



Robust Inference for Mixtures of Multivariate Wrapped Normal Distributions: The Weighted Likelihood Approach

ICORS 2026, Istanbul, Türkiye

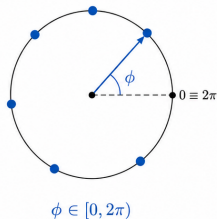
Giulia Bertagnolli*, L. Greco and C. Agostinelli

*Free University of Bozen-Bolzano

2026-07-23

Directional data

1. Univariate directional data

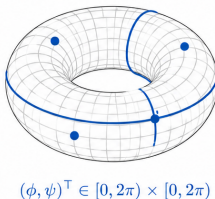


Example:

Wind direction

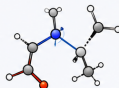


2. Bivariate directional data



Example:

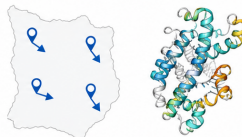
Protein dihedral angles



3. Multivariate directional data (p -dimensional torus)

$$\mathbb{T}^p = \underbrace{S^1 \times \dots \times S^1}_p$$

$$\phi = (\phi_1, \dots, \phi_p)^T \in [0, 2\pi)^p$$



Examples:

Wind directions
at different stations

RNA
(rotameric nature)

- Key feature of directional data: **periodicity**

Multivariate Wrapped Distributions I

Definition (Wrapped Distributions)

Consider $\mathbf{X} \in \mathbb{R}^p$ with density function $m(\mathbf{x}, \theta)$ and distribution function $M(\mathbf{x}, \theta)$ with $\theta \in \Theta$. The wrapped random variable $\mathbf{Y} = \mathbf{X} \bmod 2\pi$ has density function

$$m^\circ(\mathbf{y}; \theta) = \sum_{\mathbf{j} \in \mathbb{Z}^p} m(\mathbf{y} + 2\pi\mathbf{j}; \theta)$$

and distribution function

$$M^\circ(\mathbf{y}; \theta) = \sum_{\mathbf{j} \in \mathbb{Z}^p} [M(\mathbf{y} + 2\pi\mathbf{j}; \theta) - M(2\pi\mathbf{j}; \theta)],$$

where $\mathbf{y} \in \mathbb{T}^p \equiv (0, 2\pi]^p$.

Multivariate Wrapped Distributions II

Wrapped Unimodal Elliptically Symmetric distributions:

$$m(\mathbf{x}; \boldsymbol{\theta}) \propto |\Sigma|^{-1/2} h \left((\mathbf{x} - \boldsymbol{\mu})^\top \Sigma^{-1} (\mathbf{x} - \boldsymbol{\mu}) \right)$$

where here we assume that h is a non-increasing function with continuous first derivative, and $\boldsymbol{\theta} = (\boldsymbol{\mu}, \Sigma)$.

The Wrapped Normal (WN) distribution corresponds to $h(t) = \exp(-\frac{t}{2})$.

Maximum Likelihood I

Given a sample $(\mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_n)$ from $\mathbf{Y} \sim m^\circ(\cdot; \boldsymbol{\theta})$

$$\begin{aligned}\text{maximise:} \quad \ell^\circ(\boldsymbol{\theta}) &= \sum_{i=1}^n \log m^\circ(\mathbf{y}_i; \boldsymbol{\theta}) \\ &= \sum_{i=1}^n \log \sum_{\mathbf{j} \in \mathbb{Z}^p} m(\mathbf{y}_i + 2\pi\mathbf{j}; \boldsymbol{\theta}) = \dots\end{aligned}$$

$$u^\circ(\mathbf{y}; \boldsymbol{\theta}) = \nabla \log m^\circ(\mathbf{y}; \boldsymbol{\theta}) = \frac{\nabla m^\circ(\mathbf{y}; \boldsymbol{\theta})}{m^\circ(\mathbf{y}; \boldsymbol{\theta})}$$

$$\text{solve:} \quad \sum_{i=1}^n u^\circ(\mathbf{y}_i; \boldsymbol{\theta}) = 0$$

