\tau^-, \hat{g}_\tau^+$ remain suitable contour level choices for visualising $\hat{g}$. Thus Algorithm 2 can be adapted to compute contour levels for any grid-based function, and is not restricted to density functions.

This is the situation for the temperature anomaly data. Since these contour levels diverge from a neutral (zero) value to two opposing extreme values, a diverging colour scale is the most suitable for their visualisation (Zeileis et al. 2020).

We conclude with a brief description of several other contour level selection methods. The naive quantile contour levels compute the τ -quantile of all values in the grid $\hat{f}(\mathbf{g}_1), \dots, \hat{f}(\mathbf{g}_M)$. Suppose that we select m naive quantile contour levels, then for comparison, we can compute the m equal interval length contour levels, namely $\hat{f}_{\min} + j/m \cdot (\hat{f}_{\max} - \hat{f}_{\min}), j = 1, \dots, m$, or the m ‘natural’ or Jenks contour levels (Jenks 1967), which is a k -means clustering with m clusters of $\hat{f}(\mathbf{g}_1), \dots, \hat{f}(\mathbf{g}_M)$. These contour level selection methods are easy to understand and to implement. In the following section, we compare the empirical performance of these contour level selection methods to density contour levels for synthetic and experimental data.

3 Data analysis

To demonstrate the applicability of density contour levels for gridded data, we deploy the visualisation capabilities of the R statistical analysis environment. Due to the simplicity of Algorithm 2, it can be implemented in any visualisation software.

3.1 Synthetic data

We focus on two bivariate Gaussian mixture densities and a t -mixture density: Density #1 is a base case $N((-1, 0), [1/4, 0; 0, 1])$; Density #2 is trimodal with a central mode oriented obliquely from the coordinates axes $4/11N((-1, 1), 1/8[1, 0; 0, 1]) + 3/11N((0, 0), 1/8[1, 9/10; 9/10, 1]) + 4/11N((1, -1), 1/8[1, 0; 0, 1])$; Density #3 is bimodal with asymmetric modes and heavy tails $1/4t((-1, 0), [1/4, 0; 0, 1], 10) + 3/4t((1, 0), [1/4, 0; 0, 1], 10)$, where each t component has 10 d.f. For each density, we draw a Monte Carlo $N = 10^6$ random sample $\mathbf{X}_1, \dots, \mathbf{X}_N$, and we transform this random sample using the target mixture density f to obtain $f(\mathbf{X}_1), \dots, f(\mathbf{X}_N)$. The quantile-based approximation f_τ^* is the upper τ -quantile of $f(\mathbf{X}_1), \dots, f(\mathbf{X}_N)$, from which we can establish $R_\tau^* = \{\mathbf{x} : f(\mathbf{x}) \geq f_\tau^*\}$. We use f_τ^*, R_τ^* as the proxies for f_τ, R_τ . These $R_\tau^*, \tau = 0.1, 0.3, 0.5, 0.7, 0.9$ are displayed in the first row in Fig. 4.

We then draw another $n = 10,000$ random sample from the target mixture density, apply Algorithm 1 with a kernel density estimate $\hat{f}$, on an $M = 151 \times 151$ grid, to yield quantile-based estimates $\hat{f}_\tau$ and $\hat{R}_\tau$. The kernel density estimate is computed using a plug-in bandwidth (Duong and Hazelton 2003), though any sensible density estimate can be used here. To mimic the situation where we only have access to the gridded density estimate and we do not have access to the data sample, then we input only the grid of the previous kernel density estimate values $\hat{f}(\mathbf{g}_1), \dots, \hat{f}(\mathbf{g}_M)$ into Algorithm 2 to yield the grid-based approximations $\hat{f}_\tau$ and $\hat{R}_\tau$. The approximation computations are carried out by the `as.kde` function in the `ks` package (Duong 2005) or by the `tidy_as_kde/st_as_kde` functions in the `eks` package (Duong

