



ML-EM in continuum: sparse measure solutions

Camille Pouchol, Olivier Verdier

► To cite this version:

Camille Pouchol, Olivier Verdier. ML-EM in continuum: sparse measure solutions. 2019. hal-02192480v1

HAL Id: hal-02192480

<https://inria.hal.science/hal-02192480v1>

Preprint submitted on 23 Jul 2019 (v1), last revised 12 Aug 2019 (v2)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

ML-EM in continuum: sparse measure solutions

Camille Pouchol*, Olivier Verdier*[†]

2019-07-23

Linear inverse problems $A\mu = y$ with Poisson noise and non-negative unknown $\mu \geq 0$ are ubiquitous in applications where light is involved, such as Positron Emission Tomography (PET) in medical imaging. The associated maximum likelihood problem is most often solved by the celebrated iterative ML-EM algorithm, which is known to yield results which look spiky even when stopped early. This work provides an explanation for this phenomenon by going into continuum in image space when the operator A is integral, such as in PET. This is done by considering the image μ as a measure rather than an element of a Lebesgue space. We prove that if the measurements y are not in the cone $\{A\mu, \mu \geq 0\}$, which is typical of high noise or short exposure times, likelihood maximisers as well as ML-EM cluster points must be sparse, i.e., typically a sum of Dirac masses. In the better case where y is in the cone, we prove that cluster points of ML-EM will be measures without singular part, as long as the initial guess is smooth. Finally, we provide concentration bounds for the probability to be in the undesirable sparse case, proving that it is exponentially small with exposure time, with a rate controlled by the distance of the noiseless data to the boundary of the cone.

1 Introduction

In various imaging modalities, recovering the image from acquired data can be recast as solving an inverse problem of the form $A\mu = y$, where A is a linear operator, y represents noisy measurements and μ the image, with $\mu \geq 0$ usually a desirable property. The problem thus becomes $\min_{\mu \geq 0} d(y, A\mu)$ where d is some given distance or divergence.

When the model is finite-dimensional, the operator A is simply a matrix $A = (a_{ij}) \in \mathbb{R}^{d \times r}$. If we assume a Poisson noise model, i.e., $y_i \sim \mathcal{P}((A\mu)_i)$ with independent draws,

*Department of Mathematics, KTH Royal Institute of Technology, 100 44 Stockholm, Sweden.

[†]Department of Computing, Mathematics and Physics, Western Norway University of Applied Sciences, Bergen, Norway.

the corresponding (negative log) likelihood problem is equivalent to

$$\min_{\mu \geq 0} d(y||A\mu), \quad (1)$$

where d is the Kullback–Leibler divergence. As it turns out, this statistical model is similar to the familiar non-negative least-squares regression corresponding to Gaussian noise, but for a different distance functional: if $y \notin \{A\mu, \mu \geq 0\}$, it is projected onto it in the sense of d , whereas if it belongs to this image set, any $\mu \geq 0$ such that $A\mu = y$ will be optimal.

The celebrated Maximum Likelihood Expectation Maximisation algorithm (ML-EM) precisely aims at solving (1) and was introduced by Shepp and Vardi [29, 31], in the particular context of the imaging modality called Positron Emission Tomography (PET). ML-EM is iterative and writes

$$\mu_{k+1} = \frac{\mu_k}{A^T 1} A^T \left(\frac{y}{A\mu_k} \right), \quad (2)$$

starting from $\mu_0 > 0$, usually $\mu_0 = 1$.

This algorithm is an expectation-maximisation (EM) algorithm, and as such it has many desirable properties: it preserves non-negativity and the negative log-likelihood decreases along iterates [9]. It can also be interpreted in several other ways [31, 8, 3], see [22] for an overview. The expectation-maximisation point of view has also led to alternative algorithms [10], but in spite of various competing approaches, ML-EM (actually, its more numerically efficient variation OSEM [13]) has remained the algorithm used in practice in many PET scanners.

Despite its success, the ML-EM algorithm (2) is known to produce undesirable spikes along the iterates, where some pixels take increasingly high values. The toy example presented in Figure 1 is an example of such artefacts in the case of PET. The reconstruction of a torus of maximum 1 with 100 iterations of ML-EM indeed exhibits some pixels with values as high as about 6.

This phenomenon has long been noticed in the literature, and has been called “spiky” or “speckled” images [30], or “the chequerboard effect” [31]. To the best of our knowledge, a theoretical explanation for this result has however remained elusive.

The aim of the present paper is to better understand that phenomenon via the analysis *in a continuous setting* of the minimisation problem (1) and the corresponding ML-EM algorithm (2). The continuous setting here refers to the image not being discretised on a grid. Note however that we keep the data space discrete.

Informally, considering μ as an element in some function space, we consider forward operators A of the form

$$(A\mu)_i := \int_K a_i(x)\mu(x) dx,$$

where K is the compact on which one aims at reconstructing the image, and a_i is some non-negative function on K . This covers a wide range of applications, including PET.

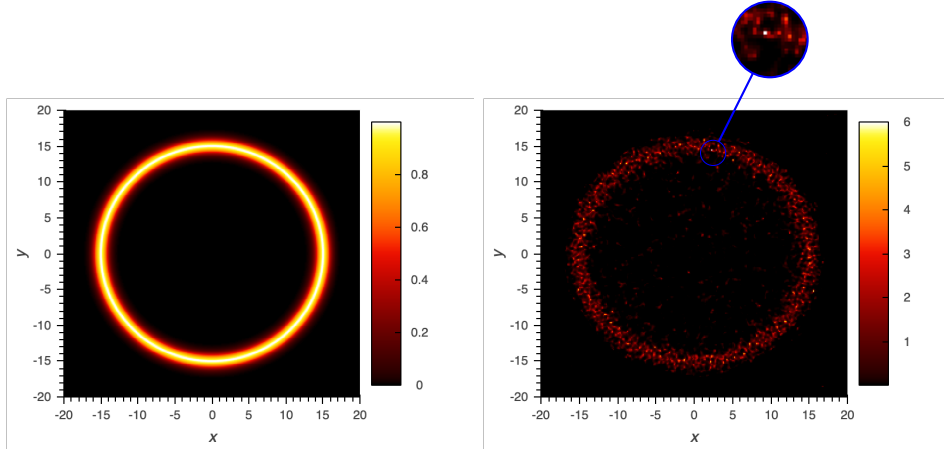


Figure 1: Phantom and reconstruction after 100 iterations of ML-EM, with a zoom on the region containing the pixel of highest value.

One of our motivations is to derive algorithms for Poisson inverse problems with movement, for example for PET acquisition of a moving organ [15]. In that case, movement can be modelled by deformations of images which do not easily carry over to discretised images (simply because interesting deformations do not preserve the grid). It is then desirable to express the problem in a continuous setting, in order to both analyse the algorithms proposed in the literature, and to derive new ones [12, 23].

The field of inverse problems for imaging, with a continuum description of the unknown image, is abundant [5, 2]. Most often, the image is taken to be a function in some appropriate Sobolev space. To the best of our knowledge, however, there are relatively few results concerning the continuum description of the Poisson likelihood and the ML-EM algorithm for solving it.

In [20, 21, 19] and [26], both the image and data are considered in continuum, with a deterministic description of noise. These authors assume that detectors a_i lie in $L^\infty(K)$ and correspondingly assume that the image μ lies in $L^1(K)$. They study the convergence properties of the corresponding ML-EM algorithm in detail. In the first series of three papers, the compact K is restricted to $K = [0, 1]$.

Our paper differs from these works in that we do not make the two following restrictive assumptions, common to [20, 21, 19, 26]. The first restriction is to assume the existence of a non-negative solution μ to the equation $A\mu = y$, assumed to lie in $L^1(K)$. The second restriction is to assume that the functions a_i are bounded away from zero. This last assumption is unrealistic for some applications such as PET [26, Remark 6.1].

The seminal paper [17] considers the optimisation problem over the set of non-negative Borel measures as we do. They obtain the corresponding likelihood function informally as the limit of the discrete one, but do not prove that it is an actual maximum likelihood problem for the PET statistical model. They then proceed to study the problem of whether minimisers can be identified with bounded functions, and not merely measures which might have a singular part. They indeed note that in some very specific cases (see

also [19]), one can prove that the minimiser should be a Dirac mass. They speculate that there might be a link with the usual spiky results from ML-EM. They, however, do not provide any general conditions for sparsity.

Working in the space of non-negative measures $\mathcal{M}_+$, our main contributions are as follows:

Continuous Framework. We prove that the continuous setting of measures is precisely the maximum likelihood problem with a Poisson point process model (Proposition 2.1), and that the natural generalisation of the ML-EM iterates (17) indeed corresponds to the expectation-maximisation method associated to that continuous statistical model (see § 2.2).

Sparsity. We give a precise sparsity criterion (sparsity means that any optimal solution has singular support): if the data y is outside the cone $A(\mathcal{M}_+)$, then all optimal solutions are necessarily sparse (Corollary 3.6); if the data y is inside the cone $A(\mathcal{M}_+)$, then there exist absolutely continuous solutions (Lemma 3.10)

Properties of ML-EM iterates. We show the expected properties of the ML-EM iterates, namely monotonicity (Proposition 4.1) and optimality of cluster points (Proposition 4.2).

Properties of ML-EM solutions. We show that in the non-sparse case, i.e., when an absolutely continuous solution exists as just mentioned, the solution picked up by ML-EM is absolutely continuous (Theorem 4.7), and we give an explicit example of ML-EM converging to a sum of point masses in the sparse case (Proposition 4.3).

Effect of Noise. We derive estimates on the probability to be in the sparse case, depending on the noise level (Proposition 5.1, Theorem 5.2).

"Spiky" artefacts. With these results, we provide an explanation for the artefacts of Figure 1: they are related to the sparsity result. By weak duality, we can indeed *certify* that the iterates should converge to a sum of point masses in that case, as detailed in § 6 dedicated to simulations.

Outline of the paper The paper is organised as follows. In § 2, we introduce the functional and ML-EM in continuum in detail, with all the necessary notations, normalisations, definitions and useful properties about Kullback–Leibler divergences. § 3 contains all results on the functional minimisers, starting from the optimality conditions to the diverging cases of the data y being inside or outside the cone $A(\mathcal{M}_+)$. § 4 is devoted to the algorithm ML-EM itself, with the proof of its usual properties in continuum together with the implications they have on the case where the data y is in the cone $A(\mathcal{M}_+)$. In § 5, we estimate the probability that the data y ends up outside the image cone $A(\mathcal{M}_+)$. In § 6, we present simulations which confirm our theoretical predictions. Finally, in § 7 we conclude with open questions and perspectives.