Maximum Likelihood II

- For the Wrapped Unimodal Elliptically Symmetric model: MLEs $\hat{\mu}, \hat{\Sigma}$ are the solution of a pair of fixed point equations (through an IRW procedure) (Agostinelli, Greco, and Saraceno 2024)
- For the WN model ($h(t) = \exp(-\frac{t}{2})$): same MLEs (approx. MLEs) from **EM** and (resp. **Classification EM**) algorithms as in (Nodehi et al. 2021), where wrapping coefficients $j \in \mathbb{Z}^p$ are treated as **latent variables**

$$\ell_c(\boldsymbol{\theta} \mid \mathbf{y}, \mathbf{j}) = \sum_{i=1}^n \sum_{\mathbf{j}_i \in \mathbb{Z}^p} v_{ij_i} \log m(\mathbf{y}_i + 2\pi \mathbf{j}_i; \boldsymbol{\theta}),$$
$$\text{with } v_{ij} = \frac{m(\mathbf{y}_i + 2\pi \mathbf{j}; \boldsymbol{\theta})}{\sum_{\mathbf{k} \in \mathbb{Z}^p} m(\mathbf{y}_i + 2\pi \mathbf{k}; \boldsymbol{\theta})}$$

Mixtures of Wrapped Normal Distributions I

Consider a finite mixture model whose K components are distributed according to a p -dimensional WN variate (Greco, Inverardi, and Agostinelli 2023):

$$f^\circ(\mathbf{y}; \boldsymbol{\theta}) = \sum_{k=1}^K \omega_k m_k^\circ(\mathbf{y}; \boldsymbol{\theta}_k)$$

with $\boldsymbol{\theta} = (\omega_1, \dots, \omega_K, \boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_K)$:

- ω_k : membership probabilities
- $\boldsymbol{\theta}_k$: component specific parameters.

The operation of mixing and wrapping commute $f^\circ(\mathbf{y}; \boldsymbol{\theta}) = \sum_{\mathbf{j} \in \mathbb{Z}^p} f(\mathbf{y} + 2\pi\mathbf{j}; \boldsymbol{\theta})$.

Estimation of Mixtures of WNs I

Efficient EM algorithm based on the complete log-likelihood

$$\ell_c(\boldsymbol{\theta}) = \sum_{i=1}^n \sum_{k=1}^K z_{ik} \left\{ \log \omega_k + \sum_{\mathbf{j} \in \mathbb{Z}^p} v_{ik}(\mathbf{j}) \log m(\mathbf{y}_i + 2\pi \mathbf{j}; \boldsymbol{\theta}_k) \right\} \quad (1)$$

- $\mathbf{z}_i = (z_{i1}, z_{i2}, \dots, z_{iK})$ denote the i -th observation component label binary vector (Assumpt.: realizations of multinomial random variables $\mathbf{Z}_i$)
- $\mathbf{v}_i = (v_{i1}, v_{i2}, \dots, v_{iK})$ denotes the i -th component wrapping index vector, i.e. $v_{ik}(\mathbf{j}) = 1$ if obs. i has wrapping coefficients $\mathbf{j}$ in group k ($\mathbf{V}_i$)
- for each $\mathbf{y}_i$ we have $(\mathbf{z}_i, \mathbf{v}_i)$ missing components.

Estimation of Mixtures of WNs II

Evaluate the conditional expected value of the complete log-likelihood in (Equation 1) w.r.t. the joint distribution of (Z, V) given the data and current parameter values:

$$Q(\theta|\theta') = \mathbb{E}_Z \left[\mathbb{E}_{V|Z} [\ell_c(\theta)|y, \theta'] \right]$$

- **E-step:** *outer* for group memberships Z_i and *inner* for V_i given the data, current parameter values, and Z_i .
- **M-step:** parameter update according to $\arg \max_{\theta} Q(\theta|\theta')$
 - ▶ For the (unimodal) elliptically symmetric model: in the M-step we solve a pair of fixed point equations for μ_k and Σ_k for each $k = 1, \dots, K$.