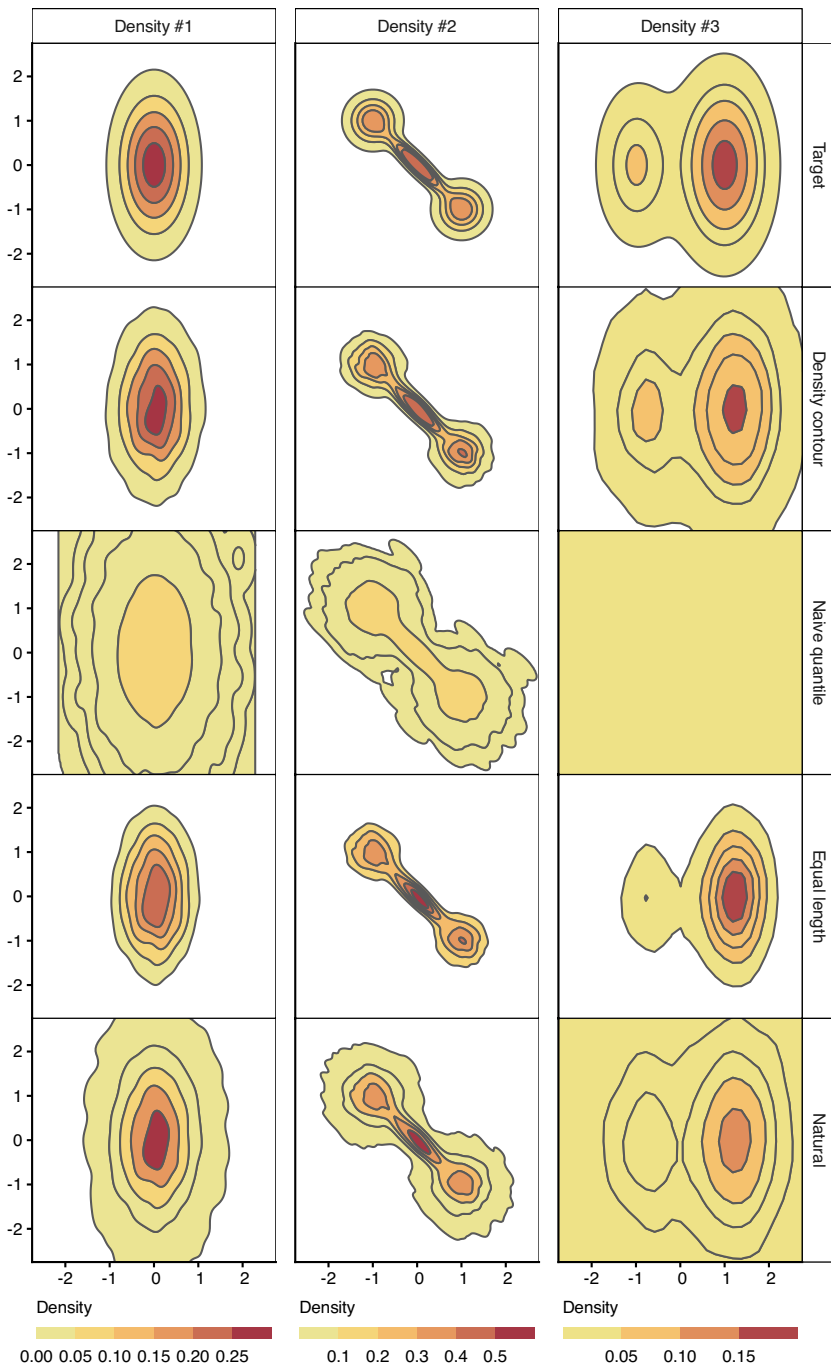


Fig. 4 Comparison of contour level selection methods for gridded density estimates from mixture densities, with discretised sequential ‘Heat’ colour scale, with $n = 10,000$ sample. (First row) Target proxy contour. (Second row) Density contour. (Third row) Naive quantile. (Fourth row) Equal length. (Fifth row) Natural (Jenks)

2022). These $\hat{R}_\tau$ are the second row in Fig. 4. We observe that the contour region estimates $\hat{R}_\tau$ generally remain close to R_τ^* . In Fig. 4, the third row is the naive quantiles, the fourth row the equal length intervals, and the fifth the natural contour levels. Overall, we observe that the density contours are the closest to the target proxy contours, followed by the natural and equal length contours, and then the naive quantile contours.

If we focus on Density #3, the proxy contour levels are $f_{0.9}^* = 0.014$, $f_{0.7}^* = 0.041$, $f_{0.5}^* = 0.071$, $f_{0.3}^* = 0.13$, $f_{0.1}^* = 0.20$. The density contour levels are $\hat{f}_{0.9} = 0.006$, $\hat{f}_{0.7} = 0.028$, $\hat{f}_{0.5} = 0.051$, $\hat{f}_{0.3} = 0.088$, $\hat{f}_{0.1} = 0.17$, and they describe the bimodality structure and account for the heavy tails well. The naive quantile contour levels are 0, 0, $1.87\text{e-}19$, $7.08\text{e-}19$, $2.26\text{e-}5$, as the vast majority of the gridded estimate values are close to zero; the contour plot with these contour regions does not reveal any structure. The equal length contour levels are 0.033, 0.066, 0.10, 0.13, 0.17, which describes both modes less well, especially the smaller mode. The natural contour levels $5.58\text{e-}5$, 0.013, 0.036, 0.072, 0.14, reveal some bimodality but suffer in comparison to the density contour and equal length levels. Natural contour levels are the results of a k -means clustering, and it well-known that this is not robust to large departures from Gaussianity and is sub-optimal in the presence of the heavy tails of a t -mixture.

Since Fig. 4 gives us good reason to believe that the proxy contours provide a good description of the structure of the underlying density, then we employ them for a quantitative comparison. That is, the target are the proxy decile contour regions R_τ^* , $\tau = 0.1, 0.3, 0.5, 0.7, 0.9$. To explore the performance of the contour level estimation methods, we consider two grid sizes $M = 51 \times 51$ and $M = 151 \times 151$, and two sample sizes $n = 1000$ and $n = 10,000$. For each combination of these parameter settings, we draw 100 replicates from the target mixture density. For each replicate, we compute the (i) density contour regions from Algorithm 2 with the same τ values, (ii) naive quantile contour regions with the same τ values, (iii) 5 equal length contour regions, and (iv) 5 natural contour regions. The equal length and natural contour levels estimates are not based on quantiles, so we estimate 5 of them each to align with the number of τ levels. Since these estimated contour regions are set estimates, the usual measure of their error to the target contour region R_τ^* is the probability of the symmetric difference between $\hat{R}_\tau$ and R_τ^* . That is, $\text{Err}(\hat{R}_\tau) = \mathbb{P}(\mathbf{X} \in \hat{R}_\tau \Delta R_\tau^*)$, where $A \Delta B = A \cup B \setminus (A \cap B)$ is the symmetric difference between sets A, B . These probabilities are approximated as Riemann sums.

These errors for the mixture densities are displayed in Table 1 ($n = 1000$) and Table 2 ($n = 10,000$). Within each sample size, we present the grid sizes 51×51 and 151×151 . Within each density, each row in the table corresponds to a contour region estimation method: density contour, naive quantile, equal length, natural. The columns are the target proxy contour regions R_τ^* , $\tau = 0.9, 0.7, 0.5, 0.3, 0.1$. Each table entry is the mean $\pm$ standard deviation of the Err of the symmetric difference of estimated and target contour regions.