2 Maximum likelihood and ML-EM in continuum

2.1 Mathematical background

Space of Radon measures. As stated in the introduction, we model the image to reconstruct as a non-negative measure μ defined on a compact set K . Some of our results require $K \subset \mathbb{R}^p$ (typically, $p = 2$ or 3).

More precisely, we will consider the set of Radon measures, denoted $\mathcal{M}(K)$ and defined as the topological dual of the set of continuous functions over K , denoted $\mathcal{C}(K)$. The space of non-negative measures will be denoted by $\mathcal{M}_+(K)$. For brevity, we will often write $\mathcal{M}$ for $\mathcal{M}(K)$ and $\mathcal{M}_+$ for $\mathcal{M}_+(K)$ when there is no ambiguity as to the underlying compact K .

We identify a linear functional $\mu \in \mathcal{M}_+$ with its corresponding Borel measure (as per the Riesz–Markov representation Theorem), using $\mu(B)$ to denote the measure of a Borel subset B of K . We will also sometimes write the dual pairing between a measure $\mu \in \mathcal{M}$ and a function $f \in \mathcal{C}(K)$ as

$$\langle \mu, f \rangle = \int_K f d\mu.$$

The support of a measure $\mu \in \mathcal{M}$ is defined as the closed set

$$\text{supp}(\mu) := \{x \in K, \mu(N) > 0, \forall N \in \mathcal{N}(x)\},$$

where $\mathcal{N}(x)$ is the set of all open neighbourhoods of x .

Finally, recall that, by the Banach–Alaoglu Theorem, bounded sets in $\mathcal{M}$ are weak-* compact [27].

Kullback–Leibler divergence. We here recall the definition of the Kullback–Leibler divergence. Instead of giving the general definition, we directly explicit the two instances that will actually be needed in this paper, for (non-normalised) non-negative vectors in $\mathbb{R}^d$, and for probability measures on K .

- *For vectors in $\mathbb{R}^d$.* For any two non-negative vectors u and v in $\mathbb{R}^d$, we define the Kullback–Leibler divergence between u and v as

$$d(u||v) := \sum_{i=1}^d \left(v_i - u_i - u_i \log \left(\frac{v_i}{u_i} \right) \right),$$

with the convention $0 \log(0) = 0$ and $d(u||v) = +\infty$ if there exists i such that $u_i = 0$ and $v_i > 0$.

- *For probability measures on K .* For any two probability measures μ and ν on K , we define the Kullback–Leibler divergence between μ and ν

$$D(\mu||\nu) := - \int_K \log \left(\frac{d\nu}{d\mu} \right) d\mu,$$

if ν is absolutely continuous with respect to μ and $\log\left(\frac{d\nu}{d\mu}\right)$ is integrable with respect to μ . Here $\frac{d\nu}{d\mu}$ stands for the Radon–Nikodym derivative of ν with respect to μ . Otherwise, we define $D(\mu||\nu) := +\infty$.

When a measure is absolutely continuous with respect to the Lebesgue measure on $K \subset \mathbb{R}^p$, we simply say that it is absolutely continuous. Any reference to the Lebesgue measure implicitly assumes that K stands for the closure of some bounded domain in $\mathbb{R}^p$ (i.e., a bounded, connected and open subset of $\mathbb{R}^p$).

2.2 Statistical model

We want to recover a measure $\mu \in \mathcal{M}_+$ from independent Poisson distributed measurements

$$N_i \sim \mathcal{P}\left(\int_K a_i d\mu\right), \quad i = 1, \dots, d, \quad (3)$$

with

$$a_i \geq 0, \quad a_i \in \mathcal{C}(K), \quad i = 1, \dots, d.$$

Positron Emission Tomography [24]. In PET, a radiotracer injected into the patient and, once concentrated into tissues, disintegrates by emitting a positron. This process is well known to be accurately modelled by a Poisson point process, itself defined by a non-negative measure. After a short travel distance called *positron range*, this positron interacts with an electron. The result is the emission of two photons in random opposite directions. Pairs of detectors around the body then detect simultaneous photons, and the data is given by the number of counts per pair of detectors.

In the case of PET, d is the number of detectors (i.e., pairs of single detectors). For a given point $x \in K$ and detector $i \in \{1, \dots, d\}$, $a_i(x)$ then denotes the probability that a positron emitted in x will be detected by detector i .

Finally, we will throughout the paper assume

$$\sum_{i=1}^d a_i > 0 \text{ on } K. \quad (4)$$

For PET, this amounts to assuming that the points in K are in the so-called field of view, namely that any emission has a non-zero probability to be detected.

Derivation of the statistical model (3) for PET. We proceed to give a proof that the statistical model (3) (and thus, the corresponding likelihood function) applies to PET. Here, we assume that the emission process of PET is modelled by a Poisson point process, defined by a measure $\mu \in \mathcal{M}_+$, and that each point drawn from the Poisson process has a probability $a_i(x)$ to be detected by detector i .

Proposition 2.1. *The statistical model (3) applies to PET.*

Proof. The proof relies on the following properties of Poisson point processes [16]:

- *law of numbers*: the number of points emitted by a Poisson process of intensity μ follows the Poisson law with parameter $\int_K d\mu = \mu(K)$.
- *thinning property*: the points that are kept with (measurable) probability $p : K \mapsto [0, 1]$ still form a Poisson point process, with intensity $p\mu$, and it is independent from that of points that are not kept. This property generalises to p_i , $1 \leq i \leq d$ with $\sum_{i=1}^d p_i(x) = 1$ for all $x \in K$.

By the thinning property, the families of points which lead to an emission detected in detector i , $i = 1, \dots, d$, are all independent Poisson processes with associated measure $a_i\mu$, for $i = 1, \dots, d$. Thus, the random variables N_i representing the number of points detected in detector i are independent and of law $\mathcal{P}(\int_K a_i d\mu)$, which proves the claim. $\square$

Maximum likelihood problem. The likelihood corresponding to the statistical model (3) reads

$$L(N_1, \dots, N_p; \mu) = \prod_{i=1}^d L(N_i; \mu) = \prod_{i=1}^d e^{-\int_K a_i d\mu} \frac{(-\int_K a_i d\mu)^{N_i}}{N_i!},$$

since $\mathbb{P}(N_i = n_i) = e^{-\int_K a_i d\mu} \frac{(-\int_K a_i d\mu)^{n_i}}{n_i!}$.

Dropping the factorial terms (they do not depend on μ and will thus play no role when maximising the likelihood), we get

$$\log(L(N_1, \dots, N_p; \mu)) = - \sum_{i=1}^d \int_K a_i d\mu + \sum_{i=1}^d N_i \log \left(\int_K a_i d\mu \right).$$

The corresponding maximum likelihood problem, written for a realisation n_i of the random variable N_i , $i = 1, \dots, d$, is given by

$$\max_{\mu \in \mathcal{M}_+} - \int_K \left(\sum_{i=1}^d a_i \right) d\mu + \sum_{i=1}^d n_i \log \left(\int_K a_i d\mu \right). \quad (5)$$

Defining the operator

$$\begin{aligned} A: \mathcal{M} &\longrightarrow \mathbb{R}^d \\ \mu &\longmapsto (\langle \mu, a_i \rangle)_{1 \leq i \leq d}, \end{aligned}$$

the optimisation problem conveniently rewrites in terms of the Kullback-Leibler divergence: upon adding constants and taking the negative log-likelihood problem, it reads

$$\min_{\mu \in \mathcal{M}_+} d(n || A\mu).$$

ML-EM iterates. We now define the ML-EM algorithm, which aims at solving the optimisation problem (5). It is given by the iterates

$$\mu_{k+1} = \frac{\mu_k}{\sum_{i=1}^d a_i} \left(\sum_{i=1}^d \frac{n_i a_i}{\int_K a_i d\mu_k} \right),$$

starting from an initial guess $\mu_0 \in \mathcal{M}_+$.

Note that this algorithm can be shown to be an EM algorithm for the continuous problem. The proof is beyond the scope of this paper, so we decide to omit it, but we just mention that the corresponding so-called *complete data* would be given by the positions of points together with the detector that has detected each of them.

2.3 Normalisations

Due to the assumption (4), we may without loss of generality assume that

$$\sum_{i=1}^d a_i = 1 \tag{6}$$

on K . Otherwise we could just define $\tilde{\mu} = (\sum_{i=1}^d a_i)\mu$ and $\tilde{a}_i = a_i/(\sum_{j=1}^d a_j)$. This normalisation now implies $0 \leq a_i \leq 1$ for all $i = 1, \dots, d$.

We further normalise the measures by dividing the functional by $n := \sum_{i=1}^d n_i$, considering $\mu := \frac{\mu}{n}$ to remove the factor. We then define

$$y_i := \frac{n_i}{n}.$$

From now on, we consider the optimisation problem (minimisation of the negative log-likelihood):

$$\min_{\mu \in \mathcal{M}_+} l(\mu), \tag{7}$$

where

$$\begin{aligned} l(\mu) &:= \int_K d\mu - \sum_{i=1}^d y_i \log \left(\int_K a_i d\mu \right) \\ &= \langle \mu, 1 \rangle - \sum_{i=1}^d y_i \log (\langle \mu, a_i \rangle), \end{aligned} \tag{8}$$

defined to be $+\infty$ for any measure such that $\langle \mu, a_i \rangle = 0$ for some $i \in \text{supp}(y)$, where

$$\text{supp}(y) := \{ i = 1, \dots, d \mid y_i > 0 \}.$$

After normalisation, the ML-EM iterates are given by

$$\mu_{k+1} = \mu_k \left(\sum_{i=1}^d \frac{y_i a_i}{\int_K a_i d\mu_k} \right) = \mu_k \left(\sum_{i=1}^d \frac{y_i a_i}{\langle \mu_k, a_i \rangle} \right). \tag{9}$$

We recall the property that ML-EM preserves the total number of counts: $\langle \mu_k, 1 \rangle = 1$ for all $k \geq 1$, which corresponds to $\langle \mu_k, 1 \rangle = n = \sum_{i=1}^d n_i$ before normalisation. We also emphasise the important property that iterations cannot increase the support of the measure, namely

$$\forall k \in \mathbb{N}, \text{supp}(\mu_{k+1}) \subset \text{supp}(\mu_k).$$

ML-EM iterates are well-defined. We assume throughout that the initial measure μ_0 fulfils

$$\langle \mu_0, a_i \rangle > 0 \quad \forall i \in \text{supp}(y). \quad (10)$$

Note that usual practice is to take μ_0 to be absolutely continuous with respect to the Lebesgue measure, typically $\mu_0 = 1$, in which case (10) is satisfied.

The following simple Lemma shows that assumption (10) ensures that the iterates are well-defined.

Lemma 2.2. *The ML-EM iterates (9) satisfy*

$$\langle \mu_k, a_i \rangle > 0 \implies \langle \mu_{k+1}, a_i \rangle > 0 \quad i \in \text{supp}(y).$$

Proof. From the Cauchy–Schwarz inequality, $\mu_k(K) \langle \mu_k, a_i^2 \rangle \geq \langle \mu_k, a_i \rangle^2$. Combined with the definition of ML-EM iterates, this entails for any $i \in \text{supp}(y)$,

$$\langle \mu_{k+1}, a_i \rangle \geq \sum_{i=1}^d y_i \frac{\langle \mu_k, a_i^2 \rangle}{\langle \mu_k, a_i \rangle} \geq y_i \frac{\langle \mu_k, a_i^2 \rangle}{\langle \mu_k, a_i \rangle} \geq y_i \frac{\langle \mu_k, a_i \rangle}{\mu_k(K)} > 0.$$

□

2.4 Adjoint and cone

Since $A: \mathcal{C}(K)^* \rightarrow \mathbb{R}^d$, we can define its adjoint $A^*: \mathbb{R}^d \rightarrow \mathcal{C}(K)$ (identifying $\mathbb{R}^d$ as a Euclidean space with its dual), which is given by

$$A^*w = \sum_{i=1}^d w_i a_i, \quad w \in \mathbb{R}^d.$$

The set $A(\mathcal{M}_+) = \{A\mu, \mu \in \mathcal{M}_+\} \subset \mathbb{R}^d$ is a closed and convex cone and, as proved in [11], its dual cone $A(\mathcal{M}_+)^*$ can be characterised as being given by the set of vectors $\lambda \in \mathbb{R}^d$ such that $\sum_{i=1}^d \lambda_i a_i \geq 0$ on K , i.e.,

$$A(\mathcal{M}_+)^* = \left\{ \lambda \in \mathbb{R}^d \mid A^*\lambda \geq 0 \text{ on } K \right\}. \quad (11)$$

As a result, the interior of the dual cone $A(\mathcal{M}_+)^*$ is given by the vectors $\lambda \in \mathbb{R}^d$ such that $A^*\lambda > 0$ on K .

The normalisation condition (6) can now be rewritten

$$A^*\mathbf{1} = \mathbf{1}, \quad (12)$$

where $\mathbf{1}$ is the vector of $\mathbb{R}^d$ which all components are one: $\mathbf{1} = (1, \dots, 1)$. Moreover, we can rewrite the ML-EM iteration (9) as

$$\mu_{k+1} = \mu_k A^* \left(\frac{y}{A\mu} \right),$$

which is the continuous analogue to the discrete case (2), taking into account the normalisation (12).

Minimisation over the cone. The problem $\min_{\mu \in \mathcal{M}_+} d(y||A\mu)$, is equivalent to the following minimisation problem over the cone $A(\mathcal{M}_+)$:

$$\min_{w \in A(\mathcal{M}_+)} d(y||w). \quad (13)$$

Indeed, if w^* is optimal for the problem (13), any μ^* such that $A\mu^* = w^*$ is optimal for the original problem. From the property $d(y, w) = 0 \Leftrightarrow y = w$, we also infer that when $y \in A(\mathcal{M}_+)$, μ^* is optimal if and only if $A\mu^* = y$.

3 Properties of minimisers

In this section, we gather results concerning the functional l and its minimisers, proving that they are sparse when the data y is not in the image cone $A(\mathcal{M}_+)$. First, we note that the functional l defined by (8) is a convex and proper function.

3.1 Characterisation of optimality

We now derive necessary and sufficient *optimality conditions*. The existence of minimisers turns out to be a byproduct of the analysis of the ML-EM algorithm, whose cluster points will be shown to satisfy the optimality conditions, see [Proposition 4.2](#).

We first prove that any optimum must have a fixed unit mass.

Proposition 3.1. *If μ^* is optimal for (7), then $\langle \mu^*, \mathbf{1} \rangle = \int_K d\mu^* = 1$.*

Proof. For any $\mu \in \mathcal{M}_+$, we have

$$l(\mu) = - \sum_{i=1}^d y_i \log \left(\frac{\langle \mu, a_i \rangle}{\langle \mu, \mathbf{1} \rangle} \right) + (\langle \mu, \mathbf{1} \rangle - \log(\langle \mu, \mathbf{1} \rangle)). \quad (14)$$

Observe that the second term depends only on the mass $\langle \mu, \mathbf{1} \rangle$, whereas the first term is scale-invariant. As a result, an optimal μ has to minimise the second term, which turns out to admit the unique minimiser $\langle \mu, \mathbf{1} \rangle = 1$. $\square$

Remark 3.2. This result follows from the optimality conditions derived later in [Proposition 3.3](#), but the proof above is simple and also highlights that the maximum likelihood estimator for μ is consistent with that for $\int_K d\mu$: the second term in (14) is none other than the negative log-likelihood of the total mass.

We now give the full optimality conditions. The convex function l defined in (8) has the following open domain:

$$\text{dom}(l) := \{\mu \in \mathcal{M}_+, \langle \mu, a_i \rangle > 0 \text{ for } i \in \text{supp}(y)\}.$$

Notice further that for any $\mu \in \text{dom}(l)$, the function l is Fréchet-differentiable (in the sense of the strong topology). Its gradient is given for $\mu \in \text{dom}(l)$ is then the element in the dual $\mathcal{M}^*$ of $\mathcal{M}$ given by

$$\nabla l(\mu) = 1 - \sum_{i \in \text{supp}(y)} \frac{y_i a_i}{\langle \mu, a_i \rangle}, \quad (15)$$

which we identify with an element of $\mathcal{C}(K)$.

For any vector $w \in \mathbb{R}^d$, we define $\lambda(w) \in \mathbb{R}^d$ by

$$\lambda_i(w) := 1 - \frac{y_i}{w_i}, \quad (16)$$

(with the convention $\lambda_i = 1$ if $y_i = 0$, that is, if $i \notin \text{supp}(y)$). Using $\sum_{i=1}^d a_i = 1$, we can rewrite (15) as

$$\nabla l(\mu) = A^* \lambda(A\mu) = \sum_{i \in \text{supp}(y)} \lambda_i(A\mu) a_i. \quad (17)$$

Proposition 3.3. *The measure $\mu^* \in \mathcal{M}$ is optimal for the problem (7) if and only if the following optimality conditions hold*

$$\begin{aligned} A^* \lambda(A\mu^*) &\geq 0 && \text{on } K, \\ A^* \lambda(A\mu^*) &= 0 && \text{on } \text{supp}(\mu^*). \end{aligned} \quad (18)$$

These conditions can be equivalently written as

$$\begin{aligned} \sum_{i \in \text{supp}(y)} \frac{y_i a_i}{\langle \mu^*, a_i \rangle} &\leq 1 && \text{on } K, \\ \sum_{i \in \text{supp}(y)} \frac{y_i a_i}{\langle \mu^*, a_i \rangle} &= 1 && \text{on } \text{supp}(\mu^*). \end{aligned} \quad (19)$$

Proof. Since f is differentiable on $\text{dom}(l)$ and convex on the convex set $\mathcal{M}_+$, a point $\mu^* \in \text{dom}(l)$ is optimal if and only if

$$\nabla l(\mu^*) \in -N_{\mathcal{M}_+}(\mu^*).$$

where

$$N_{\mathcal{M}_+}(\mu) := \{f \in \mathcal{C}(K) \mid \forall \nu \in \mathcal{M}_+, \langle \nu - \mu, f \rangle \leq 0\}$$

is the normal cone to $\mathcal{M}_+$ at μ .

From the characterisation of $N_{\mathcal{M}_+}(\mu)$ given below in Lemma 3.4, and the fact that $\nabla l(\mu) = A^* \lambda(A\mu)$, the optimality condition exactly amounts to the conditions (18). $\square$

Lemma 3.4. *The normal cone at a given $\mu \in \text{dom}(l)$ is given by*

$$N_{\mathcal{M}_+}(\mu) = \{f \in \mathcal{C}(K) \mid f \leq 0 \text{ on } K, f = 0 \text{ on } \text{supp}(\mu)\}.$$

Proof. Let $f \in \mathcal{C}(K)$ be in $N_{\mathcal{M}_+}(\mu)$, i.e., it satisfies $\langle \nu - \mu, f \rangle \leq 0$ for all ν in $\mathcal{M}_+$. First, we choose $\nu = \mu + \delta_x$ which yields $f(x) \leq 0$, so we must have $f \leq 0$ on K . Then with $\nu = 0$, we find $\langle \mu, f \rangle \geq 0$. Since we also have $f \leq 0$, $\langle \mu, f \rangle = 0$ leading to $f = 0$ on $\text{supp}(\mu)$. The reverse also holds true: if $f \leq 0$ on K and $f = 0$ on $\text{supp}(\mu)$, then $\langle \nu - \mu, f \rangle \leq 0$ for all ν in $\mathcal{M}_+$, and consequently f is in $N_{\mathcal{M}_+}(\mu)$. $\square$

Remark 3.5. An alternative proof of these optimality conditions can be obtained by considering instead the equivalent problem of minimising $d(y||w)$ with w ranging over the cone $A(\mathcal{M}_+)$. The cone has a non-empty relative interior which proves that Slater's condition is fulfilled. Since the problem is convex, KKT conditions are equivalent to optimality for $\min_{w \in A(\mathcal{M}_+)} d(y||w)$ [6].

The Lagrange dual is given by $g(\lambda) := \min d(y||w) - \langle \lambda, w \rangle$ for $\lambda \in A(\mathcal{M}_+)^*$. A straightforward computation leads to

$$g(\lambda) = \sum_{i=1}^d y_i \log(1 - \lambda_i), \quad (20)$$

for $\lambda \leq 1$, with value $-\infty$ if there exists $i \in \text{supp}(y)$ such that $\lambda_i = 1$.