Outliers

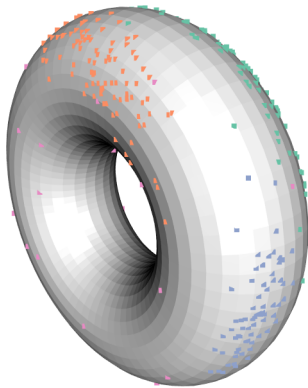


Figure 1: A three-components WN mixture and scattered outliers on a 2-torus.

The probabilistic approach:

- Outliers are “observations that are highly unlikely to occur under the assumed model” (Markatou, Basu, and Lindsay 1998; Agostinelli 2007)

Background on WLE I

Agreement with the model is quantified by Pearson's residuals (Lindsay 1994):

$$\delta(\mathbf{y}; m_{\theta}^{\circ}, f^{\circ}) = \frac{\hat{f}_n^{\circ}(\mathbf{y})}{\hat{m}^{\circ}(\mathbf{y}; \theta)} - 1$$

Solve weighted likelihood estimating equations (WLEEs):

$$\sum_{i=1}^n w(\delta_n^{\circ}(\mathbf{y}_i; \theta)) u^{\circ}(\mathbf{y}_i; \theta) = 0$$

where weights are obtained by *adjusting*¹ large residuals, so that outliers are down-weighted in the WLEE (Markatou, Basu, and Lindsay 1998).

¹through a **Residual Adjustment Function** (RAF) (Lindsay 1994)

A note on RAFs and weights

A Residual Adjustment Function (RAF) (Lindsay 1994) is a function $A(\delta)$, which is increasing and twice differentiable in $[-1, +\infty)$ and satisfies $A(0) = 0$ and $A'(0) = 1$.

- A RAF bounds the effect of large Pearson's residuals.
- RAFs connect WLE to *minimum divergence estimation*, see (Kuchibhotla and Basu 2018).

Weights are then defined as:

$$w(\delta) = \min \left\{ \frac{[A(\delta) + 1]^+}{\delta + 1}, 1 \right\}$$

- Weights decline smoothly to zero as $\delta \rightarrow \infty$ (outliers) and, depending on the RAF, also for $\delta \rightarrow -1$ (inliers).

Robust estimation of mixtures of WNs I

Weights are specific for each observations $\mathbf{y}_i$ and label k of the group

$$\delta_k^\circ(\mathbf{y}; \boldsymbol{\theta}) = \frac{\hat{f}_k^\circ(\mathbf{y})}{\hat{m}^\circ(\mathbf{y}; \boldsymbol{\theta})} - 1$$

where

$$\hat{f}_k^\circ(\mathbf{y}) = \frac{1}{\sum_{i=1}^n z_{ik}} \sum_{i=1}^n z_{ik} k^\circ(\mathbf{y}; \mathbf{y}_i, H)$$

with (kernel) density k° on the torus and bandwidth H , so that

$$w_{ik} = w(\delta_k^\circ(\mathbf{y}_i; \boldsymbol{\theta}_k))$$

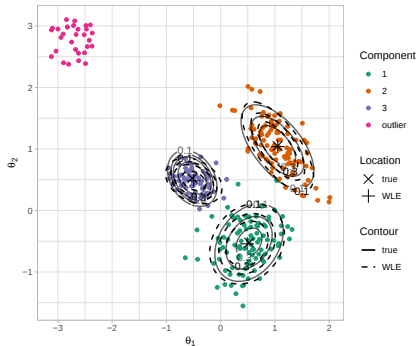
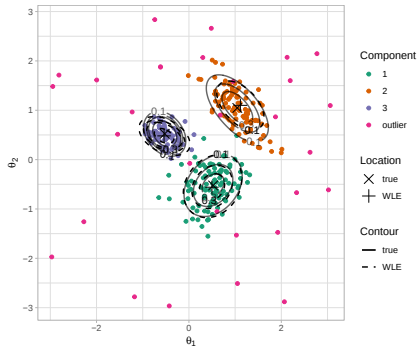
Robust estimation of mixtures of WNs II