Overall, the results for $n = 1000$ and $n = 10,000$ are broadly similar, with the larger sample size resulting in errors with lower mean and with smaller spread, all other factors being equal. The density contour estimates have the smallest errors in

Table 1 Errors of estimates of contour regions, for mixture densities with sample size $n = 1000$, and grid sizes $M = 51 \times 51, 151 \times 151$

Density	Estimator	Error of contour region estimator				
		$R_{0.9}^*$	$R_{0.7}^*$	$R_{0.5}^*$	$R_{0.3}^*$	$R_{0.1}^*$
$n = 1000, M = 51 \times 51$						
#1	Density contour	0.058±0.022	0.083±0.026	0.084±0.023	0.072±0.018	0.040±0.010
	Naive quantile	0.602±0.029	0.417±0.031	0.224±0.030	0.070±0.019	0.100±0.010
	Equal length	0.088±0.028	0.099±0.027	0.085±0.022	0.067±0.018	0.049±0.010
	Natural	0.088±0.021	0.120±0.018	0.083±0.021	0.068±0.019	0.036±0.009
#2	Density contour	0.103±0.029	0.238±0.056	0.269±0.074	0.224±0.063	0.135±0.036
	Naive quantile	0.735±0.021	0.572±0.028	0.432±0.037	0.262±0.048	0.130±0.033
	Equal length	0.082±0.025	0.216±0.041	0.259±0.067	0.217±0.063	0.123±0.035
	Natural	0.112±0.041	0.223±0.047	0.265±0.066	0.223±0.064	0.129±0.038
#3	Density contour	0.059±0.024	0.081±0.029	0.093±0.026	0.095±0.024	0.041±0.010
	Naive quantile	0.635±0.040	0.458±0.042	0.266±0.042	0.105±0.031	0.095±0.014
	Equal length	0.064±0.024	0.079±0.027	0.088±0.026	0.114±0.025	0.133±0.023
	Natural	0.070±0.023	0.100±0.026	0.132±0.031	0.112±0.029	0.059±0.019
$n = 1000, M = 151 \times 151$						
#1	Density contour	0.042±0.016	0.052±0.014	0.050±0.011	0.045±0.009	0.029±0.007
	Naive quantile	0.600±0.028	0.412±0.028	0.216±0.029	0.044±0.015	0.094±0.010
	Equal length	0.075±0.029	0.082±0.034	0.055±0.018	0.041±0.009	0.045±0.012
	Natural	0.072±0.018	0.120±0.018	0.059±0.015	0.039±0.008	0.024±0.005
#2	Density contour	0.038±0.016	0.123±0.025	0.117±0.018	0.097±0.013	0.064±0.009
	Naive quantile	0.706±0.020	0.517±0.020	0.357±0.021	0.164±0.020	0.080±0.026
	Equal length	0.035±0.013	0.136±0.025	0.136±0.036	0.094±0.021	0.060±0.013
	Natural	0.041±0.024	0.142±0.028	0.141±0.033	0.086±0.016	0.049±0.008
#3	Density contour	0.035±0.013	0.044±0.012	0.062±0.012	0.063±0.013	0.029±0.006
	Naive quantile	0.638±0.047	0.454±0.048	0.259±0.047	0.082±0.029	0.087±0.015
	Equal length	0.040±0.015	0.053±0.016	0.058±0.015	0.088±0.022	0.141±0.025
	Natural	0.043±0.014	0.068±0.026	0.107±0.025	0.112±0.029	0.065±0.020

Within each density, each row is contour region estimation method: density contour, naive quantile, equal length, natural. The columns are proxy contour regions R_{τ}^* , $\tau = 0.9, 0.7, 0.5, 0.3, 0.1$. Each table entry is the mean $\pm$ SD of error of symmetric difference of estimated and target contour regions, bold entry has smallest mean error for each R_{τ}^*

the majority of cases. In the remaining minority of cases, they are usually second by a small margin. On the other hand, the naive quantile estimates are almost always the worst, and in some cases, by a large margin. The equal length and natural contours can be effective in some cases, with a slight advantage for the natural contour regions.

The finer 151×151 grids lead to lower errors than the coarser 51×51 grids, though the magnitude of the error reduction for an increased grid size depends on the data sample. For example, the largest decreases are for Density #2 with its intricate multimodality, since a finer grid is able to more closely follow this complicated structure. Whilst a larger grid size would be better than a coarser grid in general, we recall that the grid size of gridded data is usually decided by the data provider. So the grid size for data analysis is not a tuning parameter that can be changed easily like in this simulation study.

Table 2 Errors of estimates of contour regions, for mixture densities with sample size $n = 10,000$, and grid sizes $M = 51 \times 51, 151 \times 151$

Density	Estimator	Error of contour region estimator				
		$R_{0.9}^*$	$R_{0.7}^*$	$R_{0.5}^*$	$R_{0.3}^*$	$R_{0.1}^*$
$n = 10,000, M = 51 \times 51$						
#1	Density contour	0.042±0.017	0.058±0.025	0.058±0.025	0.049±0.020	0.026±0.009
	Naive quantile	0.657±0.025	0.463±0.025	0.262±0.024	0.076±0.019	0.098±0.006
	Equal length	0.079±0.020	0.071±0.023	0.060±0.024	0.049±0.019	0.053±0.005
	Natural	0.079±0.014	0.119±0.015	0.064±0.020	0.048±0.020	0.027±0.008
#2	Density contour	0.112±0.015	0.235±0.035	0.286±0.042	0.232±0.034	0.122±0.020
	Naive quantile	0.751±0.015	0.577±0.019	0.433±0.022	0.257±0.024	0.129±0.022
	Equal length	0.085±0.015	0.244±0.033	0.262±0.034	0.229±0.034	0.122±0.016
	Natural	0.103±0.019	0.264±0.034	0.267±0.037	0.231±0.034	0.120±0.020
#3	Density contour	0.065±0.025	0.087±0.033	0.094±0.032	0.088±0.029	0.037±0.012
	Naive quantile	0.745±0.032	0.561±0.031	0.367±0.030	0.170±0.026	0.062±0.027
	Equal length	0.070±0.026	0.091±0.026	0.092±0.030	0.100±0.026	0.150±0.011
	Natural	0.075±0.026	0.105±0.027	0.127±0.027	0.119±0.028	0.061±0.014
$n = 10,000, M = 151 \times 151$						
#1	Density contour	0.027±0.011	0.032±0.010	0.029±0.009	0.025±0.006	0.015±0.003
	Naive quantile	0.644±0.024	0.447±0.024	0.246±0.024	0.053±0.020	0.093±0.006
	Equal length	0.074±0.023	0.053±0.022	0.031±0.010	0.028±0.008	0.056±0.006
	Natural	0.067±0.012	0.130±0.011	0.056±0.011	0.024±0.006	0.025±0.004
#2	Density contour	0.029±0.015	0.068±0.022	0.073±0.023	0.057±0.016	0.030±0.006
	Naive quantile	0.709±0.018	0.501±0.019	0.334±0.019	0.134±0.020	0.085±0.017
	Equal length	0.029±0.014	0.076±0.024	0.128±0.022	0.092±0.019	0.075±0.009
	Natural	0.036±0.010	0.076±0.028	0.155±0.030	0.061±0.017	0.033±0.007
#3	Density contour	0.026±0.010	0.031±0.010	0.039±0.008	0.037±0.009	0.016±0.003
	Naive quantile	0.736±0.033	0.544±0.033	0.347±0.033	0.147±0.033	0.064±0.026
	Equal length	0.030±0.012	0.049±0.011	0.043±0.009	0.051±0.012	0.165±0.013
	Natural	0.032±0.010	0.045±0.009	0.077±0.013	0.125±0.012	0.080±0.013

Within each density, each row is contour region estimation method: density contour, naive quantile, equal length, natural. The columns are proxy contour regions R_{τ}^* , $\tau = 0.9, 0.7, 0.5, 0.3, 0.1$. Each table entry is the mean $\pm$ SD of error of symmetric difference of estimated and target contour regions, bold entry has smallest mean error for each R_{τ}^*

Lastly, a comparison of the execution times for the contour level section methods is given in Table 3. Each table entry is the mean $\pm$ SD of execution time (ms) for the $N = 100$ replicates carried out on a set-up running Ubuntu 24.04 with 16Gb RAM and $4 \times$ Intel i5 CPU@3.10 GHz. Execution times are difficult to replicate, so these are given only as a guide. The sample size n and the target density have only a small effect on the execution times compared to the grid size M . The naive quantile and equal length are highly efficient, and their execution times only increase marginally from $M = 51 \times 51$ to $M = 151 \times 151$. Their disadvantage is that they rarely yield useful data visualisations. The density contour levels are more computationally intensive, and their execution times increase about 9-fold for a 9-fold increase in grid size from $M = 51 \times 51$ to $M = 151 \times 151$. The relative increase in the execution times for the natural contour levels are similar, though their baseline execution times are higher than for the density contour levels.

Table 3 Execution times for estimates of contour levels, for mixture densities with sample size $n = 1000, 10,000$, and grid sizes $M = 51 \times 51, 151 \times 151$

Density	Estimator	Execution time (ms)			
		$n = 1000$		$n = 10,000$	
		$M = 51 \times 51$	$M = 151 \times 151$	$M = 51 \times 51$	$M = 151 \times 151$
#1	Density contour	4.43±0.82	4.65±0.89	35.32± 7.01	33.30± 3.06
	Naive quantile	0.35±0.48	0.33±0.47	0.62± 0.65	0.51± 0.54
	Equal length	0.59±0.51	0.68±0.49	1.35± 0.68	1.22± 0.50
	Natural	8.50±0.91	8.73±0.81	90.44± 9.54	88.93± 7.69
#2	Density contour	4.48±0.61	4.64±0.74	33.10± 2.51	32.78± 2.01
	Naive quantile	0.32±0.47	0.37±0.48	0.56± 0.52	0.76± 0.49
	Equal length	0.68±0.47	0.63±0.50	1.28± 0.49	1.24± 0.47
	Natural	8.47±1.19	8.45±1.05	87.77±11.43	89.87±10.14
#3	Density contour	4.57±0.60	4.66±1.24	33.73± 5.65	37.37± 9.05
	Naive quantile	0.33±0.47	0.26±0.44	0.57± 0.53	0.85± 0.74
	Equal length	0.72±0.47	0.68±0.51	1.20± 0.42	1.50± 0.81
	Natural	9.11±0.86	8.94±1.33	89.06±10.21	93.08±12.78

Within each density, each row is contour region estimation method: density contour, naive quantile, equal length, natural. Each table entry is the mean ± SD of execution time (ms)

3.2 Experimental data

For our analysis of experimental data, we return to the gridded data from Figs. 1–2. For the Paris population data, there are 2078 $1 \text{ km} \times 1 \text{ km}$ square polygons in the Lambert Azimuthal Equal Area Europe (EPSG 3035) projected coordinates system. Since we rely on software which requires density estimates on uniform rectangular grids aligned to the coordinate axes, so we interpolate to make them conform. For the population data, since it is mostly a uniform rectangular grid, we are only required to create the missing grid cells with population zero: the result is a grid of $M = 57 \times 43 = 2451$ grid cells ($1 \text{ km} \times 1 \text{ km}$).

In Fig. 5a are the contour regions of the Paris population interpolated gridded data corresponding to the density contour levels: $\hat{f}_{0.1} = 30,237$ (red), $\hat{f}_{0.3} = 15,288$ (dark orange), $\hat{f}_{0.5} = 9606$ (mid orange), $\hat{f}_{0.7} = 6410$ (light orange), $\hat{f}_{0.9} = 3186$ (yellow). In comparison to the continuous colour scale, these contour regions provide a more interpretable visualisation of the population distribution. The population hotspots can be more quantitatively defined as, say, the 10% density contour region $\hat{R}_{0.1}$. That is, we can consider any area with population density of at least 30,237 per km^2 to be densely populated with respect to the rest of the ROI. The area of $\hat{R}_{0.1}$ is 22.29 km^2 , which is 0.91% of the total ROI area of 2451 km^2 . Whereas the population in $\hat{R}_{0.1}$ is 957,473.5 which is 10.09% of the total ROI population of 9,490,878. These contour regions quantify the highly unequal population spatial distribution in the ROI. In contrast, the naive quantile contour levels are 10,318 (red), 4478 (dark orange), 1446 (mid orange), 122 (light orange), 1 (yellow) in Fig. 5b. Since zero or low population grid cells are by far the most numerous, even though their sum does not rival the population sum of the less numerous densely populated grid cells. So these contour regions give the erroneous impression that the population distribution is more evenly spread throughout the ROI.

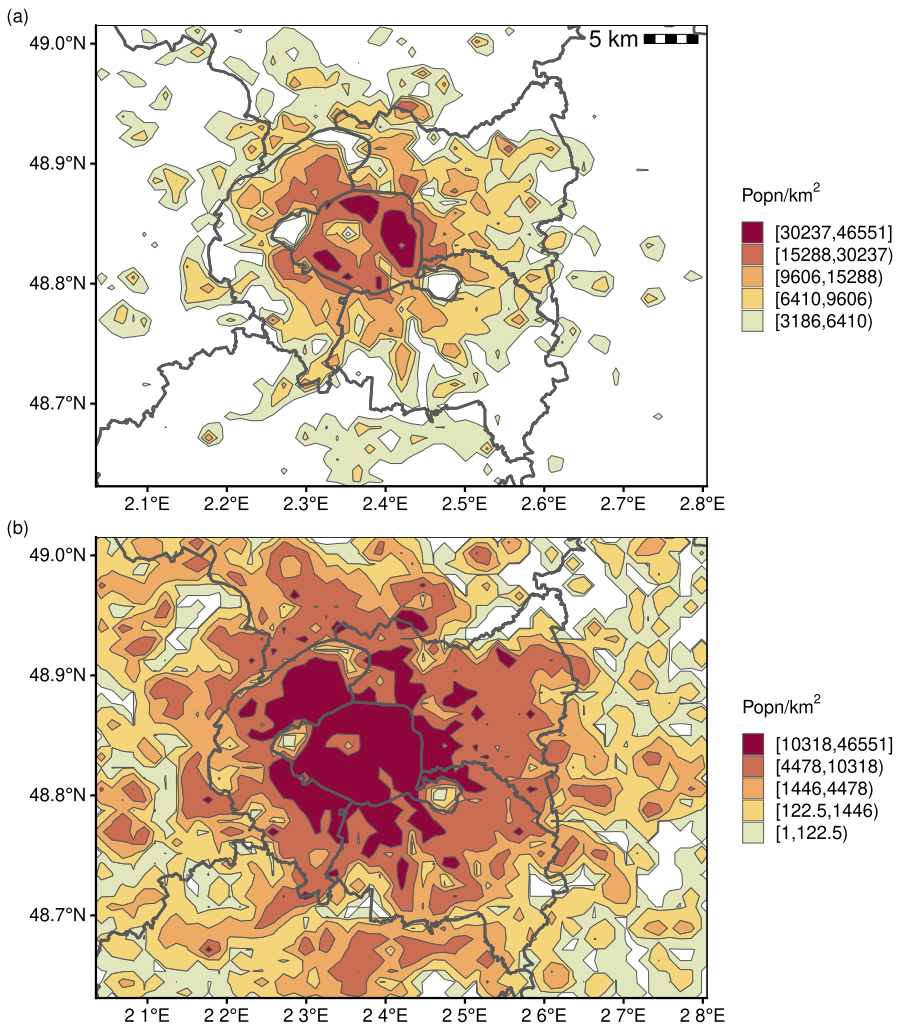


Fig. 5 Contour regions for population density for Paris capital region on $1 \text{ km} \times 1 \text{ km}$ grid, with discretised sequential ‘Heat’ colour scale. **(a)** Density contour levels. **(b)** Naive quantile contour levels

The equal length contour regions are 38,792 (red), 31,034 (dark orange), 23,276 (mid orange), 15,517 (light orange), 7758 (yellow) in Fig. 6a. These contour levels emphasise the higher population densities: the higher levels delimit well the hotspots, but the lower levels do not delimit the extent of the spatial population distribution as well as the lower density contour levels in Fig. 6a. The natural contour levels 31,532 (red), 16,273 (dark orange), 8964 (mid orange), 4518 (light orange), 520 (yellow) in Fig. 6b reveal a similar contour plot with density contour levels, though the lower levels give the visual impression that the population distribution is more evenly spread in the peripheral areas of the ROI.

For the year-monthly temperature anomaly time series in Western Australia, we are able to compute the analogous contour levels. Since these

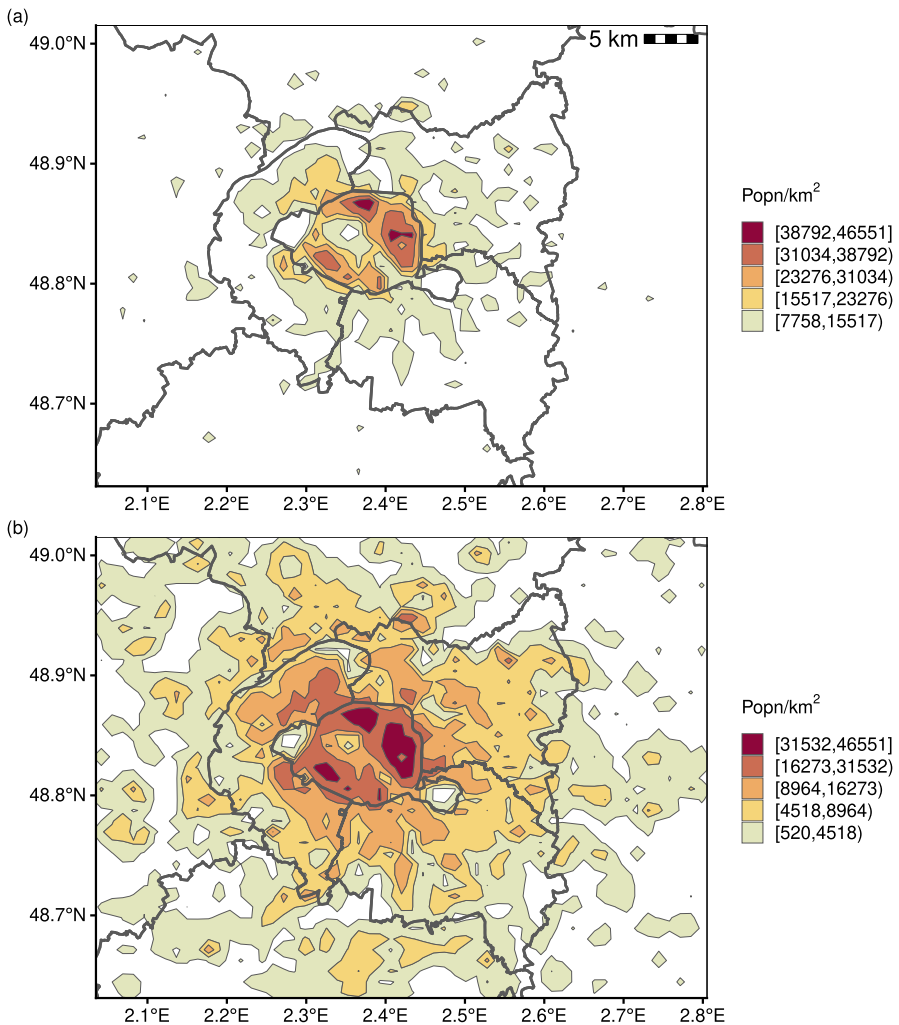


Fig. 6 Contour regions for population density for Paris capital region on 1 km × 1 km grid, with discretised sequential 'Heat' colour scale. **(a)** Equal length contour levels. **(b)** Natural (Jenks) contour levels

are non-spatial gridded data, we do not compute contour regions as above and instead we visualise the time series heat maps with colours corresponding to the contour levels. In Fig. 7a are the density contour levels, $\hat{g}_{0.25}^- = -1.04$, $\hat{g}_{0.5}^- = -0.74$, $\hat{g}_{0.75}^- = -0.49$, $\hat{g}_{0.75}^+ = 0.55$, $\hat{g}_{0.5}^+ = 0.86$, $\hat{g}_{0.25}^+ = 1.18$. There is a large interval of anomalies in $[-0.49^\circ\text{C}, 0.55^\circ\text{C}]$ that are considered to be neutral (grey), so the warming trends (above 1.18°C in dark red) since the 2000s and in the second half of the year become more prominent. For the naive quartile contour levels in Fig. 7b, most grid cells are coloured in mid/dark blue or mid/dark red, which provides a less concise summary of the warming trends. The equal length contour levels in Fig. 8a gives the impression of more gradual warming trends since most grid cells are coloured in light blue, grey or light red. The natural contour levels Fig. 8b

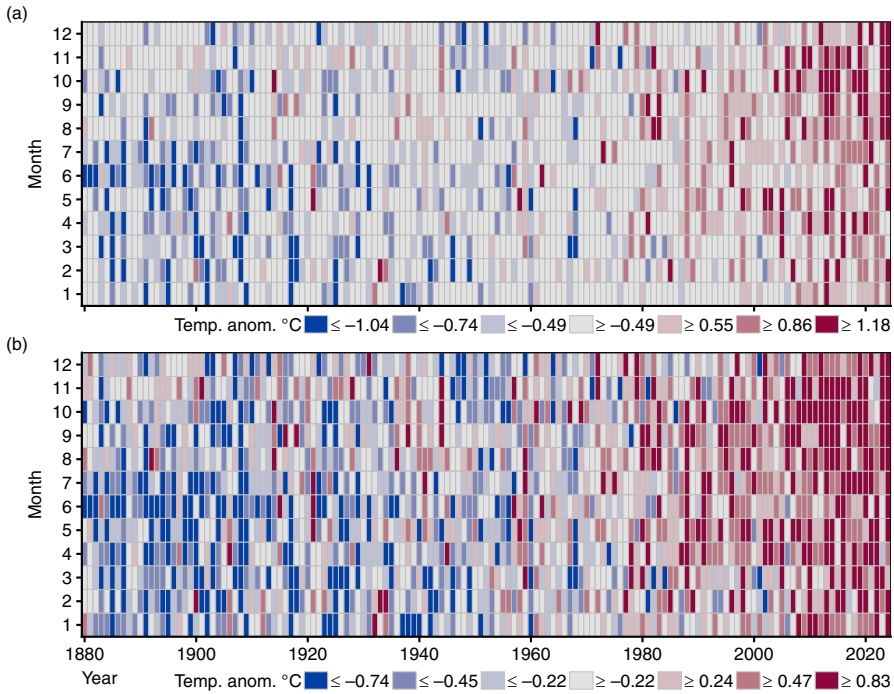


Fig. 7 Heat maps for year-monthly temperature anomaly time series, with discretised diverging ‘Red–Blue’ colour scale. **(a)** Density contour levels. **(b)** Naive quantile contour levels

is the most similar to the density contour levels, though the warming trends are less prominent due to the larger number of mid blue/mid red grid cells spread more evenly through the time series heat map.

For these experimental data sets, we do not have the proxies for the ground truth like for the simulated data to which we can compare the estimated contour regions. Nonetheless, to gain insight into the robustness of these estimation methods, we perform a simple sensitivity analysis. Let $\mathbf{X}_1, \dots, \mathbf{X}_M$ be the experimental data set with M grid cells, and let the τ -level density levels/contours be $\hat{f}_\tau, \hat{R}_\tau$. These are reasonable targets since from the simulated data, these density contours were the best performing estimates. We generate a simulated replicate in each grid cell from a model of the experimental data: for the Paris population data, this is a Poisson model with mean equal to the observed population, and for the temperature anomaly data, a Gaussian model with mean equal to the observed temperature anomaly and standard deviation equal to 0.25°C . From this data replicate, we compute the density contour, naive quantile, equal length and natural contour levels/regions. For the Paris population data, we compute the error of the symmetric difference as in the simulation study, except that the proxies are replaced by the original density contour regions $\hat{R}_{0.9}, \dots, \hat{R}_{0.1}$. For the temperature anomaly time series, contour regions are not required, so we compute the errors to the original density contour levels $\hat{g}_{0.25}^-, \hat{g}_{0.5}^-, \hat{g}_{0.75}^-, \hat{g}_{0.75}^+, \hat{g}_{0.5}^+, \hat{g}_{0.25}^+$. We carry out 100 replicates, and the errors are summarised in Tables 4–5.

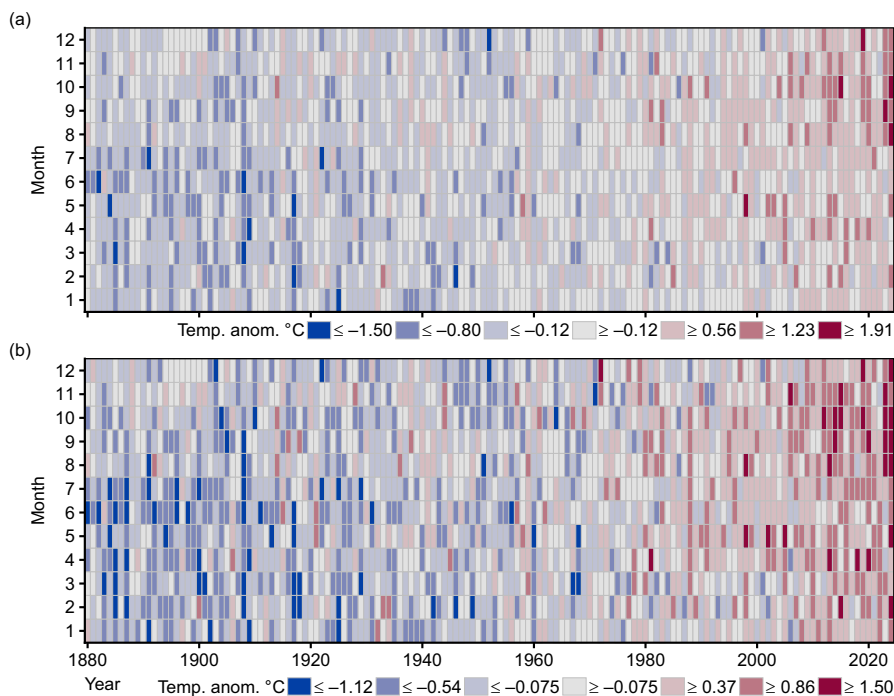


Fig. 8 Heat maps for year-monthly temperature anomaly time series, with discretised diverging ‘Red–Blue’ colour scale. **(a)** Equal length contour levels. **(b)** Natural (Jenks) contour levels

Table 4 Errors of estimates of contour regions for the Paris population data

Estimator	Error of contour region estimator				
	$\hat{R}_{0.9}$	$\hat{R}_{0.7}$	$\hat{R}_{0.5}$	$\hat{R}_{0.3}$	$\hat{R}_{0.1}$
<i>Paris population</i>					
Density contour	0.002±0.001	0.004±0.001	0.005±0.002	0.004±0.002	0.000±0.000
Naive quantile	0.074±0.000	0.126±0.000	0.030±0.000	0.170±0.000	0.371±0.000
Equal length	0.289±0.003	0.089±0.003	0.111±0.003	0.010±0.002	0.004±0.002
Natural	0.075±0.002	0.124±0.002	0.033±0.005	0.028±0.004	0.008±0.002

Each row is contour region estimation method: density contour, naive quantile, equal length, natural. The columns are original contour regions $\hat{R}_{0.9}$, $\hat{R}_{0.7}$, $\hat{R}_{0.5}$, $\hat{R}_{0.3}$, $\hat{R}_{0.1}$. Each table entry is the mean ± SD of error of symmetric difference of estimated and original contour regions

Since the targets are the density contour levels/regions of the original data set, then we expect that the mean errors of the density contour levels/regions of the replicates (first row) to be close to them. The objective of this sensitivity analysis is gleaned more from the variability of these errors of the contour levels/regions. The naive quantile contours have the lowest variation because their contour levels are highly biased towards zero for almost all cases: so they are robust but with consistently large bias. On the other hand, the equal length contours are the most variable. Next are the natural contours, which are more robust. The density

Table 5 Errors of estimates of contour levels for the temperature anomaly times series

Estimator	Error of contour level estimator					
	$\hat{g}_{0.25}^-$	$\hat{g}_{0.5}^-$	$\hat{g}_{0.75}^-$	$\hat{g}_{0.75}^+$	$\hat{g}_{0.5}^+$	$\hat{g}_{0.25}^+$
<i>Temperature anomaly</i>						
Density contour	0.089±0.026	0.066±0.016	0.036±0.014	0.035±0.013	0.053±0.019	0.085±0.026
Naive quantile	0.243±0.015	0.266±0.013	0.260±0.012	0.308±0.013	0.341±0.014	0.310±0.018
Equal length	0.509±0.158	0.131±0.104	0.366±0.129	0.110±0.077	0.444±0.143	0.829±0.164
Natural	0.180±0.055	0.120±0.054	0.341±0.060	0.234±0.068	0.065±0.049	0.333±0.097

Each row is contour level estimation method: density contour, naive quantile, equal length, natural. The columns are original contour levels $\hat{g}_{0.25}^-$, $\hat{g}_{0.5}^-$, $\hat{g}_{0.75}^-$, $\hat{g}_{0.75}^+$, $\hat{g}_{0.5}^+$, $\hat{g}_{0.25}^+$. Each table entry is the mean ± SD of difference of estimated and original contour levels

contours have the lowest variation, apart from the highly biased naive quantile contours. So they are robust to small changes in the input gridded data.

4 Conclusion

Gridded data are an important data format for the dissemination and storage of open data. Density contour levels are widely deployed to create interpretable visualisations, and we have introduced algorithms to compute density contour levels for any type of numerical gridded data. We demonstrated that for synthetic and experimental gridded data, our proposal leads to more interpretable visualisations than the current contour level selection methods. These algorithms are available in R packages on the official CRAN repository to facilitate their wider deployment for statistical data analysis.

Supplementary Information The online version contains supplementary material available at <https://doi.org/10.1007/s00180-026-01767-x>.

Funding The authors did not receive support from any organisation for the submitted work.

Data Availability The French population census gridded data are available at <https://www.insee.fr/statistiques/7655464?sommaire=7655515>. The global surface temperature anomaly gridded data (GISTEMP) are available at <https://data.giss.nasa.gov/gistemp>.

Code Availability The R packages are available on CRAN: ks <https://cran.r-project.org/package=ks> and eks <https://cran.r-project.org/package=eks>. A companion R script of the analyses carried in this manuscript is available as Online Resource 1.

Declarations

Conflict of interest The authors have no conflict of interest to declare that are relevant to the content of this article.

References

- Bowman AW, Foster P (1993) Adaptive smoothing and density-based tests of multivariate normality. *J Am Stat Assoc* 88:529–537. <https://doi.org/10.1080/01621459.1993.10476304>
- Chacón JE, Duong T (2018) Multivariate kernel smoothing and its applications. Chapman and Hall/CRC, <https://doi.org/10.1201/9780429485572>
- Duong T, Hazelton ML (2003) Plug-in bandwidth matrices for bivariate kernel density estimation. *J Non-parametric Stat* 15:17–30. <https://doi.org/10.1080/10485250306039>
- Dobson MW (1973) Choropleth maps without class intervals?: A comment. *Geogr Anal* 5:358–360. <https://doi.org/10.1111/j.1538-4632.1973.tb00498.x>
- Duong T (2005) ks: Kernel smoothing. R package version 1.15.2. <https://CRAN.R-project.org/package=ks>
- Duong T (2022) eks: Tidy and geospatial kernel smoothing. R package version 1.1.2. <https://CRAN.R-project.org/package=eks>
- Delicado P, Vieu P (2015) Optimal level sets for bivariate density representation. *J Multivar Anal* 140:1–18. <https://doi.org/10.1016/j.jmva.2015.04.005>
- GISTEMP Team: GISS Surface Temperature Analysis (GISTEMP), version 4. <https://data.giss.nasa.gov/gistemp> (2025)
- Hyndman RJ (1996) Computing and graphing highest density regions. *Am Stat* 50:120–126. <https://doi.org/10.1080/00031305.1996.10474359>
- Insee: Revenus, pauvreté et niveau de vie en 2019 - Données carroyées [In French]. <https://www.insee.fr/statistiques/7655511?sommaire=7655515> (2019)
- Jenks GF (1963) Generalization in statistical mapping. *Ann Assoc Am Geogr* 53:15–26. <https://doi.org/10.1111/j.1467-8306.1963.tb00429.x>
- Jenks GF (1967) The data model concept in statistical mapping. *Int Yearbook Cartogr* 7:186–190. <https://cir.nii.ac.jp/crid/1573668925394541312>
- Klemelä J (2009) Smoothing of multivariate data: density estimation and visualization. Wiley. <https://doi.org/10.1002/9780470425671>
- Lenssen N, Schmidt GA, Hendrickson M, Jacobs P, Menne MJ, Ruedy R (2024) A NASA GISTEMPv4 observational uncertainty ensemble. *J Geophysical Res Atmospheres* 129(17):2023-JD040179. <https://doi.org/10.1029/2023JD040179>
- Polonik W (1995) Measuring mass concentrations and estimating density contour clusters – an excess mass approach. *Ann Stat* 23:855–881. <https://doi.org/10.1214/aos/1176324626>
- Tobler WR (1973) Choropleth maps without class intervals? *Geogr Anal* 5:262–265. <https://doi.org/10.1111/j.1538-4632.1973.tb01012.x>
- Vose RS, Huang B, Yin X, Arndt D, Easterling DR, Lawrimore JH, Menne MJ, Sanchez-Lugo A, Zhang HM (2021) Implementing full spatial coverage in NOAA's global temperature analysis. *Geophys Res Lett* 48:2020–090873. <https://doi.org/10.1029/2020GL090873>
- Zeileis A, Fisher JC, Hornik K, Ihaka R, McWhite CD, Murrell P, Stauffer R, Wilke CO (2020) Color-space: a toolbox for manipulating and assessing colors and palettes. *J Stat Softw* 96(1):1–49. <https://doi.org/10.18637/jss.v096.i01>
- Zeileis A, Hornik K, Murrell P (2009) Escaping RGBland: selecting colors for statistical graphics. *Comput Stat Data Anal* 53:3259–3270. <https://doi.org/10.1016/j.csda.2008.11.033>

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.