The KKT conditions for a primal optimal w^* and dual optimal λ^* write

- (i) $w^* \in A(\mathcal{M}_+)$, $\lambda^* \in A(\mathcal{M}_+)^*$
- (ii) $\langle \lambda^*, w^* \rangle = 0$,
- (iii) $\nabla_w d(y||w^*) - w^* = 0$ (equivalent to $\lambda^* = \lambda(w^*) = 1 - \frac{y}{w^*}$)

A measure μ^* is then optimal if and only if $A\mu^* = w^*$. Since $(\lambda^*, A\mu^*) = \langle \mu^*, A^*\lambda^* \rangle$ (by definition of A^*), the condition (ii) thus becomes

$$\langle \mu^*, A^*\lambda^* \rangle = \int_K A^*\lambda^* d\mu^* = 0.$$

Since $\lambda^* \in A(\mathcal{M}_+)^*$, $A^*\lambda^* \geq 0$ over K . Thus, we must have $A^*\lambda^* = 0$ on $\text{supp}(\mu^*)$ for the above integral to vanish. All in all, we exactly recover the conditions (19), with the additional interpretation that $\lambda(A\mu^*)$ is a dual optimal variable.

3.2 Case $y \notin A(\mathcal{M}_+)$

When the data y is not in the cone $A(\mathcal{M}_+)$, optimality conditions imply sparsity of any optimal measure.

Corollary 3.6. *Assume that $y \notin A(\mathcal{M}_+)$. Then any μ^* minimiser of (7) is sparse, in the following sense*

$$\text{supp}(\mu^*) \subset \arg \min \left(\sum_{i=1}^d \lambda_i(A\mu^*) a_i \right), \quad (21)$$

and $\lambda(A\mu^*) \neq 0$.

Proof. Define $s_i := \langle \mu^*, a_i \rangle - y_i$. Note that $\sum_i s_i = 0$. Moreover, $(\forall i \quad s_i = 0) \implies y \in A(\mathcal{M}_+)$, so we conclude that for some i , $s_i \neq 0$. $\square$

Remark 3.7. Why does condition (21) imply sparsity? Let $\lambda(A\mu^*)$ be defined as (16) for an optimal measure μ^* , and define the function $\varphi := A^* \lambda(A\mu^*) = \sum_{i=1}^d \lambda_i(A\mu^*) a_i$. We know from Proposition 3.3 that both $\varphi \geq 0$ and $\text{supp}(\mu^*) \subset \arg \min(\varphi)$.

Assuming that the a_i 's are linearly independent in $\mathcal{C}(K)$, φ cannot vanish identically since $\lambda(A\mu^*) \neq 0$. Supposing further that for all i , $a_i \in \mathcal{C}^2(K)$, we have

$$\text{supp}(\mu^*) \cap \text{int}(K) \subset \mathcal{S} := \{x \in K \mid \nabla \varphi(x) = 0\}.$$

We make the final assumption that the Hessian of φ is invertible at the points $x \in \arg \min(\varphi)$, which is equivalent to its positive definiteness since these are minimum points of φ . This implies that $\mathcal{S}$ consists of *isolated points*. Consequently, the restriction to $\text{int}(K)$ of any optimal solution μ^* is a sum of Dirac masses.

Note that all the above regularity assumptions hold true for *generic* functions a_i . One case where all of them are readily satisfied is when the functions a_i are analytic with K connected.

Remark 3.8. We can exhibit a case where only Dirac masses are optimal. Suppose that only $y_{i_0} = 1$. Then the function l for measures μ such that $\mu(K) = 1$ is simply $l(\mu) = 1 - \log(\langle \mu, a_{i_0} \rangle)$. One can directly check that a minimiser μ^* necessarily satisfies $\text{supp}(\mu^*) \subset \arg \max(a_{i_0})$, in agreement with condition (21). If this set is discrete, then μ^* is a sum of Dirac masses located at these points. Note that such a data point y is outside the cone $A(\mathcal{M}_+)$ if and only if $\max(a_{i_0}) < 1$. If not, it lies on the boundary of the cone, showing that some boundary points might lead to sparse minimisers as well.

3.3 Moment matching problem and case $y \in \text{int}(A(\mathcal{M}_+))$

When the data y is in the cone $A(\mathcal{M}_+)$, searching for minimisers of (7) is equivalent to solving $A\mu = y$ for $\mu \in \mathcal{M}_+$. For the applications, we are particularly interested in the existence of absolutely continuous solutions. We make use of the results of [11], which addresses this problem.

We shall use the assumption:

$$\text{the functions } a_i, i = 1, \dots, d \text{ are linearly independent in } \mathcal{C}(K). \quad (22)$$

Under (22), $A(\mathcal{M}_+)$ has non-empty interior.

We now recall a part of Theorem 3 of [11] which will be sufficient of our purpose.

Theorem 3.9 ([11]). *Assume that $A(\mathcal{M}_+)$ and its dual cone $A(\mathcal{M}_+)^*$ have non-empty interior. Then for any $y \in \text{int}(A(\mathcal{M}_+))$, there exists μ^* which is absolutely continuous, with positive and continuous density, such that $A\mu^* = y$.*

Lemma 3.10. *Under hypothesis (22), $A(\mathcal{M}_+)$ has non-empty interior, and if $y \in \text{int}(A(\mathcal{M}_+))$, there exists an optimal measure μ^* which is absolutely continuous with positive and continuous density.*

Proof. This is a direct consequence of Theorem 3.9. We just need to check that $A(\mathcal{M}_+)^*$ has non-empty interior. Using the characterisation of the dual cone (11), this is straightforward since $\sum_{i=1}^d a_i = 1$. $\square$

3.4 Case $y \in \partial A(\mathcal{M}_+)$

The previous approach settles the case where the data y is in the interior $\text{int}(A(\mathcal{M}_+))$ of the image cone, which poses the natural question of its boundary $\partial A(\mathcal{M}_+)$. It routinely happens in practice that some components of the data y are zero, which means that the vector y lies at the border of the cone, $y \in \partial A(\mathcal{M}_+)$. Upon changing the compact, a further use of the results of [11] shows that if the support of the data $\text{supp}(y)$ is not too small (see the precise condition (24) below), the situation is the same as for $\text{int}(A(\mathcal{M}_+))$.

The idea is to remove all the zero components of the data vector y , and to consider only the positive ones and try to solve $\langle \mu^*, a_i \rangle = y_i$ for $i \in \text{supp}(y)$ but making sure that the measure μ^* has a support such that $\langle \mu^*, a_i \rangle = 0$ for $i \notin \text{supp}(y)$.

We denote $\tilde{d} := \#(\text{supp}(y))$, $\tilde{K} := K \setminus \cup_{i \notin \text{supp}(y)} a_i^{-1}(\{0\})$, $\tilde{y} = (y_i)_{i \in \text{supp}(y)}$, and finally the reduced operator,

$$\begin{aligned} \tilde{A}: \mathcal{M}(\tilde{K}) &\longrightarrow \mathbb{R}^{\tilde{d}} \\ \mu &\longmapsto \left(\int_K a_i \mu \right)_{i \in \text{supp}(y)}, \end{aligned}$$

which has an associated cone $\tilde{A}(\mathcal{M}_+(\tilde{K}))$.

We will need the assumptions

$$\text{the functions } a_i, i \in \text{supp}(y) \text{ are linearly independent in } \mathcal{C}(\tilde{K}), \quad (23)$$

and

$$\sum_{i \in \text{supp}(y)} a_i > 0 \text{ on } \tilde{K}. \quad (24)$$

Proposition 3.11. *We assume (23) and (24). $\tilde{A}(\mathcal{M}_+(\tilde{K}))$ has non-empty interior and we assume*

$$\tilde{y} \in \text{int}(\tilde{A}(\mathcal{M}_+(\tilde{K}))).$$

Then there exists an absolutely continuous solution μ^ of $A\mu = y$.*

Proof. We make use of [Theorem 3.9](#). In order to do so, we need the dual cone of $\tilde{A}(\mathcal{M}_+(\tilde{K}))$ to not have empty interior, which holds true by (24). Then we may build an absolutely continuous $\tilde{\mu}^* \in \mathcal{M}_+(\tilde{K})$ with positive and continuous density, such that $\tilde{A}\tilde{\mu}^* = \tilde{y}$. We then extend $\tilde{\mu}^*$ to a measure on the whole of K by defining μ^* to equal $\tilde{\mu}^*$ on $\tilde{K}$ with support contained in $\tilde{K}$, namely $\mu^*(B) = \tilde{\mu}^*(B \cup \tilde{K})$ for any Borel subset B of K . Then μ^* clearly solves $A\mu = y$ and thus minimises l . $\square$

Note that [Lemma 3.10](#) is a particular case of [Proposition 3.11](#), but we believe this presentation makes the role of $\text{int}(A(\mathcal{M}_+))$ and $\partial A(\mathcal{M}_+)$ clearer.

Let us now finish this section by proving that not any point of the boundary may be associated to absolutely continuous measures. We denote S the simplex in $\mathbb{R}^d$, i.e.,

$$S := \left\{ w \in \mathbb{R}^d, w \geq 0 \mid \sum_{i=1}^d w_i = 1 \right\}. \quad (25)$$

Proposition 3.12. *Assume that $y \in \partial A(\mathcal{M}_+)$ is an extremal point of $A(\mathcal{M}_+) \cap S$. Then any measure satisfying $A\mu = y$ is a Dirac mass.*

We omit the proof, which is straightforward and relies on the linearity of the operator A and the fact that the only extremal points among probability measures are the Dirac masses [\[27\]](#).

4 Properties of ML-EM

We now turn our attention to the ML-EM algorithm (9) for the minimisation of the functional l (problem (7)).

4.1 Monotonicity and asymptotics

We first proceed to prove that the algorithm is monotonous, a property stemming from it being an expectation-maximisation algorithm.

Proposition 4.1. *For any $\mu_0 \in \mathcal{M}_+$ satisfying (10), ML-EM is monotone, i.e.,*

$$l(\mu_{k+1}) \leq l(\mu_k) \quad k \in \mathbb{N}.$$

Proof. We will build a so-called *surrogate* function, i.e., a function Q_k such that $l(\mu) \leq Q_k(\mu)$ for all μ in $X_k := \{ \mu \in \mathcal{M}_+ \mid \frac{d\mu}{d\mu_k} \geq 0 \}$, with equality for $\mu = \mu_k$, where Q_k is minimised at μ_{k+1} over X_k . Note that $\mu_{k+1} \in X_k$. The claim then follows since

$$l(\mu_{k+1}) \leq Q_k(\mu_{k+1}) \leq Q_k(\mu_k) = l(\mu_k).$$

For a measure μ , we define $P_i(\mu) := \frac{a_i \mu}{\langle \mu, a_i \rangle}$. One can check that $\langle P_i(\mu), \mathbf{1} \rangle = 1$, i.e., that $P_i(\mu)$ is always a probability measure. From the non-negativity of the divergence, i.e., $D(P_i(\mu_k) || P_i(\mu)) \geq 0$, we obtain

$$\log(\langle \mu, a_i \rangle) \geq \left\langle P_i(\mu_k), \log \left(\langle \mu_k, a_i \rangle \frac{d\mu}{d\mu_k} \right) \right\rangle,$$

with equality if $\mu = \mu_k$.

Notice now that $\left\langle P_i(\mu_k), \log\left(\langle \mu_k, a_i \rangle \frac{d\mu}{d\mu_k}\right) \right\rangle = \int_K \log\left(y_i^k \frac{d\mu}{d\mu_k}\right) \frac{a_i d\mu_k}{y_i^k}$, where we use the notation

$$y_i^k := \langle \mu_k, a_i \rangle.$$

This motivates the following definition for any $\mu \in X_k$:

$$Q_k(\mu) := \int_K d\mu - \sum_{i=1}^d y_i \int_K \log\left(y_i^k \frac{d\mu}{d\mu_k}\right) \frac{a_i d\mu_k}{y_i^k},$$

which is a surrogate function of l .

Making use once more of the concavity of the logarithm through $\log(b) \leq \log(a) + \frac{b-a}{a}$ for any $a, b > 0$ we may write for $\mu \in X_k$

$$\begin{aligned} Q_k(\mu) &\geq Q_k(\mu_{k+1}) + \int_K d(\mu - \mu_{k+1}) - \sum_{i=1}^d y_i \int_K \frac{d\mu_k}{d\mu_{k+1}} \left(\frac{d\mu}{d\mu_k} - \frac{d\mu_{k+1}}{d\mu_k} \right) a_i \frac{d\mu_k}{y_i^k} \\ &= Q_k(\mu_{k+1}) + \int_K d(\mu - \mu_{k+1}) - \sum_{i=1}^d y_i \int_K \frac{d\mu_k}{d\mu_{k+1}} \frac{a_i}{y_i^k} d(\mu - \mu_{k+1}) \\ &= Q_k(\mu_{k+1}) + \int_K \left(1 - \sum_{i=1}^d \frac{y_i a_i}{y_i^k} \frac{\mu_k}{\mu_{k+1}} \right) d(\mu - \mu_{k+1}) \\ &= Q_k(\mu_{k+1}), \end{aligned}$$

the right-hand side vanishing due to the definition of μ_{k+1} . $\square$

If the initial measure μ_0 has its mass sufficiently spread over K , we now prove that all cluster points of ML-EM are optimal solutions to the minimisation problem. This also proves that minimisers exist in full generality.

Proposition 4.2. *For any μ_0 satisfying (10) and*

$$\forall x \in K, \forall \eta > 0, \quad \mu_0(K \cap B(x, \eta)) > 0, \quad (26)$$

where $B(x, \eta)$ is the closed ball of centre x and radius η , then the weak cluster points of ML-EM minimise the function l . In particular,

$$l(\mu_k) \xrightarrow{k \rightarrow +\infty} \inf_{\mu \in \mathcal{M}_+} l(\mu).$$

Proof. Let $\bar{\mu}$ be a weak cluster point of ML-EM. Such a point exists since we have $\langle \mu_k, 1 \rangle = \sum_{i=1}^d y_i = 1$ for all $k \geq 1$. Thus, (μ_k) is a bounded sequence in $L^1(K)$. By the Banach–Alaoglu Theorem, it is thus weak-* compact in $\mathcal{M}_+$.

Also note that such a limit must satisfy $\langle \bar{\mu}, a_i \rangle > 0$ for any $i \in \text{supp}(y)$ (i.e., whenever $y_i > 0$). Otherwise, l would go to $+\infty$ since it is weak-* continuous, a contradiction with

the fact that l decreases along iterates and $l(\mu_0) < +\infty$. We may then pass to the limit in the definition of ML-EM and we find

$$\bar{\mu} = \bar{\mu} \left(\sum_{i=1}^d \frac{y_i a_i}{\langle \bar{\mu}, a_i \rangle} \right).$$

In particular, we have $\sum_{i=1}^d \frac{y_i a_i}{\langle \bar{\mu}, a_i \rangle} = 1$ on $\text{supp}(\bar{\mu})$. Now, assume by contradiction that $\sum_{i=1}^d \frac{y_i a_i}{\langle \bar{\mu}, a_i \rangle} > 1$ at a point $x_0 \in K$. Defining $b_k(x) := \sum_{i=1}^d \frac{y_i a_i(x)}{\langle \mu_{k-1}, a_i \rangle}$, by assumption $\lim_{k \rightarrow +\infty} b_k(x_0) > 1$. We may then by continuity find k_0 large enough such that $b_k(x) \geq 1 + \varepsilon$ for $k \geq k_0$ and $x \in K \cap B(x_0, \delta)$, where ε, η are small enough. Note that if we had $b_j = 0$ somewhere inside $K \cap B(x_0, \delta)$, we would have $a_i = 0$ as well and thus $b_j = 0$ for all j , which would contradict $b_j \geq 1 + \varepsilon$ on $B(x_0, \delta)$ for j large enough.

Since $\mu_k = \mu_{k-1} b_k(x) = \mu_0 \prod_{j=1}^k b_j(x)$, we may bound $\int_{K \cap B(x_0, \delta)} d\mu_k$ from below as follows

$$\begin{aligned} \int_{K \cap B(x_0, \delta)} d\mu_k &\geq \left(\int_{K \cap B(x_0, \delta)} \left(\prod_{j=1}^{k_0-1} b_j \right) d\mu_0 \right) (1 + \varepsilon)^{k-k_0+1} \\ &\geq C \mu_0(K \cap B(x_0, \delta)) (1 + \varepsilon)^{k-k_0+1} \end{aligned}$$

where $C > 0$ bounds from below the positive and continuous function $\prod_{j=1}^{k_0-1} b_j$ on the compact $K \cap B(x_0, \delta)$. From (26), we also have $\mu_0(K \cap B(x_0, \delta)) > 0$: the right-hand side tends to $+\infty$ as k tends to $+\infty$, which contradicts the boundedness of the left-hand side, due to $\langle \mu_k, 1 \rangle = 1$.

Summing up, we have proved both

$$\sum_{i=1}^d \frac{y_i a_i}{\langle \bar{\mu}, a_i \rangle} \leq 1 \text{ on } K, \quad \sum_{i=1}^d \frac{y_i a_i}{\langle \bar{\mu}, a_i \rangle} = 1 \text{ on } \text{supp}(\bar{\mu}),$$

which are exactly the optimality conditions for $\bar{\mu}$.

The fact that (μ_k) asymptotically minimises the function follows from the previous results. The sequence $\{l(\mu_k) \mid k = 0, \dots\}$ is indeed bounded, and its cluster points all equal the minimum of the function l since it is weak-* continuous. Thus, the whole sequence $\{l(\mu_k) \mid k = 0, \dots\}$ converges to the minimum of l . $\square$

4.2 Case $y \notin A(\mathcal{M}_+)$

When $y \notin A(\mathcal{M}_+)$, we know from the previous result and Corollary 3.6 that cluster points of ML-EM are sparse. Note that this is also the case for boundary points which are extremal in $A(\mathcal{M}_+) \cap S$, in virtue of Proposition 3.12.

We can go further in the case where $y_i = 0$ except for $y_{i_0} = 1$, for which we saw in Remark 3.8 that any minimiser μ^* of the function l for normalised measures, i.e., $l(\mu) = 1 - \log(\langle \bar{\mu}, a_{i_0} \rangle)$, will be such that $\text{supp}(\mu^*) \subset \arg \max a_{i_0}$. The goal of this subsection is to highlight a case where one clearly identifies the limiting measure of

ML-EM, and its dependence with respect to the initial measure μ_0 . This suggests that in the sparse case $y \notin A(\mathcal{M}_+)$, the position of Dirac masses will in general depend on the initial condition μ_0 .

We recall Laplace's method (see [32]) which holds for $f \in \mathcal{C}(K)$, $g \in \mathcal{C}^2(K)$ with a single non-degenerate interior maximum point $\bar{x}$, and reads:

$$\int_K f(x) e^{g(x)t} dx \sim (2\pi)^{n/2} \frac{f(\bar{x})}{\sqrt{|\det(H(\bar{x}))|}} \frac{e^{g(\bar{x})t}}{t^{\frac{n}{2}}} \quad \text{as } t \rightarrow \infty, \quad (27)$$

where $H(\bar{x})$ is the Hessian of g at $\bar{x}$.

Proposition 4.3. *Assume that $y_i = 0$ for $i \neq i_0$, $y_{i_0} = 1$. Assume further that $\arg \max a_{i_0} = \{\bar{x}_1, \dots, \bar{x}_p\}$ with $\bar{x}_j \in \text{int}(K)$ for all $j = 1, \dots, p$, that a_{i_0} is of class $\mathcal{C}^2$ and that the maximum points $\bar{x}_j$ are non-degenerate. Under these assumptions and for μ_0 absolutely continuous with continuous positive density (still denoted μ_0), the ML-EM sequence (μ_k) satisfies*

$$\mu_k \rightharpoonup \mu^\star := C \sum_{j=1}^p \frac{\mu_0(\bar{x}_j)}{\sqrt{|\det H_j|}} \delta_{\bar{x}_j},$$

where $C > 0$ is a normalising constant such that the limit has mass one, H_j is the Hessian of a_{i_0} at the point $\bar{x}_j$, and $\delta_{\bar{x}_j}$ is the Dirac mass centred at $\bar{x}_j$.

Proof. We first remark that the ML-EM iterates are then explicitly solved as

$$\mu_k = \frac{a_{i_0}^k \mu_0}{\int_K a_{i_0}^k(x) \mu_0(x) dx}.$$

We denote $M := \max_{x \in K} \log(a_{i_0}(x))$ and let $f \in \mathcal{C}(K)$ be a generic function. For δ small enough such that, for all j , $\bar{x}_j$ is the unique maximum point of a_{i_0} in $B(\bar{x}_j, \eta)$, we split contributions in the integral of $a_{i_0}^k \mu_0$ against f as follows

$$\begin{aligned} \langle a_{i_0}^k \mu_0, f \rangle &= \int_K f(x) \mu_0(x) e^{k \log(a_{i_0}(x))} dx \\ &= \sum_{j=1}^p \int_{B(\bar{x}_j, \eta)} f(x) \mu_0(x) e^{k \log(a_{i_0}(x))} dx + \int_{K \setminus \cup_{j=1}^p B(\bar{x}_j, \eta)} f(x) \mu_0(x) e^{k \log(a_{i_0}(x))} dx \end{aligned}$$

From Laplace's method (27), each term in the first sum can be estimated as

$$\int_{B(\bar{x}_j, \eta)} f(x) \mu_0(x) e^{k \log(a_{i_0}(x))} dx \sim (2\pi)^{n/2} \frac{f(\bar{x}_j) \mu_0(\bar{x}_j)}{\sqrt{|\det H_j|}} \frac{e^{Mk}}{k^{\frac{n}{2}}} \quad \text{as } k \rightarrow \infty,$$

whereas one can check that the second term is $o(e^{Mk})$. We end up with

$$\langle a_{i_0}^k \mu_0, f \rangle \sim (2\pi)^{n/2} \sum_{j=1}^p \left(\frac{f(\bar{x}_j) \mu_0(\bar{x}_j)}{\sqrt{|\det H_j|}} \right) \frac{e^{Mk}}{k^{\frac{n}{2}}} \quad \text{as } k \rightarrow \infty.$$

Applying this equivalent for $f = 1$ yields an equivalent for the denominator $\int_K a_{i_0}^k(x) \mu_0(x) dx$ in the explicit formula for μ_k , which is of the order of $e^{Mk}/k^{\frac{n}{2}}$. All in all, we find

$$\langle \mu_k, f \rangle \longrightarrow C \sum_{j=1}^p \left(\frac{f(\bar{x}_j) \mu_0(\bar{x}_j)}{\sqrt{|\det H_j|}} \right) = \left\langle C \sum_{j=1}^p \frac{\mu_0(\bar{x}_j)}{\sqrt{|\det H_j|}} \delta_{\bar{x}_j}, f \right\rangle = \langle \mu^*, f \rangle,$$

as $k \rightarrow \infty$, with $C > 0$ normalising μ^* , which proves the claim. $\square$

4.3 Case $y \in A(\mathcal{M}_+)$

When the data y is in the cone $A(\mathcal{M}_+)$, there are infinitely many measures satisfying $A\mu = y$, and when there are absolutely continuous ones, a desirable property of ML-EM is to converge to one of them rather than to a measure having a singular part. In order to address this question, we start with a Proposition of independent interest, valid for any y . It generalises a result which holds true in the discrete case. It gives information on the divergence of ML-EM iterates to *any* minimiser of (7).

Proposition 4.4. *Let μ^* be any minimiser of (7), and μ_0 satisfy (10). We further assume that*

$$D(\mu^* || \mu_0) < \infty.$$

Then the ML-EM iterates are such that

$$\forall k \in \mathbb{N}, D(\mu^* || \mu_{k+1}) \leq D(\mu^* || \mu_k) - l(\mu_k) + l(\mu^*).$$

In particular, the KL divergence to any optimum decreases:

$$\forall k \in \mathbb{N}, D(\mu^* || \mu_{k+1}) \leq D(\mu^* || \mu_k).$$

Proof. We recursively prove that $D(\mu^* || \mu_k) < \infty$. Assuming this holds true at the step k , it easily checked that the definition of ML-EM iterates is such that if μ_k is absolutely continuous with respect to μ^* , then so is μ_{k+1} . Furthermore,

$$\begin{aligned} -\log \left(\frac{d\mu_{k+1}}{d\mu^*} \right) &= -\log \left(\frac{d\mu_k}{d\mu^*} \right) - \log \left(\sum_{i=1}^d \frac{y_i a_i}{\langle \mu_k, a_i \rangle} \right) \\ &= -\log \left(\frac{d\mu_k}{d\mu^*} \right) - \log \left(\sum_{i=1}^d \frac{y_i a_i}{\langle \mu_k, a_i \rangle} \right) \sum_{i=1}^d \frac{y_i a_i}{\langle \mu^*, a_i \rangle} \\ &= -\log \left(\frac{d\mu_k}{d\mu^*} \right) + \sum_{i=1}^d \frac{y_i a_i}{\langle \mu^*, a_i \rangle} \log \left(\sum_{i=1}^d \frac{y_i a_i}{\langle \mu^*, a_i \rangle} \middle/ \sum_{i=1}^d \frac{y_i a_i}{\langle \mu_k, a_i \rangle} \right), \end{aligned} \quad (28)$$

where we twice took advantage of $\sum_{i=1}^d \frac{y_i a_i}{\langle \mu^*, a_i \rangle} = 1$ on $\text{supp}(\mu^*)$ due to the optimality conditions (18), a property we may use since we will eventually integrate against $d\mu^*$.

Now, we use the convexity of $(u, v) \mapsto u \log(\frac{u}{v})$ on $[0, +\infty) \times (0, +\infty)$ (following an idea of [14]) to bound the second term as follows

$$\sum_{i=1}^d \frac{y_i a_i}{\langle \mu^*, a_i \rangle} \log \left(\frac{\sum_{i=1}^d \frac{y_i a_i}{\langle \mu^*, a_i \rangle}}{\sum_{i=1}^d \frac{y_i a_i}{\langle \mu_k, a_i \rangle}} \right) \leq \sum_{i=1}^d y_i \frac{a_i}{\langle \mu^*, a_i \rangle} \log \left(\frac{\langle \mu_k, a_i \rangle}{\langle \mu^*, a_i \rangle} \right).$$

When integrated against $d\mu^*$, the right hand side simplifies to

$$\sum_{i=1}^d y_i \log \left(\frac{\langle \mu_k, a_i \rangle}{\langle \mu^*, a_i \rangle} \right) = -l(\mu_k) + l(\mu^*).$$

Wrapping up, the integration of (28) against $d\mu^*$ and the above inequality exactly yield the result. $\square$

Remark 4.5. As we saw, we typically expect μ^* to be a sparse measure when $y \notin A(\mathcal{M}_+)$, which means that the assumption that μ_0 is absolutely continuous with respect to μ^* will typically not be satisfied for μ_0 chosen to be constant over K . This result will instead come in handy when $y \in A(\mathcal{M}_+)$.

Corollary 4.6. *Assume that $y \in A(\mathcal{M}_+)$, and that there exists an absolutely continuous measure μ^* such that $A\mu^* = y$. Then, for any μ_0 absolutely continuous, with a positive and continuous density, all cluster points of ML-EM are absolutely continuous.*

Proof. Let μ^* be such an absolutely continuous optimum. The assumption on the initial measure μ_0 ensures that it satisfies conditions (10) and (26), as well as $D(\mu^*||\mu_0) < \infty$. We may then use Proposition 4.4 to obtain

$$\forall k \in \mathbb{N}, D(\mu^*||\mu_{k+1}) \leq D(\mu^*||\mu_k).$$

Let $\bar{\mu}$ be a cluster point of the iterates $\{\mu_k\}$ (it is then also an optimum). Since the function $(\mu, \nu) \mapsto D(\mu||\nu)$ is weak-* lower semi-continuous [25], we may pass to the limit in the inequality to obtain

$$D(\mu^*||\bar{\mu}) \leq \lim_{k \rightarrow +\infty} D(\mu^*||\mu_k).$$

In particular, $D(\mu^*||\bar{\mu}) < +\infty$, which by definition implies that $\bar{\mu}$ is absolutely continuous with respect to μ^* , and consequently also with respect to the Lebesgue measure. $\square$

We are now in a position to prove the main result of this section, where we use the notations of Proposition 3.11.

Theorem 4.7. *Assume that conditions (23) and (24) hold, and that $\tilde{y} \in \text{int}(\tilde{A}(\mathcal{M}_+(\tilde{K})))$. Then, if the initial measure μ_0 is absolutely continuous, with a positive and continuous density, all cluster points of ML-EM are absolutely continuous.*

Proof. Under these assumptions, by [Proposition 3.11](#), there exists an absolutely continuous measure μ^* such that $A\mu^* = y$. The claim is then a consequence of [Corollary 4.6](#). $\square$

Note that these results also cover the case of [Lemma 3.10](#): if $y \in \text{int}(A(\mathcal{M}_+))$ and the functions $\{a_i\}$ are linearly independent, cluster points of ML-EM are absolutely continuous, whenever the initial measure μ_0 is absolutely continuous with a positive and continuous density.

5 Statistics

In this section, we estimate the probability that the data y stays in the image cone, i.e., $y \in A(\mathcal{M}_+)$.

Let us first go back to modelling (before any normalisation) by introducing a *time variable*. We assume that the real image is given by $\mu_r \in \mathcal{M}_+$ which represents the image in $U.s^{-1}$ where U is the relevant unit depending on the context. The acquisition during time t gives rise to independent random variables $N_i \sim \mathcal{P}(\gamma_i t)$ where $\gamma_i := \langle \mu_r, a_i \rangle$, whose sum $N = \sum_{i=1}^d N_i$ is $\mathcal{P}(\gamma t)$ with $\gamma := \sum_{i=1}^d \gamma_i$.

The expected frequencies are given by $y_r := (\frac{\gamma_1}{\gamma}, \dots, \frac{\gamma_d}{\gamma})$, which is an element of $A(\mathcal{M}_+) \cap S$, where we recall that S , given by [\(25\)](#), is the simplex in $\mathbb{R}^d$.

With our previous notations, n_i and n are thus realisations of the random variables N_i , and N , and now $y = (\frac{N_1}{N}, \dots, \frac{N_d}{N})$ is a random variable. We emphasise that it is an estimator for y_r which depends on time by using the notation $\hat{y}_t$. When conditioned on the fact that $N = n$, we will denote $\hat{y}_n := (\frac{N_1}{n}, \dots, \frac{N_d}{n})$.

By the law of large numbers, $\hat{y}_t$ tends to y_r almost surely as $t \rightarrow +\infty$, so we know that if $y_r \in \text{int}(A(\mathcal{M}_+))$, a long-enough time acquisition will ensure that $\hat{y}_t \in \text{int}(A(\mathcal{M}_+))$ as well, avoiding sparse measures as a solution to the maximum likelihood problem.

We now wish to give quantitative bounds for $\mathbb{P}(\hat{y}_t \notin A(\mathcal{M}_+))$, one with a conditioning on the number of events n , the other without such a conditioning. The aim is to address the following questions:

- a posteriori, for a given number of points n , how small is the probability that $\hat{y}_n \notin A(\mathcal{M}_+)$?
- a priori, how large should the time t be for the probability that $\hat{y}_t \notin A(\mathcal{M}_+)$ to be small enough?

The celebrated Sanov's Theorem [\[28\]](#) states that the empirical distribution has an exponentially small probability of being in a set which does not contain the real distribution, where the exponential is controlled by the Kullback–Leibler distance from the real distribution to the set. We thus define

$$\varepsilon := \inf_{p \in S \cap A(\mathcal{M}_+)^c} d(p, y_r),$$

which is the Kullback–Leibler divergence of y_r to the boundary of the set $A(\mathcal{M}_+)$ (intersected with the simplex S).

In both cases, we shall give two different bounds which might be relevant in different regimes in the parameters (d, n) and (d, t) , respectively.

Proposition 5.1. *The following concentration bounds hold:*

$$\mathbb{P}(\hat{y}_n \notin A(\mathcal{M}_+)) \leq \begin{cases} (n+1)^d e^{-n\varepsilon}, \\ 2d e^{-n\varepsilon/d}. \end{cases}$$

Proof. Conditioned on $N = n$, the random vector $\hat{y}_n$ follows the multinomial distribution of parameters n and y_r . The first inequality is then nothing but a direct application of Sanov's Theorem [28]. The second is more recent and given in Lemma 6 of [18]. $\square$

We now proceed to the case with time:

Theorem 5.2. *The following concentration bounds hold:*

$$\mathbb{P}(\hat{y}_t \notin A(\mathcal{M}_+)) \leq \begin{cases} C(d)(1 + (\gamma t)^d) e^{-\gamma t\varepsilon}, \\ 2d e^{-\gamma t\varepsilon/d}, \end{cases}$$

where $C(d)$ is a combinatorial constant which depends only on d and satisfies $C(d) \leq \left(\frac{a(d+1)}{\log(d+2)}\right)^{d+1}$, with $a = 0.792$.

Proof. We may write

$$\mathbb{P}(\hat{y}_t \notin A(\mathcal{M}_+)) = \mathbb{E}(\mathbb{P}(\hat{y}_N \notin A(\mathcal{M}_+)) | N) \leq \mathbb{E}(g(N))$$

where $g(n) = (n+1)^d e^{-n\varepsilon}$ or $2d e^{-n\varepsilon/d}$ from the previous proposition. It is now a matter of estimating this expectation with $N \sim \mathcal{P}(\gamma t)$. In the second case,

$$\begin{aligned} \mathbb{E}(g(N)) &= 2d e^{-\gamma t} \sum_{n=0}^{+\infty} e^{-n\varepsilon/d} \frac{(\gamma t)^n}{n!} \\ &= 2d e^{-\gamma t(1 - \exp(-\varepsilon/d))} \leq 2d e^{-\gamma t\varepsilon/d}, \end{aligned}$$

from $1 - e^{-u} \geq u$.

In the first case, $\mathbb{E}(g(N)) = e^{-\gamma t} \varphi_d(\gamma t e^{-\varepsilon})$, with

$$\varphi_d(x) := \sum_{n=0}^{+\infty} (n+1)^d \frac{x^n}{n!},$$

and $x := \gamma t e^{-\varepsilon}$, which we now estimate. We may integrate to find

$$\int_0^x \varphi_d(u) du = \sum_{n=1}^{+\infty} n^d \frac{x^n}{n!} =: T_d(x) e^x,$$

where T_d is the so-called Touchard polynomial of order d , which has degree d .

This allows to go back to $\varphi_d(x)$ as $\varphi_d(x) = (T'_d(x) + T_d(x)) e^x = \frac{T_{d+1}(x)}{x} e^x$, using a well-known property of Touchard polynomials. The Touchard polynomial of order d has integer coefficients (the Stirling numbers), whose sum is given by the so-called Bell number B_d .

Using the crude bound $P(x) = \sum_{k=0}^d a_k x^k \leq (\sum_{k=0}^d a_k)(1+x^d)$ valid for all $x \geq 0$ when P is a polynomial with non-negative coefficients, we may write $\frac{T_{d+1}(x)}{x} \leq B_{d+1}(1+x^d)$ for $x \geq 0$.

Summing up, we have

$$\mathbb{E}(g(N)) \leq B_{d+1} e^{-\gamma t} (1+x^d) e^x = B_{d+1} (1+(\gamma t e^{-\varepsilon})^d) e^{-\gamma t(1-\exp(-\varepsilon))} \leq C(d) (1+(\gamma t)^d) e^{-\gamma t \varepsilon},$$

where $C(d) := B_{d+1}$. The bound about Bell numbers such as $C(d)$, stated in the proposition, can be found in [4]. We use them here as bounds on the moments of a Poisson random variable. $\square$

Remark 5.3. Although the bound coming from Sanov's theorem is sharper at the limit $n \rightarrow +\infty$ or $t \rightarrow +\infty$, it may be that the other one is relevant for realistic values d , n or t . If these bounds are taken as functions of ε , it can be checked that the alternative bound becomes more stringent in the regime where $\varepsilon \ll \frac{d}{n} \log(n)$.

6 Numerical simulations

We perform simulations using the Python library Operator Discretization Library [1]. All of them are run with a 2D PET operator A having 90 views and 128 tangential positions, leading to a number of (pairs of) detectors $d = 11520$. The image resolution is 256×256 .

We choose as phantom a torus μ_r , as displayed in Figure 1, which has a maximum of 1. We compute $y_t \sim \frac{1}{t} \mathcal{P}(tA\mu_r)$ for different acquisition times t , so that the higher t , the lower the noise level.

In what follows, we will consider three different noise levels, associated to different values of acquisition times. For each of these values, we are interested in seeing whether iterations lead to sparse measures or not. From our results, this is equivalent to testing if $y_t \in A(\mathcal{M}_+)$. A first crude estimate of this problem is to plot $d(y_t || A\mu_k)$ and check whether this quantity converges to zero, which by theory is equivalent to $y_t \in A(\mathcal{M}_+)$.

Dual certificates. As it is difficult to decide against or for a convergence towards zero, we will also look for *dual certificates* $\lambda \in A(\mathcal{M}_+)^*$ such that the dual function g defined by (20) satisfies $g(\lambda) > 0$. Indeed, weak duality ensures that $\min d(y_t || A\mu) \geq g(\lambda)$, so that any dual certificate proves that $y_t \notin A(\mathcal{M}_+)$.

In order to find a good choice of λ , we will actually compute $\lambda_k = 1 - \frac{y_t}{A\mu_k}$ along iterates, which we know should converge to λ^* such that $A^* \lambda^* \geq 0$, i.e., $\lambda^* \in A(\mathcal{M}_+)^*$. For a fixed k , there is no reason that $\lambda_k \in A(\mathcal{M}_+)^*$, so we will just add a small appropriate constant c to λ_k and check that it provides a dual certificate, that is, we check if $g(\lambda_k + c) > 0$.

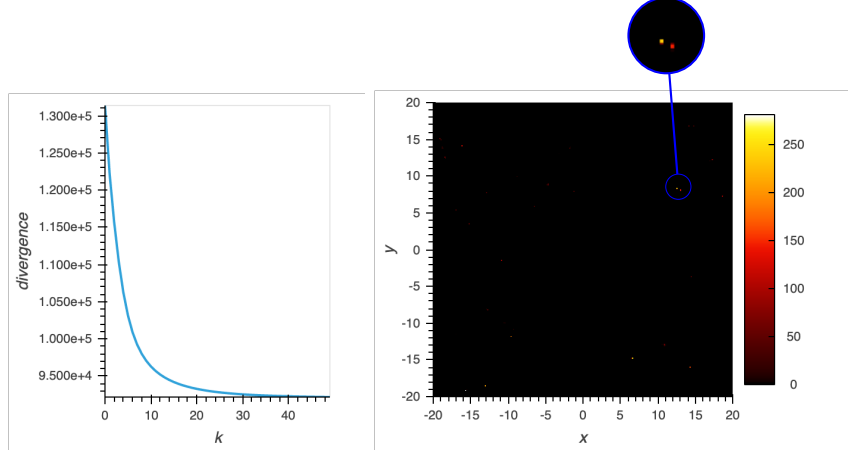


Figure 2: Divergence $d(y_t||A\mu_k)$ up to $k_0 = 50$ and reconstruction μ_{k_0} , for $t = 10^{-3}$, with a zoom on some pixels having high values.

High noise. For $t = 10^{-3}$, we quickly obtain very sparse results, as shown by the image obtained after 50 iterations in Figure 2. We also plot the evolution of the functional along iterates $d(y_t||A\mu_k)$, which seems to converge to a positive value of the order of 9.1×10^4 , hinting at the fact that $y_t \notin A(\mathcal{M}_+)$, a result we could confirm by finding a dual certificate.

Intermediate noise. For $t = 1$, the resulting image is noisier as shown with the result of 50 iterations in Figure 3, with a maximum at about 5. We also plot the evolution of the functional along iterates $d(y_t||A\mu_k)$, which seems to converge to a positive value, hinting at the fact that $y_t \notin A(\mathcal{M}_+)$.

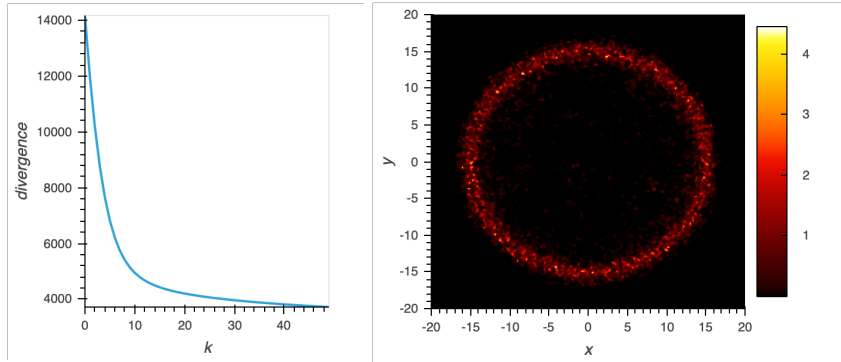


Figure 3: Divergence $d(y_t||A\mu_k)$ up to $k_0 = 50$ and reconstruction μ_{k_0} , for $t = 1$.

If we keep iterating up to 1000 iterations, some pixels take increasingly high values, with a maximum of about 10, see Figure 4. We are also able to provide a dual certificate

showing that $y_t \notin A(\mathcal{M}_+)$. The functional $d(y_t||A\mu_k)$ has kept decreasing, but the dual certificate shows that optimality has almost been reached. Our hypothesis for these results is that convergence to the sum of Dirac masses is very slow.

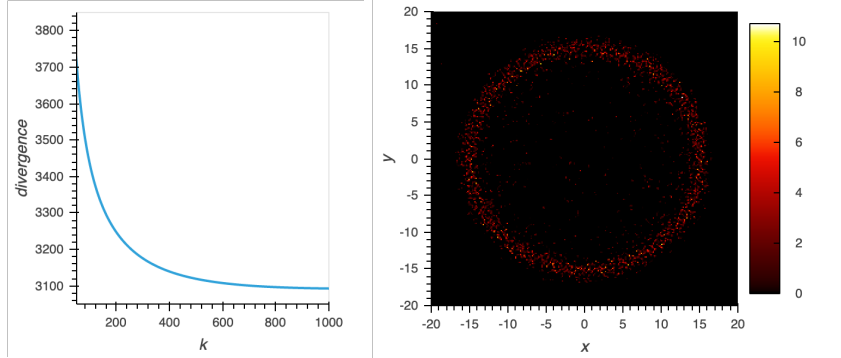


Figure 4: Divergence $d(y_t||A\mu_k)$ from $k = 50$ up to $k_0 = 1000$ and reconstruction μ_{k_0} , for $t = 1$.

Low noise. For $t = 10^3$, the image obtained remains smooth and very close to the phantom, as shown with the result of 1000 iterations in Figure 5. We also plot the evolution of the functional along iterates $d(y_t||A\mu_k)$, which seems to converge to zero, suggesting that $y_t \in A(\mathcal{M}_+)$. We also mention that we could not certify that $y_t \notin A(\mathcal{M}_+)$.

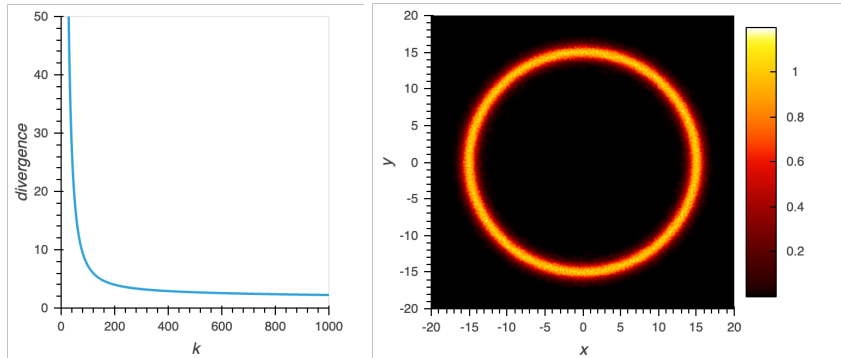


Figure 5: Divergence $d(y_t||A\mu_k)$ up to $k_0 = 1000$ and reconstruction μ_{k_0} , for $t = 10^3$.

7 Open problems and perspectives

Convergence of iterates. As shown in Proposition 4.2, cluster points of the weak-* compact iterates of ML-EM are optimal. A problem left open is that of the convergence of the whole sequence to a single point, which is a tall order since there are in general many optimal points. The discrete equivalent to Proposition 4.4 allows to prove the

full convergence of iterates in the discrete case, thanks to the continuity of the discrete Kullback–Leibler divergence [31]. Its mere lower semi-continuity in continuum does not allow us to obtain a similar result, although we postulate this holds true as well.

Regularisation. Another interesting issue in the light of our sparsity results is that of regularisation. How should one choose appropriate additional regularisation terms to alleviate the problem? Similarly and as is done in [26] for continuous data, analysing the alternative strategy of regularising by early stopping is worthy of interest, since this is usual practice in PET.

Going further in the case of PET. For PET, the functions a_i are actually close to being singular measures concentrated on a line. Studying the effect of this near-singularity on our results is of practical interest. More precisely, one could for specific geometries analyse the typical minimum set of a function of the form $A^*\lambda = \sum_{i=1}^d \lambda_i a_i$.

It would also be natural to look at the effect of binning (i.e., aggregating detectors) on the constant $\inf_{p \in S \cap A(\mathcal{M}_+)^c} d(p, y_r)$. Indeed, [Theorem 5.2](#) shows its importance when it comes to sparsity, justifying to try and make this distance as large as possible.

Sparsity results in general. We intend to investigate the generality of these sparsity results in the context of other divergences. The squared-distance is a popular one, but generalisations have recently been advocated for in the literature, such as the β -divergences [7].

Acknowledgements. We are grateful to Sebastian Banert for fruitful discussions about optimisation in Banach spaces, as well as Axel Ringh and Johan Karlsson for bringing the moment matching problem to our attention. We acknowledge support from the Swedish Foundation of Strategic Research grant AM13-004.

References

- [1] ADLER, J., KOHR, H., AND ÖKTEM, O. ODL-a Python framework for rapid prototyping in inverse problems. *Royal Institute of Technology* (2017).
- [2] BAUMEISTER, J., AND LEITÃO, A. Topics in inverse problems.
- [3] BENVENUTO, F., AND PIANA, M. Regularization of multiplicative iterative algorithms with nonnegative constraint. *Inverse Problems* 30, 3 (2014), 035012.
- [4] BEREND, D., AND TASSA, T. Improved bounds on Bell numbers and on moments of sums of random variables. *Probability and Mathematical Statistics* 30, 2 (2010), 185–205.
- [5] BERTERO, M., AND BOCCACCI, P. *Introduction to inverse problems in imaging*. CRC press, 1998.
- [6] BOYD, S., AND VANDENBERGHE, L. *Convex optimization*. Cambridge university press, 2004.

- [7] CAVALCANTI, Y. C., OBERLIN, T., DOBIGEON, N., FÉVOTTE, C., STUTE, S., RIBEIRO, M.-J., AND TAUBER, C. Factor analysis of dynamic PET images: beyond Gaussian noise. *IEEE transactions on medical imaging* (2019).
- [8] CSISZÁR, I. Information geometry and alternating minimization procedures. *Statistics and decisions* 1 (1984), 205–237.
- [9] DEMPSTER, A. P., LAIRD, N. M., AND RUBIN, D. B. Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)* 39, 1 (1977), 1–22.
- [10] FESSLER, J. A., CLINTHOME, N. H., AND ROGERS, W. L. On complete-data spaces for PET reconstruction algorithms. *IEEE Trans. Nuc. Sci* 40, 4 (1993), 1055–61.
- [11] GEORGIU, T. T. Solution of the general moment problem via a one-parameter imbedding. *IEEE transactions on automatic control* 50, 6 (2005), 811–826.
- [12] HINKLE, J., SZEGEDI, M., WANG, B., SALTER, B., AND JOSHI, S. 4D CT image reconstruction with diffeomorphic motion model. *Medical image analysis* 16, 6 (2012), 1307–1316.
- [13] HUDSON, H. M., AND LARKIN, R. S. Accelerated image reconstruction using ordered subsets of projection data. *IEEE transactions on medical imaging* 13, 4 (1994), 601–609.
- [14] IUSEM, A. N. A short convergence proof of the EM algorithm for a specific poisson model. *Brazilian Journal of Probability and Statistics* (1992), 57–67.
- [15] JACOBSON, M., AND FESSLER, J. A. Joint estimation of image and deformation parameters in motion-corrected PET. In *2003 IEEE Nuclear Science Symposium. Conference Record (IEEE Cat. No. 03CH37515)* (2003), vol. 5, IEEE, pp. 3290–3294.
- [16] LAST, G., AND PENROSE, M. *Lectures on the Poisson process*, vol. 7. Cambridge University Press, 2017.
- [17] MAIR, B., RAO, M., AND ANDERSON, J. Positron emission tomography, Borel measures and weak convergence. *Inverse Problems* 12, 6 (1996), 965.
- [18] MARDIA, J., JIAO, J., TÁNCZOS, E., NOWAK, R. D., AND WEISSMAN, T. Concentration inequalities for the empirical distribution. *arXiv preprint arXiv:1809.06522* (2018).
- [19] MÜLTHEI, H. Iterative continuous maximum-likelihood reconstruction method. *Mathematical methods in the applied sciences* 15, 4 (1992), 275–286.
- [20] MÜLTHEI, H., SCHORR, B., AND TÖRNIG, W. On an iterative method for a class of integral equations of the first kind. *Mathematical methods in the applied sciences* 9, 1 (1987), 137–168.
- [21] MÜLTHEI, H., SCHORR, B., AND TÖRNIG, W. On properties of the iterative maximum likelihood reconstruction method. *Mathematical Methods in the Applied Sciences* 11, 3 (1989), 331–342.
- [22] NATTERER, F., AND WÜBBELING, F. *Mathematical methods in image reconstruction*, vol. 5. Siam, 2001.
- [23] ÖKTEM, O., POUCHOL, C., AND VERDIER, O. Spatiotemporal PET reconstruction using ML-EM with learned diffeomorphic deformation. *submitted* (2019).
- [24] OLLINGER, J. M., AND FESSLER, J. A. Positron-emission tomography. *IEEE Signal Processing Magazine* 14, 1 (1997), 43–55.

- [25] POSNER, E. Random coding strategies for minimum entropy. *IEEE Transactions on Information Theory* 21, 4 (1975), 388–391.
- [26] RESMERITA, E., ENGL, H. W., AND IUSEM, A. N. The expectation-maximization algorithm for ill-posed integral equations: a convergence analysis. *Inverse Problems* 23, 6 (2007), 2575.
- [27] RUDIN, W. *Functional analysis*, second ed. International Series in Pure and Applied Mathematics. McGraw-Hill, Inc., New York, 1991.
- [28] SANOV, I. N. On the probability of large deviations of random variables. *Selected Translations in Mathematical Statistics and Probability* 1 (1961), 213–244.
- [29] SHEPP, L. A., AND VARDI, Y. Maximum likelihood reconstruction for emission tomography. *IEEE transactions on medical imaging* 1, 2 (1982), 113–122.
- [30] SILVERMAN, B., JONES, M., WILSON, J., AND NYCHKA, D. A smoothed EM approach to indirect estimation problems, with particular reference to stereology and emission tomography. *Journal of the Royal Statistical Society: Series B (Methodological)* 52, 2 (1990), 271–303.
- [31] VARDI, Y., SHEPP, L., AND KAUFMAN, L. A statistical model for positron emission tomography. *Journal of the American statistical Association* 80, 389 (1985), 8–20.
- [32] WONG, R. *Asymptotic approximations of integrals*, vol. 34. SIAM, 2001.