In the M-step we solve **weighted** fixed point equations:

$$\boldsymbol{\mu}_k = \frac{\sum_{i=1}^n z_{ik} w_{ik} \sum_{\mathbf{j} \in \mathbb{Z}^p} v_{ik}(\mathbf{j})(\mathbf{y}_i + 2\pi \mathbf{j})}{\sum_{i=1}^n z_{ik} w_{ik} \sum_{\mathbf{a} \in \mathbb{Z}^p} v_{ik}(\mathbf{a})}$$

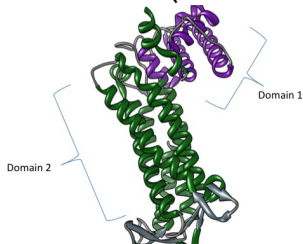
$$\boldsymbol{\Sigma}_k = -\frac{2}{\sum_{i=1}^n z_{ik} w_{ik}} \sum_{i=1}^n z_{ik} w_{ik} \sum_{\mathbf{j} \in \mathbb{Z}^p} v_{ik}(\mathbf{j})(\mathbf{y}_i + 2\pi \mathbf{j} - \boldsymbol{\mu}_k)(\mathbf{y}_i + 2\pi \mathbf{j} - \boldsymbol{\mu}_k)^\top$$

Synthetic Examples



Empirical Example I

Protein Data (Nodehi et al. 2021)



- A protein domain from the Structural Classification of Proteins (SCOP) database

- Dihedral angles, $(\phi, \psi)^T$, determine how the protein chain bends locally

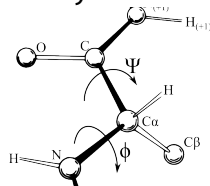


Figure 2: Adapted from: Dcrjsr, Protein backbone dihedral angles, Wikimedia Commons, CC BY 3.0.

Empirical Example II

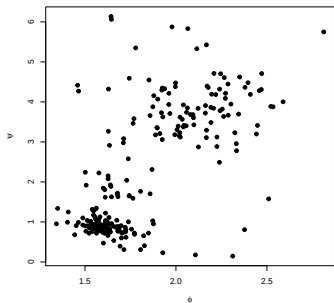


Figure 3: Ramachandran plot of (ϕ, ψ) . $n = 244$ observations: residues (selected positions along the protein chain)

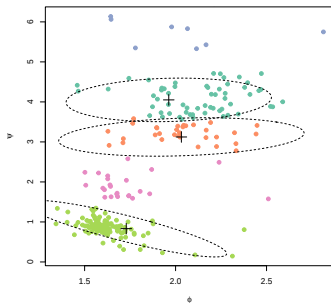


Figure 4: Hierarchical clustering (angular separation distance)

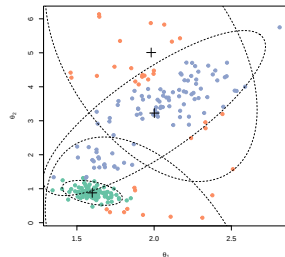


Figure 5: Robust estimates of μ_k, Σ_k for $k = 1, 2, 3$ through an E-CEM.

Thanks for your attention!

Giulia Bertagnolli

giulia.bertagnolli@unibz.it

[gbertagnolli.github.io](https://github.com/gbertagnolli)

References I

- Agostinelli, C. 2007. "Robust Estimation for Circular Data." *Computational Statistics and Data Analysis* 51 (12): 5867–75.
- Agostinelli, C., L. Greco, and G. Saraceno. 2024. "Weighted Likelihood Methods for Robust Fitting of Wrapped Models for p-Torus Data." *AStA Advances in Statistical Analysis* 108: 853–88. <https://doi.org/10.1007/s10182-024-00494-2>.
- Cressie, N., and T. R. C. Read. 1984. "Multinomial Goodness-of-Fit Tests." *Journal of the Royal Statistical Society Series B: Statistical Methodology* 46 (3): 440–64. <https://doi.org/10.1111/j.2517-6161.1984.tb01318.x>.
- Greco, L., P. L. Novi Inverardi, and C. Agostinelli. 2023. "Finite Mixtures of Multivariate Wrapped Normal Distributions for Model Based Clustering of p-Torus Data." *Journal of Computational and Graphical Statistics* 32 (3): 1215–28. <https://doi.org/10.1080/10618600.2022.2128808>.

References II

- Kuchibhotla, A. K., and A. Basu. 2018. “A Minimum Distance Weighted Likelihood Method of Estimation.” Interdisciplinary Statistical Research Unit (ISRU), Indian Statistical Institute, Kolkata, India. <https://faculty.wharton.upenn.edu/wp-content/uploads/2018/02/attemptv4p1.pdf>.
- Lindsay, B. G. 1994. “Efficiency Versus Robustness: The Case for Minimum Hellinger Distance and Related Methods.” *The Annals of Statistics* 22: 1018–114.
- Markatou, M., A. Basu, and B. G. Lindsay. 1998. “Weighted Likelihood Equations with Bootstrap Root Search.” *Journal of the American Statistical Association* 93 (442): 740–50.
- Nodehi, A., M. Golalizadeh, M. Maadooliat, and C. Agostinelli. 2021. “Estimation of Parameters in Multivariate Wrapped Models for Data on Ap-Torus.” *Computational Statistics* 36: 193–215.

References III

Park, C., A. Basu, and B. G. Lindsay. 2002. "The Residual Adjustment Function and Weighted Likelihood: A Graphical Interpretation of Robustness of Minimum Disparity Estimators." *Computational Statistics & Data Analysis* 39 (1): 21–33.
[https://doi.org/10.1016/S0167-9473\(01\)00047-0](https://doi.org/10.1016/S0167-9473(01)00047-0).

Appendix

Appendix I

RAFs

We choose a RAF stemming from the Generalised Kullback–Leibler (GKL) divergence

$$A(\delta, \tau) = \frac{\log(\tau\delta + 1)}{\tau}, \quad 0 \leq \tau \leq 1$$

But one can choose a RAF corresponding to Symmetric Chi-squared divergence (SCHI) (Markatou, Basu, and Lindsay 1998) or the class of Power Divergences (PWD) (Cressie and Read 1984; Park, Basu, and G. Lindsay 2002).

Appendix II

Approximated MLEs and Classification EM algorithm

An approximate MLE can be obtained using crispy assignments after the computation of $v_i(\mathbf{j})$, that is we let

$$\hat{\mathbf{j}}_i = \arg \max_{\mathbf{j} \in \mathbb{Z}^p} v_i(\mathbf{j})$$

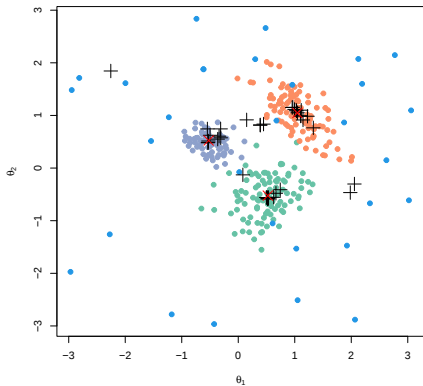
and solve the estimating equation

$$\sum_{i=1}^n u(\hat{\mathbf{x}}_i; \boldsymbol{\theta}) = \mathbf{0} \tag{2}$$

based on the **unwrapped** (fitted) linear data $\hat{\mathbf{x}}_i = \mathbf{y}_i + 2\pi\hat{\mathbf{j}}_i$.

Appendix III

Inizialisation



1. Hierarchical clustering on circular distances (angular separation) allowing for “more” groups
2. Subsampling inside of each *initial component*
3. For each subsample of cluster k we obtain MLE of μ_k and Σ_k
4. For each component k we select the root $(\hat{\mu}_k, \hat{\Sigma}_k)$ minimising the GLK divergence between the data distribution and the WN distribution with parameters $(\hat{\mu}_k, \hat{\Sigma}_k)$.