

Data-Driven Priors in the Maximum Entropy on the Mean Method for Linear Inverse Problems.

Matthew King-Roskamp, Rustum Choksi, and Tim Hoheisel

Abstract. We establish the theoretical framework for implementing the maximum entropy on the mean (MEM) method for linear inverse problems in the setting of approximate (data-driven) priors. We prove a.s. convergence for empirical means and further develop general estimates for the difference between the MEM solutions with different priors μ and ν based upon the epigraphical distance between their respective log-moment generating functions. These estimates allow us to establish a rate of convergence in expectation for empirical means. We illustrate our results with denoising on MNIST and Fashion-MNIST data sets.

1. Introduction. Linear inverse problems are pervasive in data science. A canonical example (and our motivation here) is denoising and deblurring in image processing. Machine learning algorithms, particularly neural networks trained on large data sets, have proven to be a game changer in solving these problems. However, most machine learning algorithms suffer from the lack of a foundational framework upon which to rigorously assess their performance. Thus, there is a need for mathematical models which are on one end, data driven, and on the other end, open to rigorous evaluation. In this article, we devise and analyze one such model based upon what is known as *Maximum Entropy on the Mean* (MEM) (described in some detail in [subsection 1.2](#)).

1.1. History and state of the art of the MEM method. Emerging from ideas of E.T. Jaynes in 1957 [23, 24], various forms and interpretations of MEM (see [7, 20, 30, 13, 29]) have appeared in the literature. Applications have occurred in different disciplines such as earth sciences [17, 33, 43], crystallography [31, 32], and medical imaging [1, 10, 11, 21, 22]. Recently, the MEM method has been shown to be a powerful tool for blind deblurring of images that possess some form of symbology (for example, UPC and QR barcodes) [35, 34]. However, MEM methods are not widely used and have yet to become a modern tool for solving contemporary data-driven inverse problems in image processing and machine learning.

MEM methods require problem-specific knowledge in the form of a *statistical prior*. For a prior based upon a known distribution, the theory is well understood, with dedicated algorithms for its implementation [44]. This work is the first to incorporate data in a systematic way into the MEM framework for linear inverse problems. While the incorporation is natural, we provide theoretical guarantees of convergence and upper bounds on rates of convergence without requiring model assumptions. In addition to providing the theoretical framework, we present several numerical examples for denoising images from *MNIST* [15] and *Fashion-MNIST* [47] data sets; showcasing that our work results in a data-driven model with numerical implementation via standard optimization routines.

1.2. Brief overview of the MEM method. Let us now provide some details, summarizing the MEM method for linear inverse problems. Full details will be provided in the next section. Our canonical inverse problem takes the following form

$$(1.1) \quad b = C\bar{x} + \eta.$$

The unknown solution $\bar{x}$ is a vector in $\mathbb{R}^d$; the observed data is $b \in \mathbb{R}^m$; $C \in \mathbb{R}^{m \times d}$, and $\eta \sim \mathcal{Z}$ is an random noise vector in $\mathbb{R}^m$ drawn from noise distribution $\mathcal{Z}$. In the setting of image processing, $\bar{x}$ denotes the ground truth image with d pixels, C is a blurring matrix with typically $d = m$, and the observed noisy (and blurred image) is b . For known C , we seek to recover the ground truth $\bar{x}$ from b . In certain classes of images, the case where C is also unknown (blind deblurring) can also be solved with the MEM framework (cf. [35, 34]) but we will not focus on this here. In fact, our numerical experiments will later focus purely on denoising, i.e., $C = I$. The power of MEM is to exploit the fact that there exists a prior distribution μ for the space of admissible ground truths.

48 The basis of the method is the *MEM function* $\kappa_\mu : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{+\infty\}$ defined as

$$49 \quad \kappa_\mu(x) := \inf \{ \text{KL}(Q \| \mu) : Q \in \mathcal{P}(\mathcal{X}), \mathbb{E}_Q = x \},$$

50 where $\text{KL}(Q \| \mu)$ denotes the Kullback-Leibler (KL) divergence between the probability distributions
51 μ and Q (see [Subsection 2.2](#) for the definition). With κ_μ in hand, our proposed solution to (1.1) is

$$52 \quad (1.2) \quad \bar{x}_\mu = \underset{x \in \mathbb{R}^d}{\operatorname{argmin}} \{ \alpha g_b(Cx) + \kappa_\mu(x) \},$$

53 where g_b is any (closed, proper) convex function that measures *fidelity* of Cx to b . The function
54 g_b depends on b and can in principle be adapted to the noise distribution $\mathcal{Z}$. For example, as was
55 highlighted in [44], one can take the *MEM estimator* (an alternative to the well-known *maximum*
56 *likelihood estimator*) based upon a family of distributions (for instance, if the noise is Gaussian,
57 then the MEM estimator is the familiar $g_b(\cdot) = \frac{1}{2} \|(\cdot) - b\|_2^2$). Finally $\alpha > 0$ is a fidelity parameter.

58 The variational problem (1.2) is solved via its Fenchel dual. As we explain in [Subsection 2.2](#),
59 we exploit the well-known connection in the large deviations literature that, under appropriate
60 assumptions, the MEM function κ_μ is simply the Cramér rate function defined as the Fenchel
61 conjugate of the log-moment generating function (LMGF)

$$62 \quad L_\mu(y) := \log \int_{\mathcal{X}} \exp\langle y, \cdot \rangle d\mu.$$

63 Under certain assumptions on g_b (cf. [Subsection 2.2](#)) we obtain strong duality

$$64 \quad (1.3) \quad \min_{x \in \mathbb{R}^d} \alpha g_b(Cx) + \kappa_\mu(x) = - \min_{z \in \mathbb{R}^m} \alpha g^*(-z/\alpha) + L_\mu(C^T z),$$

65 and, more importantly, a primal-dual recovery is readily available: If $\bar{z}_\mu$ is a solution to the dual
66 problem (the argmin of the right-hand-side of (1.3)) then

$$67 \quad \bar{x}_\mu := \nabla L_\mu(C^T \bar{z})$$

68 is the unique solution of the primal problem. This is the MEM method in a nutshell.

69 **1.3. Contributions and outline of the article.** In this article, we address the following ques-
70 tions: Suppose we do not have full access to the underlying prior distribution μ ; rather we have
71 access to an approximation sequence μ_n which in a suitable sense (e.g. weak convergence of mea-
72 sures) converges to μ . Does the approximate MEM solution $\bar{x}_{\mu_n}$ converge to the solution $\bar{x}_\mu$, and
73 if so, at which rate? A key feature of the MEM approach is that one does not need to quantify
74 the convergence of μ_n to μ , but rather only approximate the LMGF L_μ from data. Hence our
75 analysis is based on the *closeness* of L_{μ_n} to L_μ . This results in the *closeness* of the dual solutions
76 $\bar{z}_n$ and in turn the primal solutions $\bar{x}_{\mu_n}$. Here, we leverage the fundamental work of Wets et al. on
77 *epigraphical distances*, *epigraphical convergence*, and *epi-consistency* ([37],[40],[26]).

78 Our results are presented in four sections. In [Section 3](#), we work with a general g_b satisfying
79 standard assumptions. We consider the simplest way of approximating μ via empirical means of n
80 i.i.d. samples from μ . In [Theorem 3.9](#), we prove that the associated MEM solutions $\bar{x}_{\mu_n}$ converge
81 almost surely to the solution $\bar{x}_\mu$ with full prior. In fact, we prove a slightly stronger result pertaining
82 to ε_n -solutions as $\varepsilon_n \searrow 0$. This result opens the door to two natural questions: (i) At which rate
83 do the solutions converge? (ii) Empirical means is perhaps the simplest way of approximating μ
84 and what is the corresponding rate? Given that the MEM method rests entirely on the LMGF of
85 the prior, it is natural to ask how the rate depends on an approximation to the LMGF. So, if we
86 used a different way of approximating μ , what would the rate look like? We address these questions
87 for the case $g_b = \frac{1}{2} \|(\cdot) - b\|_2^2$. In [Section 4](#) we provide insight into the second question first via
88 a deterministic estimate which controls the difference in the respective solutions associated with

two priors ν and μ based upon the epigraphical distance between their respective LMGFs. We again prove a general result for ε -solutions associated with prior μ (cf. [Theorem 4.7](#)). In [Section 5](#), we apply this bound to the particular case of the empirical means approximation, proving a $\frac{1}{n^{1/4}}$ convergence rate (cf. [Theorem 5.5](#)) in expectation.

Finally, in [Section 6](#), we present several numerical experiments for denoising based upon a finite MNIST data set. These serve not to compete with any of the state-of-the-art machine learning-based denoising algorithm, but rather to highlight the effectiveness of our data-driven mathematical model which is fully supported by theory.

Remark 1.1 (Working at the higher level of the probability distribution of the solution). As in [\[35, 34\]](#), an equivalent formulation of the MEM problem is to work not at the level of the x , but rather at the level of the probability distribution of the ground truth, i.e., we seek to solve

$$\bar{Q} = \operatorname{argmin}_{Q \in \mathcal{P}(\mathcal{X})} \{ \alpha g_b(C\mathbb{E}_Q) + \operatorname{KL}(Q\|\mu) \},$$

where one can recover the previous image-level solution as $\bar{x}_\mu = \mathbb{E}_{\bar{Q}}$. As shown in [\[34\]](#), under appropriate assumptions this reformulated problem has exactly the same dual formulation as in the right-hand-side of [\(1.3\)](#). Because of this one has full access to the entire probability distribution of the solution, not just its expectation. This proves useful in our MNIST experiments where the optimal ν is simply a weighted sum of images uniformly sampled from the MNIST set. For example, one can do thresholding (or masking) at the level of the optimal ν (cf. the examples in [Section 6](#)).

Notation: $\bar{\mathbb{R}} := \mathbb{R} \cup \{\pm\infty\}$ is the extended real line. The standard inner product on $\mathbb{R}^n$ is $\langle \cdot, \cdot \rangle$ and $\|\cdot\|$ is the Euclidean norm. For $C \in \mathbb{R}^{m \times d}$, $\|C\| = \sqrt{\lambda_{\max}(C^T C)}$ is its spectral norm, and analogously $\sigma_{\min}(C) = \sqrt{\lambda_{\min}(C^T C)}$ is the smallest singular value of C . The trace of C is denoted $\operatorname{Tr}(C)$. For smooth $f : \mathbb{R}^d \rightarrow \mathbb{R}$, we denote its gradient and Hessian by ∇f and $\nabla^2 f$, respectively.

2. Tools from convex analysis and the MEM method for solving the problem (1.1) .

2.1. Convex analysis. We present here the tools from convex analysis essential to our study. We refer the reader to the standard texts by Bauschke and Combettes [\[5\]](#) or Chapters 2 and 11 of Rockafellar and Wets [\[37\]](#) for further details. Let $f : \mathbb{R}^d \rightarrow \bar{\mathbb{R}}$. The domain of f is $\operatorname{dom}(f) := \{x \in \mathbb{R}^d \mid f(x) < +\infty\}$. We call f proper if $\operatorname{dom}(f)$ is nonempty and $f(x) > -\infty$ for all x . We say that f is lower semicontinuous (lsc) if $f^{-1}([-\infty, a])$ is closed (possibly empty) for all $a \in \mathbb{R}$. We define the (Fenchel) conjugate $f^* : \mathbb{R}^d \rightarrow \bar{\mathbb{R}}$ of f as $f^*(x^*) := \sup_{x \in \mathbb{R}^d} \{\langle x, x^* \rangle - f(x)\}$. A proper f is said to be convex, if $f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y)$ for every $x, y \in \operatorname{dom}(f)$ and all $\lambda \in (0, 1)$. If the former inequality is strict for all $x \neq y$, then f is said to be strictly convex. Finally, if f is proper and there is a $c > 0$ such that $f - \frac{c}{2}\|\cdot\|^2$ is convex we say f is c -strongly convex. In the case where f is (continuously) differentiable on $\mathbb{R}^d$, then f is c -strongly convex if and only if

$$(2.1) \quad f(y) - f(x) \geq \nabla f(x)^T (y - x) + \frac{c}{2} \|y - x\|_2^2 \quad \forall x, y \in \mathbb{R}^d.$$

The subdifferential of a convex function $f : \mathbb{R}^d \rightarrow \bar{\mathbb{R}}$ at $\bar{x} \in \operatorname{dom}(f)$ is $\partial f(\bar{x}) = \{x^* \in \mathbb{R}^d \mid \langle x - \bar{x}, x^* \rangle \leq f(x) - f(\bar{x}), \forall x \in \mathbb{R}^d\}$. A function $f : \mathbb{R}^d \rightarrow \bar{\mathbb{R}}$ is said to be level-bounded if for every $\alpha \in \mathbb{R}$, the set $f^{-1}([-\infty, \alpha])$ is bounded (possibly empty). A function f is (level) coercive if it is bounded below on bounded sets and satisfies

$$\liminf_{\|x\| \rightarrow +\infty} \frac{f(x)}{\|x\|} > 0.$$

In the case f is proper, lsc, and convex, level-boundedness is equivalent to level-coerciveness [\[37, Corollary 3.27\]](#). A function f is said to be supercoercive if $\liminf_{\|x\| \rightarrow +\infty} \frac{f(x)}{\|x\|} = +\infty$.

A point $\bar{x}$ is said to be an ε -minimizer of a proper function f if $f(\bar{x}) \leq \inf_{x \in \mathbb{R}^d} f(x) + \varepsilon$ for some $\varepsilon > 0$. We denote the set of all such points as $S_\varepsilon(f)$. Correspondingly, the solution set of proper function f is denoted as $\operatorname{argmin}(f) = S_0(f) =: S(f)$.

133 The epigraph of a function $f : \mathbb{R}^d \rightarrow \overline{\mathbb{R}}$ is the set $\text{epi}(f) := \{(x, \alpha) \in \mathbb{R}^d \times \overline{\mathbb{R}} \mid \alpha \geq f(x)\}$.
 134 A sequence of functions $f_n : \mathbb{R}^d \rightarrow \overline{\mathbb{R}}$ epigraphically converges (epi-converges)¹ to f , written
 135 $f_n \xrightarrow[n \rightarrow +\infty]{e} f$, if and only if

$$136 \quad (i) \forall z, \forall z_n \rightarrow z : \liminf f_n(z_n) \geq f(z), \quad (ii) \forall z \exists z_n \rightarrow z : \limsup f_n(z_n) \leq f(z).$$

137 **2.2. Maximum Entropy on the Mean Problem.** For basic concepts of measure and probability,
 138 we follow most closely the standard text of Billingsley [6, Chapter 2]. Globally in this work, μ will
 139 be a Borel probability measure defined on compact $\mathcal{X} \subset \mathbb{R}^d$. Precisely, we work on the probability
 140 space $(\mathcal{X}, \mathcal{B}_{\mathcal{X}}, \mu)$, where $\mathcal{X} \subset \mathbb{R}^d$ is compact and $\mathcal{B}_{\mathcal{X}} = \{B \cap \mathcal{X} : B \in \mathbb{B}_d\}$ where $\mathbb{B}_d$ is the σ -
 141 algebra induced by the open sets in $\mathbb{R}^d$.² We will denote the set of all probability measures on the
 142 measurable space $(\mathcal{X}, \mathcal{B}_{\mathcal{X}})$ as $\mathcal{P}(\mathcal{X})$, and refer to elements of $\mathcal{P}(\mathcal{X})$ as probability measures on $\mathcal{X}$,
 143 with the implicit understanding that these are always Borel measures. For $Q, \mu \in \mathcal{P}(\mathcal{X})$, we say Q
 144 is absolutely continuous with respect to μ (and write $Q \ll \mu$) if for all $A \in \mathcal{B}_{\mathcal{X}}$ with $\mu(A) = 0$, then
 145 $Q(A) = 0$. [6, p. 422]. For $Q \ll \mu$, the Radon-Nikodym derivative of Q with respect to μ is defined
 146 as the (a.e.) unique function $\frac{dQ}{d\mu}$ with the property $Q(A) = \int_A \frac{dQ}{d\mu} d\mu$ for $A \in \mathcal{B}_{\mathcal{X}}$ [6, Theorem 32.3].

147 The Kullback-Leibler (KL) divergence [28] of $Q \in \mathcal{P}(\mathcal{X})$ with respect to $\mu \in \mathcal{P}(\mathcal{X})$ is defined as
 148

$$149 \quad (2.2) \quad \text{KL}(Q \parallel \mu) := \begin{cases} \int_{\mathcal{X}} \log(\frac{dQ}{d\mu}) d\mu, & Q \ll \mu, \\ +\infty, & \text{otherwise.} \end{cases}$$

150 For $\mu \in \mathcal{P}(\mathcal{X})$, the expected value $\mathbb{E}_{\mu} \in \mathbb{R}^d$ and moment generating function $M_{\mu} : \mathbb{R}^d \rightarrow \mathbb{R}$ function
 151 of μ are defined as [6, Ch.21]

$$152 \quad \mathbb{E}_{\mu} := \int_{\mathcal{X}} x d\mu(x), \quad M_{\mu}(y) := \int_{\mathcal{X}} \exp\langle y, x \rangle d\mu(x),$$

153 respectively. The log-moment generating function of μ is defined as

$$154 \quad L_{\mu}(y) := \log M_{\mu}(y) = \log \int_{\mathcal{X}} \exp\langle y, x \rangle d\mu(x).$$

155 As $\mathcal{X}$ is bounded, M_{μ} is finite-valued everywhere. By standard properties of moment generating
 156 functions (see e.g. [41, Theorem 4.8]) it is then analytic everywhere, and in turn so is L_{μ} .

157 Given $\mu \in \mathcal{P}(\mathcal{X})$, the Maximum Entropy on the Mean (MEM) function [44] $\kappa_{\mu} : \mathbb{R}^d \rightarrow \overline{\mathbb{R}}$ is

$$158 \quad \kappa_{\mu}(y) := \inf\{\text{KL}(Q \parallel \mu) : \mathbb{E}_Q = y, Q \in \mathcal{P}(\mathcal{X})\}.$$

159 The functions κ_{μ} and L_{μ} are paired in duality in a way that is fundamental to this work. We
 160 will flesh out this connection, as well as give additional properties of κ_{μ} for our setting; a Borel
 161 probability measure μ on compact $\mathcal{X}$. A detailed discussion of this connection under more general
 162 assumptions is the subject of [44].

163 For any $\mu \in \mathcal{P}(\mathcal{X})$ we have a vacuous tail-decay condition of the following form: for any $\sigma > 0$,

$$164 \quad \int_{\mathcal{X}} e^{\sigma \|x\|} d\mu(x) \leq \max_{x \in \mathcal{X}} \|x\| e^{\sigma \|x\|} < +\infty.$$

165 Consequently, by [16, Theorem 5.2 (iv)]³ we have that

$$166 \quad \kappa_{\mu}(x) = \sup_{y \in \mathbb{R}^d} \left[\langle y, x \rangle - \log \int_{\mathcal{X}} e^{\langle y, x \rangle} d\mu(x) \right] (= L_{\mu}^*(x)).$$

¹This is one of many equivalent conditions that characterize epi-convergence, see e.g. [37, Proposition 7.2].

²Equivalently, we could work with a Borel measure μ on $\mathbb{R}^d$ with support contained in $\mathcal{X}$.

³Applied to μ considered as a measure over $\mathbb{R}^d$ with support in $\mathcal{X}$.

Note that the conjugate L_μ^* is known in the large deviations literature as the (Cramér) rate function. For a more full development and alternative derivations of this conjugacy we refer to [29, 31].

Returning to the setting of interest with our standing assumption that $\mathcal{X}$ is compact, $\kappa_\mu = L_\mu^*$. This directly implies the following properties of κ_μ : (i) As L_μ is proper, lsc, and convex, so is its conjugate $L_\mu^* = \kappa_\mu$. (ii) Reiterating that L_μ is proper, lsc, convex, we may assert $(L_\mu^*)^* = L_\mu$ via Fenchel-Moreau ([40, Theorem 5.23]), and hence $\kappa_\mu^* = L_\mu$. (iii) As $\text{dom}(L_\mu) = \mathbb{R}^d$ we have that κ_μ is supercoercive [37, Theorem 11.8 (d)]. (iv) Recalling that L_μ is everywhere differentiable, κ_μ is strictly convex on every convex subset of $\text{dom}(\partial\kappa_\mu)$, which is also referred to as essentially strictly convex [36, p. 253].

With these preliminary notions, we can (re-)state the problem of interest in full detail. We work with images represented as vectors in $\mathbb{R}^d$, where d is the number of pixels. Given observed image $b \in \mathbb{R}^m$ which may be blurred and noisy, and known matrix $C \in \mathbb{R}^{m \times d}$, we wish to recover the ground truth $\hat{x}$ from the linear inverse problem $b = C\hat{x} + \eta$, where $\eta \sim \mathcal{Z}$ is an unknown noise vector in $\mathbb{R}^m$ drawn from noise distribution $\mathcal{Z}$. We remark that, in practice, it is usually the case that $m = d$ and C is invertible, but this is not necessary from a theoretical perspective. We assume the ground truth $\hat{x}$ is the expectation of an underlying image distribution - a Borel probability measure - μ on compact set $\mathcal{X} \subset \mathbb{R}^d$. Our best guess of $\hat{x}$ is then obtained by solving

$$(P) \quad \bar{x}_\mu = \underset{x \in \mathbb{R}^d}{\operatorname{argmin}} \alpha g(Cx) + \kappa_\mu(x).$$

where $g = g_b$ is a proper, lsc, convex function which may depend on b and serves as a fidelity term, and $\alpha > 0$ a parameter. For example, if $g = \frac{1}{2}\|b - (\cdot)\|_2^2$ one recovers the so-called reformulated MEM problem, first seen in [29].

Lemma 2.1. *For any lsc, proper, convex g , the primal problem (P) always has a solution.*

Proof. By the global assumption of compactness of $\mathcal{X}$, we have κ_μ is proper, lsc, convex and supercoercive, following the discussion above. As $g \circ C$ and κ_μ are convex, so is $\alpha g \circ C + \kappa_\mu$ for $\alpha > 0$. Further as both $\alpha g \circ C$ and κ_μ are proper and lsc, and κ_μ is supercoercive, the summation $\alpha g \circ C + \kappa_\mu$ is supercoercive, [37, Exercise 3.29, Lemma 3.27]. A supercoercive function is, in particular, level-bounded, so by [37, Theorem 1.9] the solution set $\operatorname{argmin}(\alpha g \circ C + \kappa_\mu)$ is nonempty. ■

We make one restriction on the choice of g , which will hold globally in this work:

Assumption 2.2. $0 \in \operatorname{int}(\operatorname{dom}(g) - C \operatorname{dom}(\kappa_\mu))$.

We remark that this property holds vacuously whenever g is finite-valued, e.g., $g = \frac{1}{2}\|b - (\cdot)\|_2^2$.

Instead of solving (P) directly, we use a dual approach. As $\kappa_\mu^* = L_\mu$ (by compactness of $\mathcal{X}$), the primal problem (P) has Fenchel dual (e.g., [5, Definition 15.19]) given by

$$(D) \quad (\arg)\min_{z \in \mathbb{R}^m} \alpha g^*(-z/\alpha) + L_\mu(C^T z).$$

We will hereafter denote the dual objective associated with $\mu \in \mathcal{P}(\mathcal{X})$ as

$$(2.3) \quad \phi_\mu(z) := \alpha g^*(-z/\alpha) + L_\mu(C^T z).$$

We remark that our sign convention and use of minimization in the dual agrees with [5], but the dual problem appears elsewhere in the literature as $\max -\phi_\mu(z)$, see e.g. [48, Corollary 2.8.5]. We record the following result which highlights the significance of Assumption 2.2 to our study.

Theorem 2.3. *The following are equivalent:*

(i) Assumption 2.2 holds; (ii) $\operatorname{argmin} \phi_\mu$ is nonempty and compact; (iii) ϕ_μ is level-coercive. In particular, under Assumption 2.2, the primal problem (P) has a unique solution given by

$$(2.4) \quad \bar{x}_\mu = \nabla L_\mu(C^T \bar{z}),$$

where $\bar{z} \in \operatorname{argmin} \phi_\mu$ is any solution of the dual problem (D).

212 *Proof.* As ϕ_μ is proper, convex and lsc, the equivalence of (i)-(iii) is exactly [4, Proposition
 213 3.1.3 (a),(c),(d)] as $\text{int}(\text{dom}(\phi_\mu^*)) = \text{int}(C \text{dom}(\kappa_\mu) - \text{dom}(g))$. Furthermore [4, Theorem 5.2.1]⁴,
 214 yields the primal-dual recovery formula (2.4) using the differentiability of L_μ . ■

215 **2.3. Approximate and Empirical Priors, Random Functions, and Epi-consistency.** If one has
 216 access to the true underlying image distribution μ , then the solution recipe is complete: solve (D)
 217 and use the primal-dual recovery formula (2.4) to find a solution to (P). But in practical situations,
 218 such as the imaging problems of interest here, it is unreasonable to assume full knowledge of μ and
 219 instead one models μ from domain specific knowledge or data set e.g. the discussions of [14, 42].
 220 That is, one specifies a prior $\nu \in \mathcal{P}(\mathcal{X})$ with $\nu \approx \mu$, and solves the approximate dual problem

$$221 \quad (2.5) \quad \min_{z \in \mathbb{R}^m} \phi_\nu(z).$$

222 Given $\varepsilon > 0$ and any ε -solution to (2.5), i.e. given any $z_{\nu,\varepsilon} \in S_\varepsilon(\nu)$, we define

$$223 \quad (2.6) \quad \bar{x}_{\nu,\varepsilon} := \nabla L_\nu(C^T z_{\nu,\varepsilon}),$$

224 with the hope, inspired by the recovery formula (2.4), that with a “reasonable” choice of $\nu \approx \mu$,
 225 and small ε , then also $\bar{x}_{\nu,\varepsilon} \approx \bar{x}_\mu$. The remainder of this work is dedicated to formalizing how well
 226 $\bar{x}_{\nu,\varepsilon}$ approximates $\bar{x}_\mu$ under various assumptions on g and ν .

227 A natural first approach is to construct ν from sample data. Let $(\Omega, \mathcal{F}, \mathbb{P})$ be a probability
 228 space. We model image samples as i.i.d. $\mathcal{X}$ -valued random variables $\{X_1, \dots, X_n, \dots\}$ with shared
 229 law $\mu := \mathbb{P} X_1^{-1}$. That is, each $X_i : \Omega \rightarrow \mathcal{X}$ is an $(\Omega, \mathcal{F}) \rightarrow (\mathcal{X}, \mathcal{B}_\mathcal{X})$ measurable function with the
 230 property that $\mu(B) = \mathbb{P}(\omega \in \Omega : X_1(\omega) \in B)$, for any $B \in \mathcal{B}_\mathcal{X}$. In particular, the law μ is by
 231 construction a Borel probability measure on $\mathcal{X}$. Intuitively, a random sample of n images is a given
 232 sequence of realizations $\{X_1(\omega), \dots, X_n(\omega), \dots\}$, from which we take only the first n vectors. In
 233 practice, such a sequence could arise from sampling a fixed dataset uniformly at random, in which
 234 case n is dictated by the amount of data that is available and is feasible to compute with. We then
 235 approximate μ via the empirical measure

$$236 \quad \mu_n^{(\omega)} := \frac{1}{n} \sum_{i=1}^n \delta_{X_i(\omega)}.$$

237 With this choice of $\nu = \mu_n^{(\omega)}$, we have the approximate dual problem

$$238 \quad (2.7) \quad \min_{z \in \mathbb{R}^m} \phi_{\mu_n^{(\omega)}}(z) \quad \text{with} \quad \phi_{\mu_n^{(\omega)}}(z) = \alpha g^* \left(\frac{-z}{\alpha} \right) + \log \frac{1}{n} \sum_{i=1}^n e^{\langle C^T z, X_i(\omega) \rangle}.$$

239 And exactly analogous to (2.6), given an ε -solution $\bar{z}_{n,\varepsilon}(\omega)$ of (2.7), we define

$$240 \quad (2.8) \quad \bar{x}_{n,\varepsilon}(\omega) := \nabla L_{\mu_n^{(\omega)}}(C^T \bar{z}_{n,\varepsilon}(\omega)) = \frac{\sum_{i=1}^n C X_i(\omega) e^{\langle C^T \bar{z}_{n,\varepsilon}(\omega), X_i(\omega) \rangle}}{\sum_{i=1}^n e^{\langle C^T \bar{z}_{n,\varepsilon}(\omega), X_i(\omega) \rangle}}.$$

241 Clearly, while the measure $\mu_n^{(\omega)}$ is well-defined and Borel for any given ω , the convergence
 242 properties of $\bar{z}_{n,\varepsilon}(\omega)$ and $\bar{x}_{n,\varepsilon}(\omega)$ should be studied in a stochastic sense over Ω . To this end, we
 243 leverage a probabilistic version of epi-convergence for random functions known as epi-consistency
 244 [26].

245 Let $(T, \mathcal{A})$ be a measurable space. A function $f : \mathbb{R}^m \times T \rightarrow \overline{\mathbb{R}}$ is called a random⁵ lsc function
 246 (with respect to $(T, \mathcal{A})$) [40, Definition 8.50] if the (set-valued) map $S_f : T \rightrightarrows \mathbb{R}^{m+1}$, $S_f(t) =$
 247 $\text{epi } f(\cdot, t)$ is closed-valued and measurable in the sense $S_f^{-1}(O) = \{t \in T : S_f(x) \cap O \neq \emptyset\} \in \mathcal{A}$.

⁴Note that there is a sign error in equation (5.3) in the reference.

⁵The inclusion of the word ‘random’ in this definition need not imply a priori any relation to a random process; we simply require measurability properties of f . Random lsc functions are also known as normal integrands in the literature, see [37, Chapter 14].

Our study is fundamentally interested in random lsc functions on $(\Omega, \mathcal{F})$, in service of proving convergence results for $\bar{x}_{n,\varepsilon}(\omega)$. But we emphasize that random lsc functions with respect to $(\Omega, \mathcal{F})$ are tightly linked with random lsc functions on $(X, \mathcal{B}_X)$. Specifically, if $X : \Omega \rightarrow \mathcal{X}$ is a random variable and $f : \mathbb{R}^m \times \mathcal{X} \rightarrow \bar{\mathbb{R}}$ is a random lsc function with respect to $(\mathcal{X}, \mathcal{B}_X)$, then the composition $f(\cdot, X(\cdot)) : \mathbb{R}^m \times \Omega \rightarrow \bar{\mathbb{R}}$ is a random lsc function with respect to the measurable space $(\Omega, \mathcal{F})$, see e.g. [37, Proposition 14.45 (c)] or the discussion of [39, Section 5]. This link will prove computationally convenient in the next section.

While the definition of a random lsc function is unwieldy to work with directly, it is implied by a host of easy to verify conditions [40, Example 8.51]. We will foremost use the following one: Let $(T, \mathcal{A})$ be a measurable space. If a function $f : \mathbb{R}^m \times T \rightarrow \bar{\mathbb{R}}$ is finite valued, with $f(\cdot, t)$ continuous for all t , and $f(z, \cdot)$ measurable for all z , we say f is a Carathéodory function. Any function which is Carathéodory is random lsc [38, Example 14.26].

Immediately, we can assert $\phi_{\mu_n^{(\cdot)}}$ is a random lsc function from $\mathbb{R}^d \times \Omega \rightarrow \bar{\mathbb{R}}$, as it is Carathéodory. In particular, by [37, Theorem 14.37] or [39, Section 5], the ε -solution mappings

$$\omega \mapsto \left\{ z : \phi_{\mu_n^{(\omega)}}(z) \leq \inf \phi_{\mu_n^{(\omega)}} + \varepsilon \right\}$$

are measurable (in the set valued sense defined above), and for all $\varepsilon \geq 0$ there always exists a $\mathbb{P}$ -measurable selection $z_{n,\varepsilon}(\omega) \in S_\varepsilon(\mu_n^{(\omega)})$ [37, Theorem 14.37, Theorem 14.33].

We conclude with the definition of epi-consistency as seen in [26, p. 86]; a sequence of random lsc functions $h_n : \mathbb{R}^m \times \Omega \rightarrow \bar{\mathbb{R}}$ is said to be epi-consistent with limit function $h : \mathbb{R}^m \rightarrow \bar{\mathbb{R}}$ if

$$(2.9) \quad \mathbb{P} \left(\left\{ \omega \in \Omega \mid h_n(\cdot, \omega) \xrightarrow[n \rightarrow +\infty]{e} h \right\} \right) = 1.$$

3. Epigraphical convergence and convergence of minimizers. The goal of this section is to prove convergence of minimizers in the empirical case, i.e., that $\bar{x}_{n,\varepsilon}(\omega)$ as defined in (2.8) converges to $\bar{x}_\mu$, the solution of (P), for $\mathbb{P}$ -almost every $\omega \in \Omega$ as $\varepsilon \searrow 0$. To do so, we prove empirical approximations of the moment generating function are epi-consistent with M_μ , and leverage this to prove epi-consistency of $\phi_{\mu_n^{(\omega)}}$ with limit ϕ_μ . Via classic convex analysis techniques, this guarantees the desired convergence of minimizers with probability one.

3.1. Epi-consistency of the empirical moment generating functions. Given $\{X_1, \dots, X_n, \dots\}$ i.i.d. with shared law $\mu = \mathbb{P} X_1^{-1} \in \mathcal{P}(\mathcal{X})$, we denote the moment generating function of $\mu_n^{(\omega)}$ as $M_n(y, \omega) := \frac{1}{n} \sum_{i=1}^n e^{\langle y, X_i(\omega) \rangle}$. Define $f : \mathbb{R}^m \times \mathbb{R}^d \rightarrow \bar{\mathbb{R}}$ as $f(z, x) = e^{\langle C^T z, x \rangle}$. Then

$$\begin{aligned} M_\mu(C^T z) &= \int_{\mathcal{X}} e^{\langle C^T z, \cdot \rangle} d\mu = \int_{\mathcal{X}} f(z, \cdot) d\mu, \\ M_n(C^T z, \omega) &= \frac{1}{n} \sum_{i=1}^n e^{\langle C^T z, X_i(\omega) \rangle} = \frac{1}{n} \sum_{i=1}^n f(z, X_i(\omega)). \end{aligned}$$

This explicit decomposition is useful to apply a specialized version of the main theorem of King and Wets [26, Theorem 2], which we restate without proof.

Proposition 3.1. *Let $f : \mathbb{R}^m \times \mathcal{X} \rightarrow \bar{\mathbb{R}}$ be a random lsc function such that $f(\cdot, x)$ is convex and differentiable for all x . Let $X_1, \dots, X_n$ be i.i.d. $\mathcal{X}$ -valued random variables on $(\Omega, \mathcal{F}, \mathbb{P})$ with shared law $\mu \in \mathcal{P}(\mathcal{X})$. If there exists $\bar{z} \in \mathbb{R}^m$ such that*

$$\int_{\mathcal{X}} f(\bar{z}, \cdot) d\mu < +\infty, \quad \text{and} \quad \int_{\mathcal{X}} \|\nabla_z f(\bar{z}, \cdot)\| d\mu < +\infty,$$

then the sequence of (random lsc) functions $S_n : \mathbb{R}^m \times \Omega \rightarrow \bar{\mathbb{R}}$ given by

$$S_n(z, \omega) := \frac{1}{n} \sum_{i=1}^n f(z, X_i(\omega))$$

287 is epi-consistent with limit $S_\mu : z \mapsto \int_{\mathcal{X}} f(z, \cdot) d\mu$, which is proper, convex, and lsc.

288 Via a direct application of the above we have the following.

289 **Corollary 3.2.** *The sequence $M_n(C^T(\cdot), \cdot)$ is epi-consistent with limit $M_\mu \circ C^T$.*

290 *Proof.* Define $f(z, x) = e^{\langle C^T z, x \rangle}$. For any x , $\langle C^T(\cdot), x \rangle$ is a linear function, and $e^{(\cdot)}$ is convex -
 291 giving that the composition $f(\cdot, x)$ is convex. As f is differentiable (hence continuous) in z for fixed
 292 x and vice-versa, it is Carathéodory and thus a random lsc function (with respect to $(\mathcal{X}, \mathcal{B}_{\mathcal{X}})$).

293 Next we claim $\bar{z} = 0$ satisfies the conditions of the proposition. First, by direct computation

$$294 \quad \int_{\mathcal{X}} e^{\langle 0, x \rangle} d\mu(x) = \int_{\mathcal{X}} d\mu(x) = 1 < +\infty$$

295 as μ is a probability measure on $\mathcal{X}$. As $f(\cdot, x)$ is differentiable, we can compute $\nabla_z f(\bar{z}, x) =$
 296 $Cx e^{\langle C^T 0, x \rangle} = Cx$. Hence

$$297 \quad \int_{\mathcal{X}} \|\nabla_z f(\bar{z}, x)\| d\mu(x) = \int_{\mathcal{X}} \|Cx\| d\mu(x) \leq \|C\| \max_{x \in \mathcal{X}} \|x\| < +\infty,$$

298 where we have used the boundedness of $\mathcal{X}$, and once again that μ is a probability measure. Thus
 299 we satisfy the assumptions of [Proposition 3.1](#), and can conclude that the sequence of random
 300 lsc functions S_n given by $S_n(z, \omega) = \frac{1}{n} \sum_{i=1}^n f(z, X_i(\omega))$ are epi-consistent with limit $S_\mu : z \mapsto$
 301 $\int_{\mathcal{X}} f(z, \cdot) d\mu$. But,

$$302 \quad S_n(z, \omega) = \frac{1}{n} \sum_{i=1}^n e^{\langle C^T z, X_i(\omega) \rangle} = M_n(C^T z, \omega) \quad \text{and} \quad S_\mu(z) = \int_{\mathcal{X}} e^{\langle C^T z, \cdot \rangle} d\mu = M_\mu(C^T z),$$

303 and so we have shown the sequence $M_n(C^T(\cdot), \cdot)$ is epi-consistent with limit $M_\mu \circ C^T$. ■

304 **Corollary 3.3.** *The sequence $L_{\mu_n^{(\omega)}} \circ C^T$ is epi-consistent with limit $L_\mu \circ C^T$.*

305 *Proof.* Let

$$306 \quad \Omega_e = \left\{ \omega \in \Omega \mid M_n(C^T(\cdot), \omega) \xrightarrow[n \rightarrow +\infty]{e} M_\mu \circ C^T(\cdot) \right\},$$

307 which has $\mathbb{P}(\Omega_e) = 1$ by [Corollary 3.2](#), and let $\omega \in \Omega_e$. Both M_n and M_μ are finite valued and
 308 strictly positive, and furthermore the function $\log : \mathbb{R}_{++} \rightarrow \mathbb{R}$ is continuous and increasing. Hence,
 309 by a simple extension of [\[36, Exercise 7.8\(c\)\]](#), it follows, for all $\omega \in \Omega_e$, that

$$310 \quad L_{\mu_n^{(\omega)}} \circ C^T = \log M_n(C^T(\cdot), \omega) \xrightarrow[n \rightarrow +\infty]{e} \log M_\mu \circ C^T = L_\mu \circ C^T. \quad \text{■}$$

311 **3.2. Epi-consistency of the dual objective functions.** We now use the previous lemma to
 312 obtain an epi-consistency result for the entire empirical dual objective function. This is not an
 313 immediately clear, as epi-convergence is not generally preserved by even simple operations such as
 314 addition, see, e.g., the discussion in [\[37, p. 276\]](#) and the note [\[8\]](#) that eludes to subtle difficulties
 315 when dealing with extended real-valued arithmetic in this context.

316 We recall the following pointwise convergence result for compact $\mathcal{X}$, which is classical in the
 317 statistics literature.

319 **Lemma 3.4.** *If $\mu \in \mathcal{P}(\mathcal{X})$, for almost every $\omega \in \Omega$, and all $z \in \mathbb{R}^m$*

$$320 \quad M_n(C^T z, \omega) \rightarrow M_\mu \circ C^T(z),$$

321 *namely pointwise convergence in z .*

We remark that the literature contains stronger uniform convergence results, observed first in Csörgö [?] without proof, and later proven in [18] and [12, Proposition 1]. Noting that both $M_n(z, \omega), M_\mu(z) > 0$ are strictly positive for all $z \in \mathbb{R}^m$, and that the logarithm is continuous on the strictly positive real line, we have an immediate corollary:

Corollary 3.5. *For almost every $\omega \in \Omega$, for all $z \in \mathbb{R}^m$*

$$L_{\mu_n^{(\omega)}}(C^T z) = \log M_n(C^T z, \omega) \rightarrow \log M_\mu(C^T z) = L_\mu(C^T z).$$

Using this we prove the first main result:

Theorem 3.6. *For any lsc, proper, convex function g , the empirical dual objective function $\phi_{\mu_n^{(\omega)}}$ is epi-consistent with limit ϕ_μ*

Proof. Define

$$\Omega_e = \left\{ \omega \in \Omega \mid L_{\mu_n^{(\omega)}} \circ C^T(\cdot) \xrightarrow[n \rightarrow +\infty]{e} L_\mu \circ C^T(\cdot) \right\}.$$

By Corollary 3.3, $\mathbb{P}(\Omega_e) = 1$. Similarly denote

$$\Omega_p = \left\{ \omega \in \Omega \mid L_{\mu_n^{(\omega)}} \circ C^T(\cdot) \rightarrow L_\mu \circ C^T(\cdot) \text{ pointwise} \right\}.$$

By Corollary 3.5, we also have $\mathbb{P}(\Omega_p) = 1$. In particular we observe that $\mathbb{P}(\Omega_e \cap \Omega_p) = 1$.

On the other hand we have vacuously that the constant sequence of convex, proper, lsc functions $\alpha g^* \circ (-\text{Id}/\alpha)$ converges to $\alpha g^* \circ (-\text{Id}/\alpha)$ both epigraphically and pointwise.

Thus for any fixed $\omega \in \Omega_p \cap \Omega_e$ we have constructed two sequences, namely $g_n \equiv \alpha g^* \circ (-\text{Id}/\alpha)$ and $L_n = L_{\mu_n^{(\omega)}} \circ C^T$, which both converge epigraphically and pointwise for all $\omega \in \Omega_e \cap \Omega_p$. Therefore, by [37, Theorem 7.46(a)], for all $\omega \in \Omega_e \cap \Omega_p$

$$\alpha g^* \circ (-\text{Id}/\alpha) + L_{\mu_n^{(\omega)}} \circ C^T \xrightarrow[n \rightarrow +\infty]{e} \alpha g^* \circ (-\text{Id}/\alpha) + L_\mu \circ C^T.$$

As $\mathbb{P}(\Omega_e \cap \Omega_p) = 1$, this proves the result. ■

3.3. Convergence of minimizers. We now use epi-consistency to prove convergence of minimizers. At the dual level this can be summarized in the following lemma, essentially [26, Proposition 2.2]; which was stated therein without proof.⁶

Lemma 3.7. *There exists a subset $\Xi \subset \Omega$ of measure one, such that for any $\omega \in \Xi$ we have: Let $\{\varepsilon_n\} \searrow 0$ and $z_n(\omega)$ such that*

$$\phi_{\mu_n^{(\omega)}}(z_n(\omega)) \leq \inf_z \phi_{\mu_n^{(\omega)}}(z) + \varepsilon_n.$$

Let $\{z_{n_k}(\omega)\}$ be any convergent subsequence of $\{z_n(\omega)\}$. Then $\lim_{k \rightarrow +\infty} z_{n_k}(\omega)$ is a minimizer of ϕ_μ . If ϕ_μ admits a unique minimizer $\bar{z}_\mu$, then $z_n \rightarrow \bar{z}_\mu$.

Proof. Denote

$$\Xi = \left\{ \omega \in \Omega \mid \phi_{\mu_n^{(\omega)}} \xrightarrow[n \rightarrow +\infty]{e} \phi_\mu \right\}.$$

By Theorem 3.6, $\mathbb{P}(\Xi) = 1$. Fix any $\omega \in \Xi$.

⁶We remark that (as observed in [26]) epigraphical convergence of a (multi-)function depending on a parameter (such as ω) guarantees convergence of minimizers in much broader contexts, see e.g. [3, Theorem 1.10] or [38, Theorem 3.22]. Here we include a first principles proof.

As we have fixed $\omega \in \Xi$, we have that the sequence $\phi_{\mu_n}^{(\omega)} \xrightarrow[n \rightarrow +\infty]{e} \phi_\mu$ epi-converges. Also, by Theorem 2.3, our global Assumption 2.2 holds if and only if ϕ_μ is level-bounded. These two observations together imply by [37, Theorem 7.32 (c)] that the sequence $\phi_{\mu_n}^{(\omega)}$ is eventually level-bounded.⁷ Altogether, this means the sequence of lsc, proper, eventually level-bounded functions $\phi_{\mu_n}^{(\omega)}$ epi-converge to ϕ_μ - which is also lsc and proper. This set of properties is precisely the necessary assumptions of [37, Theorem 7.33], which then asserts any sequence of approximate minimizers $\{z_n(\omega)\}$ is bounded with all cluster points belonging to $\operatorname{argmin} \phi_\mu$. Namely, any convergent subsequence $\{z_{n_k}(\omega)\}$ has the property that its limit $\lim_{k \rightarrow +\infty} z_{n_k} \in \operatorname{argmin} \phi_\mu$. Lastly, if we also have $\operatorname{argmin} \phi_\mu = \{\bar{z}_\mu\}$, then from the same result [37, Theorem 7.33], then necessarily $z_n(\omega) \rightarrow \bar{z}_\mu$. ■

We now push this convergence to the primal level by using, in essence, Attouch's Theorem [2], [3, Theorem 3.66], in the form of a corollary of Rockafellar and Wets [37, Theorem 12.40].

Lemma 3.8. *Let $\hat{z} \in \mathbb{R}^m$, and let $z_n \rightarrow \hat{z}$ be any sequence converging to $\hat{z}$. Then for almost every ω ,*

$$\lim_{n \rightarrow +\infty} \nabla L_{\mu_n}^{(\omega)}(C^T z_n) = \nabla L_\mu(C^T \hat{z}).$$

Proof. We first observe that $\operatorname{dom}(L_\mu \circ C^T) = \mathbb{R}^m$ so that $\hat{z} \in \operatorname{int}(\operatorname{dom}(L_\mu \circ C^T))$. Also as M_μ is everywhere finite-valued, $L_\mu(C^T \hat{z}) = \log M_\mu(C^T \hat{z}) < +\infty$. Furthermore for all n , the function $L_{\mu_n}^{(\omega)} \circ C^T$ is proper, convex, and differentiable. