

# Note on the John Transform for Object Classification

Jonas Bresch\*

Technische Universität Berlin

based on Inspiration during the IPMS at Malta,

*inspired by the works with M. Quellmalz, R. Beinert, M. Beckmann*

## Abstract

The John transform, or also called X-ray transform, is a useful tool for analyzing data in imaging beside the closely related Radon transform. We prove a closed-form expression for the John transform of an axis-aligned rectangular voxel and use the relation of John and Radon transform via one- and multi-stage projection to introduce the John-CDT (J-CDT) and its normalized variant (NJ-CDT). The latter has theoretically provable invariance properties for translated and isotropically scaled measures in arbitrary dimensions. Numerically, a multi-slicing strategy that iteratively applies John projections before CDT produces compact, invariant features that are highly effective in limited-data regimes. Comparisons to related Radon-CDT embeddings, LOT, and momentum-based methods demonstrate competitive or superior accuracy with simple classifiers.

## 1 Introduction

We study a family of integral transforms and associated feature maps for object and image classification based on one- and multi-stage projections. Classical tomography uses the Radon transform to convert a function or a measure on  $\mathbb{R}^d$  into its integrals over hyperplanes; this transform and its computational and stability properties are well understood and form the mathematical background of many imaging modalities [16]. For certain low-data or geometric applications it is natural to replace hyperplane integrals by lower-dimensional projections: the John (or X-ray) transform integrates along parallel lines to a given direction and records the resulting measures on the orthogonal complement.

In this work we exploit the interplay between the Radon and John transforms to build stable, low-dimensional representations that are tailored to classification tasks with limited training data. Furthermore, we derive a closed form of the John transform of a hypercube (Thm. 2), similar to [5: Thm. 1] (cf. Thm. 1), for the numerical realizations. The approach is motivated by recent work on *Radon Cumulative Distribution Transforms* (R-CDT) and their normalized variants. Those are demonstrating

---

\*[Webpage](#) and [Email](#)

17th July 2026

invariance and competitive empirical performance for limited data image classification [2–4]. Here the study of affine measure classes in the sense that each class originates from one (unique) template measure under affine transformations is considered. The latter are crucial, for instance, in the domains of isotropic position of convex bodies [11; 14; 17], since every full-dimensional log-concave probability measure has a unique isotropic representative, or for John and Löwner positions [1; 10; 12], in statistics and machine learning for *Independent Component Analysis* (ICA) [6; 9], in optimal transport and measure geometry by the standard correspondence in the Wasserstein-2 distance [18; 19] or filigranology, the study of historic watermarks [8; 13].

Our central idea is to use iterative John projections to reduce the ambient dimension and then apply one-dimensional Radon and CDT techniques on the projected measures. This “multi-slicing” strategy yields a *John-CDT* (J-CDT) and a *Normalized J-CDT* (NJ-CDT) that (i) admit closed-form expressions for simple templates (rectangular voxels), (ii) are equivariant under translations and isotropic scalings, and (iii) produce compact features that are effective in limited-data classification experiments. The analytical relations between affine changes of variables and the induced action on John and Radon transforms that underlie these properties are elementary but, to our knowledge, have not been exploited systematically for design of classification features. We therefore connect classical transform theory [16] with recent optimal-transport-inspired feature maps [2–4; 15; 20] and with moment-based descriptors used in shape recognition [7].

## 2 Radon and John transform

Let  $\mu \in \mathcal{M}(\mathbb{R}^d)$  be a signed, regular, finite measure. We define the slicing map

$$S_\theta : \mathbb{R}^d \rightarrow \mathbb{R}, \quad x \mapsto \langle x, \theta \rangle \quad \text{in direction} \quad \theta \in \mathbb{S}^{d-1} := \{x \in \mathbb{R}^d : \|x\| = 1\}$$

with the Euclidean norm  $\|\cdot\|$  on  $\mathbb{R}^d$  induced by the standard inner product  $\langle \cdot, \cdot \rangle$  and the restricted Radon transform

$$\mathcal{R}_\theta : \mathcal{M}(\mathbb{R}^d) \rightarrow \mathcal{M}(\mathbb{R}), \quad \mu \mapsto (S_\theta)_\# \mu =: (\mathcal{B}(\mathbb{R}) \ni A \rightarrow \mu(S_\theta^{-1}(A))). \quad (1)$$

Here  $\mathcal{B}$  denotes the Borel set. By definition, the Radon transform is mass preserving, i.e. it holds  $\mathcal{R}_\theta[\mu](\mathbb{R}) = \mu(\mathbb{R}^d)$  for any direction  $\theta \in \mathbb{S}^{d-1}$ . The Radon transform for  $\mu \in \mathcal{M}(\mathbb{R}^d)$  may be defined by integrating along all directions  $\theta \in \mathbb{S}^{d-1}$ . We define the gluing operator

$$\mathcal{I}_\mathcal{R} : \mathbb{S}^{d-1} \times \mathbb{R}^d \rightarrow \mathbb{S}^{d-1} \times \mathbb{R}, \quad (\theta, x) \mapsto (\theta, S_\theta(x))$$

and the Radon transform by

$$\mathcal{R} : \mathcal{M}(\mathbb{R}^d) \rightarrow \mathcal{M}(\mathbb{S}^{d-1} \times \mathbb{R}), \quad \mu \mapsto (\mathcal{I}_\mathcal{R})_\# [u_{\mathbb{S}^{d-1}} \times \mu] =: (u_{\mathbb{S}^{d-1}} \times \mu) \circ (\text{id}, S_\theta^{-1}),$$

where  $u_{\mathbb{S}^{d-1}}$  denotes the uniform measure on  $\mathbb{S}^{d-1}$  and  $\text{id} : X \rightarrow X$  the identity on arbitrary spaces  $X$ .

*Remark 1* (Radon transform for functions). The definition in (1) is equivalent to the definition of the Radon transform of functions  $f \in L^1(\mathbb{R}^d)$  given by the integral restricted to the hyperplane, i.e.

$$\mathcal{R}_\theta[f] : \mathbb{R} \rightarrow \mathbb{R}, \quad t \mapsto \int_{H_\theta(t)} f(x) \, d\mathcal{H}^{d-1}(x),$$

and

$$H_\theta(t) := \{x \in \mathbb{R}^d : \langle x, \theta \rangle = t\} = S_\theta^{-1}(t)$$

the  $(d-1)$ -dimensional hyperplane, where  $\mathcal{H}^{d-1}$  denotes the Hausdorff measure on  $H_\theta(t)$ .  $\circ$

The John transform uses the orthogonal projection

$$\pi_\theta : \mathbb{R}^d \mapsto \theta^\perp, \quad x \mapsto x - S_\theta(x)\theta = (I_d - \theta\theta^\top)x,$$

where

$$\theta^\perp := \{x \in \mathbb{R}^d : \langle x, \theta \rangle = 0\} = S_\theta^{-1}(0)$$

denotes the  $(d-1)$ -dimensional space orthogonal to the direction  $\theta \in \mathbb{S}^{d-1}$ . The restricted John transform, or also called (restricted) X-ray transform, of a measure  $\mu \in \mathcal{M}(\mathbb{R}^d)$  is defined by

$$\mathcal{X}_\theta : \mathcal{M}(\mathbb{R}^d) \rightarrow \mathcal{M}(\theta^\perp), \quad \mu \mapsto (\pi_\theta)_\# \mu =: (\mathcal{B}(\theta^\perp) \ni A \mapsto \mu(\pi_\theta^{-1}(A))). \quad (2)$$

Similar to the Radon transform (1), the John transform (2) is mass preserving, i.e. it holds  $\mathcal{X}_\theta[\mu](\theta^\perp) = \mu(\mathbb{R}^d)$ . In analogy to the definition of the Radon transform, we introduce for the John transform the gluing map

$$\mathcal{I}_\mathcal{X} : \mathbb{S}^{d-1} \times \mathbb{R}^d \rightarrow \mathbb{S}^{d-1} \times \mathbb{R}^d, \quad (\theta, x) \mapsto (\theta, \pi_\theta(x)).$$

Note that the projection is by definition in the multi-dimensional Euclidean space  $\mathbb{R}^d$ ; but for each direction  $\theta \in \mathbb{S}^{d-1}$  it can be restricted to  $\theta^\perp$  to avoid overparameterization. The John transform is defined by

$$\mathcal{X} : \mathcal{M}(\mathbb{R}^d) \rightarrow \mathcal{M}(\mathbb{S}^{d-1} \times^* \mathbb{R}^d), \quad \mu \mapsto (\mathcal{I}_\mathcal{X})_\# (u_{\mathbb{S}^{d-1}} \times \mu) =: (u_{\mathbb{S}^{d-1}} \times \mu) \circ (\text{id}, \pi_\theta^{-1}),$$

where  $\mathbb{S}^{d-1} \times^* \mathbb{R}^d := \{(\theta, x) \in \mathbb{S}^{d-1} \times \mathbb{R}^d : x \in \theta^\perp\}$  reducing the overparametrization; we have  $\mathbb{S}^{d-1} \times^* \mathbb{R}^d \simeq \mathbb{S}^{d-1} \times \mathbb{R}^{d-1}$ .

*Remark 2* (John transform for functions). The definition in (2) is equivalent to the definition for functions  $f \in L^1(\mathbb{R}^d)$  given by the integral restricted to  $\text{span}(\theta) := \{t \cdot \theta \mid t \in \mathbb{R}\}$ , cf. [16],  $\mathcal{X}_\theta[f] : \theta^\perp \rightarrow \mathbb{R}, x \mapsto \int_{\mathbb{R}} f(x + t\theta) dt$ .  $\circ$

*Remark 3* (Correspondence of Radon and John transform). If  $d = 2$ , the hyperplane  $H_\theta(t)$  is given by the line  $t\theta \oplus \theta^\perp$  and hence the John transform, modulo the overparametrization discussed in Rem. 2 equals the Radon transform.  $\circ$

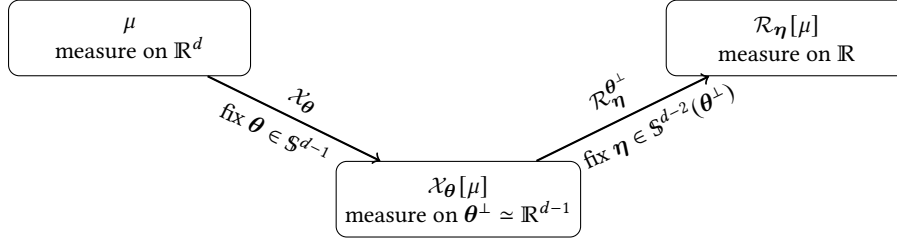
We can represent the  $d$  dimensional Radon transform by applying first the John transform followed by a  $(d-1)$  dimensional Radon transform, cf. Fig. 1. Denote the Radon transform on an  $\ell$ -dimensional vector space  $X \subset \mathbb{R}^d$  by

$$\mathcal{R}_\theta^X : \mathcal{M}(X) \rightarrow \mathcal{M}(\mathbb{R}), \quad \theta \in \mathbb{S}^{\ell-1}(X) := \{\eta \in X : \|\eta\| = 1\}.$$

**Lemma 1.** Let  $\mu \in \mathcal{M}(\mathbb{R}^d)$  and  $\theta \in \mathbb{S}^{d-1}$ . Then, it holds  $\mathcal{R}_\eta^{\theta^\perp}[\mathcal{X}_\theta[\mu]] = \mathcal{R}_\eta[\mu]$  for any  $\eta \in \mathbb{S}^{d-2}(\theta^\perp)$ .

*Proof.* Let  $\eta \in \mathbb{S}^{d-2}(\theta^\perp) \subset \mathbb{R}^d$  and  $B \in \mathcal{B}(\mathbb{R})$ . By definition, we obtain

$$\begin{aligned} \mathcal{R}_\eta^{\theta^\perp}[\mathcal{X}_\theta[\mu]](B) &= \mathcal{X}_\theta[\mu](\{y \in \theta^\perp : \langle \eta, y \rangle \in B\}) \\ &= \mu(\pi_\theta^{-1}(\{y \in \theta^\perp : \langle \eta, y \rangle \in B\})). \end{aligned}$$



**Figure 1:** Visualization of the procedure to capture the (restricted) Radon transform for measure on  $\mathbb{R}^d$  via Lem. 1.

Now, for any  $x \in \mathbb{R}^d$  we have  $\langle \pi_\theta(x), \eta \rangle = \langle (x - \langle x, \theta \rangle \theta), \eta \rangle = \langle x, \eta \rangle$ , since  $\langle \eta, \theta \rangle = 0$ . Hence, we obtain

$$\begin{aligned} \pi_\theta^{-1}(\{y \in \theta^\perp : \langle \eta, y \rangle \in B\}) &= \{x \in \mathbb{R}^d : \pi_\theta(x) \in \theta^\perp, \langle \pi_\theta(x), \eta \rangle \in B\} \\ &= \{x \in \mathbb{R}^d : \langle x, \eta \rangle \in B\} \end{aligned}$$

such that

$$\mathcal{R}_\eta^{\theta^\perp}[\chi_\theta[\mu]](B) = \mu(\{x \in \mathbb{R}^d : S_\eta(x) \in B\}) = \mathcal{R}_\eta[\mu](B), \quad \forall B \in \mathcal{B}(\mathbb{R}). \quad \square$$

**Corollary 1.** Let  $\mu \in \mathcal{M}(\mathbb{R}^d)$  and  $\{\theta_0, \dots, \theta_{d-1}\} \subset \mathbb{S}^{d-1}$  be an orthonormal basis of  $\mathbb{R}^d$ . Define  $\mathcal{X}_{\theta_1}^{(1)}[\mu] := \chi_{\theta_1}[\mu]$ , and iteratively  $\mathcal{X}_{\theta_k}^{(k)}[\mu] := \chi_{\theta_k}[\mathcal{X}_{\theta_{k-1}}^{(k-1)}[\mu]]$  for  $k \in \{2, \dots, d-1\}$ . Then

- i)  $\mathcal{X}_{\theta_k}^{(k)}[\mu]$  is a measure on  $\text{span}(\theta_1, \dots, \theta_k)^\perp$ ,
- ii)  $\mathcal{X}_{\theta_{d-1}}^{(d-1)}[\mu] = \mathcal{R}_{\theta_0}[\mu]$ .

*Proof.* Apply Lem. 1 iteratively and utilize Rem. 3.  $\square$

**Transformations of uniform measure** For  $a \in \mathbb{R}^d$ , let  $u_a$  be the Lebesgue measure  $\lambda^d$  restricted to the  $d$ -dimensional cube  $(-a, a] := \times_{i=1}^d (-a_i, a_i]$ , i.e. we have  $u_a = \chi_a \lambda^d$  for some density function  $\chi_a|_{(-a, a]} = 1$  and zero elsewhere. Furthermore, we denote by  $x \odot y$  the componentwise product in  $\mathbb{R}^d$ , by  $P : \mathbb{R}^d \rightarrow \mathbb{R}$ ,  $x \mapsto \prod_{i=1}^d x_i$  the product of all entries of a vector, and by  $x^\circ := x|_{\text{supp}(x)} \in \mathbb{R}^\ell$  the vector whose  $i$ -th component is the  $i$ -th non-zero component of  $x$ , where  $\ell = \|x\|_0 := |\text{supp}(x)|$  denotes the  $\ell_0$  pseudo norm. In the following, we provide closed forms for the Radon and John transform of a rectangular voxel. This view-point is based on Rem. 1 and Rem. 2 and is a basic first step for numerical computations afterwards. Notably, the following is the density function of  $\mathcal{R}_\theta[u_a]$  w.r.t. the Lebesgue measure  $\lambda^1$ .

**Theorem 1** (Radon transform [5: Thm. 1]). Let  $\theta \in \mathbb{S}^{d-1}$ ,  $t \in \mathbb{R}$ , and  $a \in \mathbb{R}_{>0}^d$ . Furthermore, let  $\ell := \|\theta\|_0$ . It holds that

$$\mathcal{R}_\theta[\chi_a](t) = \frac{2^{d-\ell} P(a)}{P((a \odot \theta)^\circ) (\ell-1)!} \sum_{k \in \{-1, 1\}^\ell} P(k) (t + \langle k, (a \odot \theta)^\circ \rangle_+)^{\ell-1}, \quad t \in \mathbb{R}. \quad (3)$$

The following is the density function of  $\chi_\theta[u_a]$  w.r.t. the Lebesgue Measure  $\lambda^{d-1}$  i.e.  $\lambda^d|_{\theta^\perp}$  restricted to  $\theta^\perp \simeq \mathbb{R}^{d-1}$ .

**Theorem 2** (John transform). *Let  $\theta \in \mathbb{S}^{d-1}$ ,  $x \in \theta^\perp \subset \mathbb{R}^d$ , and  $a \in \mathbb{R}_{\geq 0}^d$ . It holds that*

$$\mathcal{X}_\theta[\chi_a](x) = \sum_{k \in \{-a, a\}} P(\text{sign}(k)) \max \left\{ 0, \min_{i: \theta_i < 0} \left\{ \frac{k_i - x_i}{\theta_i} \right\} - \max_{i: \theta_i > 0} \left\{ \frac{k_i - x_i}{\theta_i} \right\} \right\}.$$

*Proof.* Denote by  $H(x)$  the Heavyside function, which is  $H(x) = 0$  for all  $x \leq 0$  and  $H(x) = 1$  for all  $x > 0$ . The indicator function satisfies the identity

$$\chi_a(x) = H(x + a) - H(x - a) \quad \text{for } x \in \mathbb{R} \quad \text{and} \quad a > 0.$$

Therefore, it is

$$\chi_a(x) = \prod_{i=1}^d \chi_{a_i}(x_i) = \prod_{i=1}^d H(x + a) - H(x - a) = \sum_{k \in \{-a, a\}^d} (-1)^{|k|} \prod_{i=1}^d H(x_i + k_i).$$

By definition, we obtain

$$\mathcal{X}_\theta[\chi_a](x) = \int_{\mathbb{R}} \chi_a(x + t\theta) \, dt = \sum_{k \in \{-a, a\}^d} (-1)^{|k|} \int_{\mathbb{R}} \prod_{i=1}^d H(x_i + t\theta_i + k_i) \, dt.$$

Note that the product of  $H(\cdot + k_i)$  defines the area, after substituting  $x \mapsto x + t\theta$ , where

$$\begin{aligned} x_i + t\theta_i \geq k_i &\Leftrightarrow \begin{cases} t \geq \frac{k_i - x_i}{\theta_i} & : \theta_i > 0, \\ t \leq \frac{k_i - x_i}{\theta_i} & : \theta_i < 0 \end{cases} \\ &\Leftrightarrow \max_{i: \theta_i > 0} \left\{ \frac{k_i - x_i}{\theta_i} \right\} \leq t \leq \min_{i: \theta_i < 0} \left\{ \frac{k_i - x_i}{\theta_i} \right\}, \quad 1 \leq i \leq d. \quad \square \end{aligned}$$

*Remark 4.* More generally, i.e. for the rectangular cube  $(a, b] := \times_{i=1}^d (a_i, b_i]$  where  $a < b$  componentwise, holds

$$\mathcal{X}_\theta[\chi_{a,b}](x) = \max\{0, t^\uparrow(x)_a^b - t^\downarrow(x)_a^b\},$$

where  $\chi_{a,b}|_{(a,b]} := 1$ , and zero elsewhere with

$$t^\uparrow(x)_a^b := \min_{i \in \{1, \dots, d\}} t_i^\uparrow(x)_a^b \quad \text{and} \quad t^\downarrow(x)_a^b := \max_{i \in \{1, \dots, d\}} t_i^\downarrow(x)_a^b,$$

and for any  $i \in \{1, \dots, d\}$

$$t_i^\uparrow(x)_a^b := \begin{cases} \max \left\{ \frac{a_i - x_i}{\theta_i}, \frac{b_i - x_i}{\theta_i} \right\} & \text{if } \theta_i \neq 0 \\ +\infty, & \text{if } \left\{ \theta_i = 0 \text{ and } x_i \in [a_i, b_i) \right\} \\ -\infty, & \text{else,} \end{cases}$$

and

$$t_i^\downarrow(x)_a^b := \begin{cases} \min \left\{ \frac{a_i - x_i}{\theta_i}, \frac{b_i - x_i}{\theta_i} \right\} & \text{if } \theta_i \neq 0 \\ -\infty & \text{if } \left\{ \theta_i = 0 \text{ and } x_i \in [a_i, b_i) \right\} \\ +\infty & \text{else.} \end{cases} \quad \circ$$

**Affine measure classes** In the manner of [2-4], we define an affinely transformed measure  $\mu_{A,y}$  in terms of a template measure  $\mu \in \mathcal{M}(\mathbb{R}^d)$ , a regular matrix  $A \in \text{GL}(d)$  and bias  $y \in \mathbb{R}^d$  by

$$\mu_{A,y} := (A \cdot + y)_\# \mu = \mu(A^{-1}(\cdot - y)). \quad (4)$$

It is well known [3: Prop. 3] that the Radon transform of an affinely transformed measure  $\mu_{A,y}$  has a close correspondence to a Radon transform of the underlying template measure  $\mu$ :

$$\mathcal{R}_\theta[\mu_{A,y}] = (\|A^\top \theta\| \cdot + \langle y, \theta \rangle)_\# \mathcal{R}_{\frac{A^\top \theta}{\|A^\top \theta\|}}[\mu] = \mathcal{R}_{\frac{A^\top \theta}{\|A^\top \theta\|}}[\mu] \circ \left( \frac{\cdot - \langle y, \theta \rangle}{\|A^\top \theta\|} \right). \quad (5)$$

The following similar correspondence holds for the John transform.

**Lemma 2.** *Let  $\theta \in \mathbb{S}^{d-1}$ ,  $A \in \text{GL}(d)$ ,  $y \in \mathbb{R}^d$  and define  $\theta_A := A^{-1}\theta/\|A^{-1}\theta\| \in \mathbb{S}^{d-1}$ . Furthermore, let  $M_\theta: \theta_A^\perp \rightarrow \theta^\perp$ ,  $x \mapsto \pi_\theta(A\pi_{\theta_A}(x))$ . Then the John transform satisfies  $\mathcal{X}_\theta[\mu_{A,y}] = (z \mapsto M_\theta(z) + \pi_\theta(y))_\# \mathcal{X}_{\theta_A}[\mu]$ .*

*Proof.* We consider the direction  $\theta_A$  and observe that  $\pi_\theta Ax = 0$  if and only if  $\pi_{\theta_A} x = 0$ . This means that  $\ker(\pi_\theta \circ (x \mapsto Ax)) = \ker(\pi_{\theta_A})$ . By definition and linearity, we have for all  $x \in \mathbb{R}^d$  that

$$\begin{aligned} M_\theta(x) &= \pi_\theta(A\pi_{\theta_A}(x)) = \pi_\theta(A(x - \langle x, \theta_A \rangle \theta_A)) = \pi_\theta(Ax) - \langle x, \theta_A \rangle \pi_\theta(A\theta_A) \\ &= \pi_\theta(Ax) - \langle x, \theta_A \rangle \frac{1}{\|A^{-1}\theta\|} \pi_\theta(\theta) = \pi_\theta(Ax), \end{aligned}$$

and hence  $M_\theta = \pi_\theta \circ (z \mapsto Az)$ . Then, we have by the linearity of  $\pi_\theta$

$$\begin{aligned} \mathcal{X}_\theta[\mu_{A,y}] &= (\pi_\theta)_\#[(z \mapsto Az + y)_\# \mu] \\ &= (z \mapsto \pi_\theta(Az + y))_\# \mu \\ &= (z \mapsto z + \pi_\theta(y))_\#[(z \mapsto \pi_\theta(Az))_\# \mu] \end{aligned}$$

and hence

$$\mathcal{X}_\theta[\mu_{A,y}] = (z \mapsto z + \pi_\theta(y))_\# \underbrace{(M_\theta)_\# (\pi_{\theta_A})_\# \mu}_{=\mathcal{X}_{\theta_A}[\mu]} = (z \mapsto M_\theta(z) + \pi_\theta(y))_\# \mathcal{X}_{\theta_A}[\mu]. \quad \square$$

**Corollary 2** (Special case for  $A \in \text{O}(d)$ ). *If  $A \in \text{O}(d)$ , which means that  $AA^\top = A^\top A = I_d$ , then it holds*

$$\mathcal{X}_\theta[\mu_{A,y}] = (A \cdot + \pi_\theta(y))_\# \mathcal{X}_{A^\top \theta}[\mu] = \mathcal{X}_{A^\top \theta}[\mu] \circ \left( A^\top (\cdot - \pi_\theta(y)) \right).$$

*Proof.* It holds

$$\begin{aligned} \mathcal{X}_\theta[\mu_{A,y}] &= (\pi_\theta)_\#[(A \cdot + y)_\# \mu] = (\pi_\theta(A \cdot + y))_\# \mu = (\pi_\theta(A \cdot) + \pi_\theta(y))_\# \mu \\ &= (A\pi_{A^\top \theta} \cdot + \pi_\theta(y))_\# \mu = (A \cdot + \pi_\theta(y))_\# \mathcal{X}_{A^\top \theta}[\mu], \end{aligned}$$

using that  $\|A^\top \theta\| = \|\theta\| = 1$  and

$$\begin{aligned} \pi_\theta(Ax) &= (I_d + \theta\theta^\top)Ax = Ax + \underbrace{AA^\top}_{=I_d} \theta(A^\top \theta)^\top x \\ &= A(I_d - A^\top \theta(A^\top \theta)^\top)x = A\pi_{A^\top \theta}(x) \end{aligned}$$

for any  $x \in \mathbb{R}^d$ .  $\square$

**Multi slicing for multi-dimensional data** Motivated by Cor. 1, we aim to combine for applications the John transform twice, i.e. first the John transform with respect to a direction  $\theta \in \mathbb{S}^{d-1}$  for reducing the dimensions of the directions and afterwards applying the Radon transform on the resulting measure on the orthogonal hyperplane  $\text{span}\{\theta\}^\perp \simeq \mathbb{R}^{d-1}$ , which has a closed form and is much easier to evaluate than the 3D Radon transform.

**Corollary 3** (3D John transform and Projecting Twice). *Let  $\{\theta_0, \theta_1, \theta_2\} \subset \mathbb{S}^2$  be an orthonormal basis of  $\mathbb{R}^3$ . Then, for any  $A \in \text{GL}(3)$  and  $y \in \mathbb{R}^3$ , it holds*

$$\mathcal{X}_{\theta_2}[\mathcal{X}_{\theta_1}[\mu_{A,y}]] = \mathcal{X}_{\eta_2}[\mathcal{X}_{\eta_1}[\mu]] \circ \left( \frac{\cdot - \langle y, \theta_0 \rangle}{\|A^\top \theta_0\|} \right),$$

where

$$\eta_0 := \frac{A^\top \theta_0}{A^\top \theta_0}, \quad \eta_1 := \frac{\pi_{\eta_0}(\theta_1)}{\|\pi_{\eta_0}(\theta_1)\|}, \quad \text{and} \quad \eta_2 := \eta_0 \times \eta_1.$$

*Proof.* Since  $x = \sum_{j=0}^2 \theta_j \theta_j^\top x = \sum_{j=0}^2 S_{\theta_j}(x) \theta_j$ , we have

$$(\pi_{\theta_1} \circ \pi_{\theta_2})(x) = (I_d - \sum_{j=1}^2 \theta_j \theta_j^\top) x = S_{\theta_0}(x) \theta_0, \quad \forall x \in \mathbb{R}^d.$$

The assertion follows by applying Cor. 1 and (5).  $\square$

In the following, we denote  $\mathcal{X}_\Theta[\mu] := \mathcal{X}_{\theta_{d-1}}[\dots[\mathcal{X}_{\theta_1}[\mu]]]$  for the iterative application of the John transform from Cor. 1.

**Corollary 4.** *Let  $A \in \{c \cdot I_d : c > 0\}$  and  $y \in \mathbb{R}^d$ , then  $\mathcal{X}_\Theta[\mu_{A,y}] = (\|A\theta_0\| \cdot + \langle \theta_0, y \rangle)_\# \mathcal{X}_\Theta$  on  $\Theta^\perp = \text{span}(\theta_0)$ , where  $\{\theta_0\} \cup \Theta \subset \mathbb{S}^{d-1}$  is an orthonormal basis of  $\mathbb{R}^d$ .*

*Proof.* Utilize Lem. 2, with  $\theta_A = \theta$ , the idempotency of  $\pi_{\theta_i}$ , and the decompositions from the proof of Cor. 3.  $\square$

### 3 Classification in John and Radon CDT spaces

Due to the multi slicing procedure and Cor. 3 and 4, the analogue to the Radon Cumulative Distribution Transform (R-CDT) [3], say John Cumulative Distribution Transform (J-CDT), can be defined via the *cumulative distribution transform* (CDT)  $\hat{\nu} := F_\nu^{[-1]}$ , with the *cumulative distribution function*  $F_\nu := \nu((-\infty, \cdot])$  for  $\nu \in \mathcal{P}(\mathbb{R})$ , and the *generalized inverse*  $g^{-1} := \inf\{s : g(s) > \cdot\}$  applied on  $\mathcal{X}_\Theta[\mu] \in \mathcal{P}(\mathbb{R})$  where  $\mu \in \mathcal{P}(\mathbb{R}^d)$  by Cor. 1. Notably, it holds

$$\hat{\nu} = \arg \max_{T_\# u_{[0,1]} = \nu} \int_0^1 c(s - T(s)) \, ds, \quad \text{for some convex cost function } c : \mathbb{R} \rightarrow [0, \infty),$$

where  $T : \mathbb{R} \rightarrow \mathbb{R}$  is measurable and  $u_{[0,1]}$  denotes the uniform measure on  $[0, 1]$ . By [4], we have for  $\mu \in \mathcal{P}_c^*(\mathbb{R}^d)$ , where

$$\begin{aligned} \mathcal{P}_c^*(\mathbb{R}^d) &:= \{\mu \in \mathcal{P}(\mathbb{R}^d) \mid \dim(\text{conv}(\text{supp}(\mu))) > d - 1 \\ &\quad \wedge \text{conv}(\text{supp}(\mu)) \subset \mathbb{R}^d\} \subset \mathcal{P}_2(\mathbb{R}^2), \end{aligned}$$

and  $\subset$  denotes a compact subset, that normalizing yields features, say *Normalized John-CDT* (NJ-CDT), which are invariant under isotropic scaling as well as shifting given by

$$\mathcal{J}_\Theta : \mathcal{P}(\mathbb{R}^d) \rightarrow L^\infty([0, 1]), \quad \mu \mapsto (\widehat{\mathcal{X}_\Theta[\mu]} - \text{mean}(\widehat{\mathcal{X}_\Theta[\mu]})) \text{std}(\widehat{\mathcal{X}_\Theta[\mu]})^{-1},$$

where  $\text{mean}(g) := \int_0^1 g(s) \, ds$  and  $\text{std}(g)^2 := \int_0^1 (g(s) - \text{mean}(g))^2 \, ds$  for  $g \in L^\infty([0, 1])$ .

**Proposition 1.** Let  $\theta_0 \in \mathbb{S}^{d-1}$  and  $\Theta \subset \mathbb{S}^{d-1}$  such that  $\{\theta_0\} \cup \Theta$  is an orthonormal basis of  $\mathbb{R}^d$ . Further, let  $A \in \{\text{diag}(c \cdot \mathbf{1}) : c \in \mathbb{R} \setminus \{0\}\}$  and  $y \in \mathbb{R}^d$ . Then, it holds  $\mathcal{J}_\Theta[\mu] = \mathcal{J}_\Theta[(A \cdot + y)_\# \mu]$ .

*Proof.* The well-definiteness of the NJ-CDT is given by the continuity of  $\theta \rightarrow \text{mean}(\widehat{\mathcal{X}_\Theta[\mu]})$  and  $\theta \rightarrow \text{std}(\widehat{\mathcal{X}_\Theta[\mu]})$ , cf. [3: Lem. 1]. Furthermore, for  $\mu \in \mathcal{P}_c^*(\mathbb{R}^d)$  it is  $\text{std}(\widehat{\mathcal{X}_\Theta[\mu]}) \geq c$  for some  $c > 0$  for any  $\Theta \subset \mathbb{S}^{d-1}$  with  $|\Theta| = d - 1$ . Translating Lem. 2, i.e. Cor. 4, into the CDT space, we obtain

$$\widehat{\mathcal{X}_\Theta[\mu_{A,y}]}(t) = \|A\theta_0\| \widehat{\mathcal{X}_\Theta[\mu]}(t) + \|\prod_{i=1}^{d-1} \pi_{\theta_i} y\| = c_A \widehat{\mathcal{X}_\Theta[\mu]}(t) + \langle \theta_0, y \rangle, \quad t \in \mathbb{R},$$

where  $A = c_A I_d$ , such that

$$\text{mean}(\widehat{\mathcal{X}_\Theta[\mu_{A,y}]}) = \|A\theta_0\| \text{mean}(\widehat{\mathcal{X}_\Theta[\mu]}) + \langle \theta_0, y \rangle,$$

and

$$\text{std}(\widehat{\mathcal{X}_\Theta[\mu_{A,y}]}) = \|A\theta_0\| \text{std}(\widehat{\mathcal{X}_\Theta[\mu]}) + \langle \theta_0, y \rangle,$$

and hence  $\mathcal{J}_\Theta[\mu_{A,y}] = \mathcal{J}_\Theta[\mu]$ , which yields the assertion.  $\square$

We directly obtain the (linear) separability of arbitrary finitely many measure classes of the form

$$\mathbb{C}_i := \{(c \cdot I_d + y)_\# \mu_i \mid c \in \mathbb{R} \setminus \{0\}, y \in \mathbb{R}^d, \mu_i \in \mathcal{P}_c^*(\mathbb{R}^d), \text{ with } \mathcal{J}_\Theta[\mu_i] \neq \mathcal{J}_\Theta[\mu_j],$$

for  $i \neq j$ .

*Remark 5* (Normalized Radon-CDT). For arbitrary  $A \in \text{GL}(d)$  and  $y \in \mathbb{R}^d$ , the so-called *Normalized Radon-CDT* (NR-CDT) [2] provides features

$$\mathcal{N}_{\theta_0}[(A \cdot + y)_\# \mu] = \mathcal{J}_\Theta[(A \cdot + y)_\# \mu] = \mathcal{N}_{A\theta_0/\|A\theta_0\|}[\mu]$$

and resorting of the angles—by taking the maximum—yields the assertion [2; 3] as

$$\mathcal{N}_m[\mu](t) := \sup_{\theta \in \mathbb{S}^d} \mathcal{N}_\theta[\mu](t), \quad t \in \mathbb{R},$$

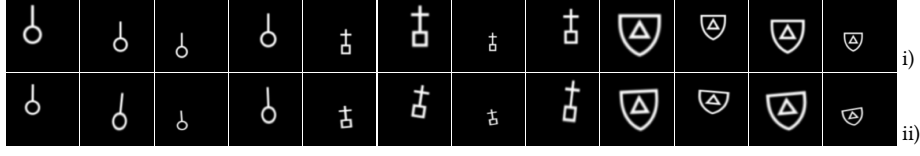
such that  $\mathcal{N}_m[(A \cdot + y)_\# \mu] = \mathcal{N}_m[\mu]$ .  $\circ$

## 4 Numerical results

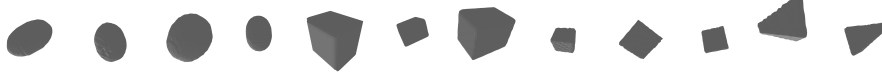
We start with simple classification techniques for data recognition in limited data regimes. For numerical evidence of the theoretical results we consider a 2D and 3D dataset of template images and objects which become translated and isotropically scaled to fulfill the conditions of Prop. 1, similar to [15; 20]. We further add rotations and anisotropic scaling violating the theoretical assumptions to study the robustness of the feature extractors.

We implemented the multidimensional discrete John transform and the experiments in Julia.<sup>1</sup> They are performed on a MacBookPro 2020 with an Intel Core i5 CPU (4 cores, 1.4 GHz) and 8 GB RAM.

<sup>1</sup>The code is available at <https://github.com/JJEWBresch/3dxNRCDT>.



**Figure 2:** The 2D academic dataset. i): Idealized setting; ii): Non-idealized setting (with additional rotating and anisotropic scaling).



**Figure 3:** The 3D academic dataset. Non-idealized setting (with additional rotating and anisotropic scaling).

**Image classification** We consider a dataset of three template images of size  $128 \times 128$  which get randomly isotropically scaled within a range of  $[0.5, 1.5]$  and translated by pixel positions within the range  $[-20, 20]$ , see for visualization Fig. 2. Afterwards the John transform is applied with a random angle  $\theta_0 \in \mathbb{S}^1$ , uniformly chosen from the circle, with a uniform grid of size 128 in the tangent space  $\mathbb{R} \leftrightarrow (\theta_0)^\perp \subset \mathbb{R}^2$  inside the box (centered image) of (pixel)size  $[-64, 64]^2$ . The generalized inverse is calculated by applying a linear interpolation on a uniform grid of size 128.

The resulting features  $\mathcal{J}_{\theta_0^\perp}$  can be compared: here, we, firstly, use the underlying three template measures as the reference to calculate the nearest template for the NJ-CDT transformations of each sample, see Tab. 1 (second columns, respectively); secondly, we only use the NJ-CDT transformations of the samples and perform a nearest neighbor calculation, with an random split of  $\{1, 3, 5\}$  samples out of 10 for each class for training, against  $\{9, 7, 5\}$ , respectively, for testing, and repeating it 20 times. The means and standard deviations are reported in Tab. 1 (third to fives columns, respectively, top). For both examples, we utilize different norms in the feature space for calculating the distances. The same experiments are performed on the non-idealized setting, i.e. Fig. 2 ii). As the assumptions in Prop. 1 are similar to the once on the *linear OT* (LOT) [15; 20], a comparison with our proposed feature extractor is provided in Tab. 1 (bottom).

We observe perfect classification results on the idealized dataset, and fairly good accuracies (up to 93%) in the non-idealized setting substantiating significant robustness. In both settings, the LOT embedding becomes clearly outperformed.

**Object recognition** In the second part, we consider the task of 3D shape recognition. Here, the experiments consists of one academic dataset based three synthetic

**Table 1:** Top: NT and 1-NN calculation on  $\mathcal{J}_{\theta_0^\perp}$  embedding, where  $\theta_0 \in \mathbb{S}^1$  is randomly sampled, uniformly. Bottom: NT and 1-NN calculation on LOT embedding [15; 20].

| dataset            | Fig. 2 i) |               |               |               | Fig. 2 ii)    |                      |                      |                      |
|--------------------|-----------|---------------|---------------|---------------|---------------|----------------------|----------------------|----------------------|
| method             | NT        | 1-NN (10%)    | 1-NN (33%)    | 1-NN (50%)    | NT            | 1-NN (10%)           | 1-NN (33%)           | 1-NN (50%)           |
| $\ \cdot\ _1$      | 1.0000    | 1.0000±0.0000 | 1.0000±0.0000 | 1.0000±0.0000 | 0.8333        | 0.7241±0.1246        | 0.8881±0.0864        | 0.9434±0.0817        |
| $\ \cdot\ _2$      | 1.0000    | 1.0000±0.0000 | 1.0000±0.0000 | 1.0000±0.0000 | 0.8333        | 0.7074±0.1529        | 0.8905±0.0726        | 0.9100±0.0845        |
| $\ \cdot\ _\infty$ | 1.0000    | 1.0000±0.0000 | 1.0000±0.0000 | 1.0000±0.0000 | <b>0.9333</b> | <b>0.8259±0.1147</b> | <b>0.9190±0.0466</b> | <b>0.9434±0.0622</b> |
| $\ \cdot\ _1$      | 0.5660    | 0.5324±0.1390 | 0.7380±0.0780 | 0.8130±0.0932 | 0.5000        | 0.5314±0.1122        | 0.5833±0.0858        | 0.6266±0.0687        |
| $\ \cdot\ _2$      | 0.5377    | 0.5574±0.1198 | 0.7119±0.1064 | 0.7966±0.1091 | 0.5333        | 0.5351±0.1268        | 0.5833±0.0999        | 0.6495±0.1121        |
| $\ \cdot\ _\infty$ | 0.5000    | 0.4570±0.0949 | 0.4328±0.0983 | 0.5566±0.1150 | 0.4666        | 0.4370±0.0896        | 0.4309±0.0920        | 0.5300±0.0929        |

**Table 2:** Left: 1-NN calculation for the non-idealized dataset on  $\mathcal{J}_{(\theta_1, \theta_2)^+}$ , where  $\theta_1 \in \mathbb{S}^2$  and  $\theta_2 \in \mathbb{S}^1(\theta_1^\perp)$  is randomly sampled. Middle: NT and 1-NN calculation on  $\mathcal{N}_m$  calculated by  $\mathcal{R}_\Theta[\mathcal{X}_{\theta_1}]$  where  $\Theta = \mathbb{S}^1(\theta_1^\perp)$  and  $\theta_1 \in \mathbb{S}^2$  is randomly chosen. Right: NT and 1-NN calculation with momentum based features on  $\mathcal{X}_{\theta_1^\perp}$ .

| method             | 1-NN (20%)           | 1-NN (40%)           | 1-NN (60%)           | 1-NN (20%)           | 1-NN (40%)           | 1-NN (60%)           | 1-NN (20%)    | 1-NN (40%)    | 1-NN (60%)    |
|--------------------|----------------------|----------------------|----------------------|----------------------|----------------------|----------------------|---------------|---------------|---------------|
| $\ \cdot\ _1$      | 0.6500±0.1420        | 0.7222±0.2555        | 0.8333±0.1529        | 0.9541±0.0425        | 0.9722±0.0493        | 0.9666±0.0418        | 0.4916±0.1642 | 0.6333±0.1616 | 0.6333±0.1999 |
| $\ \cdot\ _2$      | 0.6375±0.1124        | 0.8222±0.1503        | 0.8750±0.1417        | <b>0.9208±0.0632</b> | <b>0.9833±0.0407</b> | <b>1.0000±0.0000</b> | 0.5208±0.1452 | 0.6388±0.1833 | 0.7083±0.1863 |
| $\ \cdot\ _\infty$ | <b>0.7875±0.0875</b> | <b>0.8500±0.1263</b> | <b>0.9000±0.0837</b> | 0.9208±0.0875        | 0.9722±0.0493        | 0.9916±0.0372        | 0.5125±0.1123 | 0.5055±0.1549 | 0.7983±0.2529 |

template objects (sphere, cube, pyramid). First it is generated in an idealized way, i.e. we randomly scaled the object, isotropically, within the range of  $[0.5, 1.5]$  and translated in all three directions within the range of  $[-20, 20]$ . Additionally, we consider rotations within  $[-180^\circ, 180^\circ]$  for all Eulerian angles, and add anisotropic scaling effect within the range of  $[0.75, 1.25]$ .

For classification, we use, as in the previous paragraph, 1-NN calculation, firstly, on the NJ-CDT, via the John transform w.r.t. the angle pairs  $(\theta_1, \theta_2) \in \mathbb{S}^2 \times \mathbb{S}^1(\theta_1^\perp)$ , which corresponds to  $\mathcal{R}_{\theta_0}$  in the manner of Cor. 1, and accuracy results in Tab. 2 (left block). Secondly, we perform the same techniques on the  $m$ NR-CDT via the coupling from Cor. 1; the accuracy results are reported in Tab. 2 (middle block) Furthermore, we combine the John transform  $\mathcal{X}_{\theta_1}$  with the momentum-based method from [7]; the accuracy results are reported in Tab. 2 (right block).

Theoretically, for the moment-based feature extractors on the measures  $\mathcal{X}_{\theta_1}[\mu] \in \mathcal{P}_c^*(\mathbb{R}^2)$  the setting would be idealized, at least for the idealized dataset, since isotropic scaling and shifting result into affine transformations of the measures after applying the John transform once; however, we observe accuracies up to 71%, and the other proposed combinations yield perfect classification results in all settings. In the non-idealized setting, cf. Fig. 3, slight uncertainties can be observed, cf. Tab. 2, which can be significantly reduced via the NN calculation with an appropriate rate of test samples.

## References

- [1] BALL, K., “An elementary introduction to modern convex geometry”, in S. Levy (ed.), *Flavors of Geometry*, xxxi (MSRI Publications; Cambridge University Press, 1997), 1–58.
- [2] BECKMANN, M., BEINERT, R., and BRESCH, J., *Generalizations of the Normalized Radon Cumulative Distribution Transform for Limited Data Recognition*, arXiv:2512.08099, 2025. doi: [10.48550/arXiv.2512.08099](https://doi.org/10.48550/arXiv.2512.08099).
- [3] — “Max-Normalized Radon Cumulative Distribution Transform for Limited Data Classification”, in *International Conference on Scale Space and Variational Methods in Computer Vision (SSVM)* (Springer, 2025), 241–54. doi: [10.1007/978-3-031-92366-1\\_19](https://doi.org/10.1007/978-3-031-92366-1_19).
- [4] — “Normalized Radon Cumulative Distribution Transforms for Invariance and Robustness in Optimal Transport Based Image Classification”, *SIAM Journal on Imaging Sciences*, 19/2 (2026), 677–709. doi: [10.1137/25M1767741](https://doi.org/10.1137/25M1767741).
- [5] BEINERT, R., BRESCH, J., and QUELLMALZ, M., *A Discrete Radon Transform Based on the Area of Cube-Plane Intersection*, arXiv:2603.13503, 2026. doi: [10.48550/arXiv.2603.13503](https://doi.org/10.48550/arXiv.2603.13503).

- [6] COMON, P., “Independent component analysis, a new concept?”, *Signal Processing*, 36 (1994), 287–314.
- [7] FLUSSER, J., and SUK, T., “Pattern recognition by affine moment invariants”, *Pattern Recognition*, 26/1 (1993), 167–74. doi: [10.1016/0031-3203\(93\)90098-H](https://doi.org/10.1016/0031-3203(93)90098-H).
- [8] HAUSER, D., BECKMANN, M., KOLIANDER, G., and STIEHL, H. S., “On Image Processing and Pattern Recognition for Thermograms of Watermarks in Manuscripts – A First Proof-of-Concept”, in *International Conference on Document Analysis and Recognition (ICDAR)* (Springer, 2024), 91–107. doi: [10.1007/978-3-031-70543-4\\_6](https://doi.org/10.1007/978-3-031-70543-4_6).
- [9] HYVÄRINEN, A., KARHUNEN, J., and OJA, E., *Independent Component Analysis* (Wiley, 2001).
- [10] JOHN, F., “Extremum problems with inequalities as subsidiary conditions”, in *Studies and Essays Presented to R. Courant on His 60th Birthday* (Interscience, 1948), 187–204.
- [11] KLARTAG, B., “On convex perturbations with a bounded isotropic constant”, *Geom. Funct. Anal.* 16 (2006), 1274–90.
- [12] LÖWNER, K., “Über monotone Matrixfunktionen”, *Mathematische Zeitschrift*, 38 (1934), 177–216.
- [13] MAZZOLA, R., and MOSCHINI, U., “Automatic watermark recognition and classification”, *Digital Scholarship in the Humanities* (2012).
- [14] MILMAN, V. D., and PAJOR, A., “Isotropic position and inertia ellipsoids and zonoids of the unit ball of a normed  $n$ -dimensional space”, in J. Lindenstrauss and V. Milman (eds.), *GAGA Seminar 1987–89*, mcccclxxvi (Lecture Notes in Mathematics; Springer, 2006), 64–104.
- [15] MOOSMÜLLER, C., and CLONINGER, A., “Linear optimal transport embedding: provable Wasserstein classification for certain rigid transformations and perturbations”, *Information and Inference: A Journal of the IMA* 12/1 (2023), 363–89. doi: [10.1093/imaiai/iaac023](https://doi.org/10.1093/imaiai/iaac023).
- [16] NATTERER, F., and WÜBBELING, F., *Mathematical Methods in Image Reconstruction* (Philadelphia, PA: SIAM, 2001). doi: [10.1137/1.9780898718324](https://doi.org/10.1137/1.9780898718324).
- [17] VERSHYNIN, R., *High-Dimensional Probability* (Cambridge University Press, 2018).
- [18] VILLANI, C., *Optimal Transport: Old and New* (Springer, 2009).
- [19] ———, *Topics in Optimal Transportation* (American Mathematical Society, 2003).
- [20] WANG, W., SLEPČEV, D., BASU, S., OZOLEK, J. A., and ROHDE, G. K., “A linear optimal transportation framework for quantifying and visualizing variations in sets of images”, *International Journal of Computer Vision*, 101/2 (2013), 254–69. doi: [10.1007/s11263-012-0566-z](https://doi.org/10.1007/s11263-012-0566-z).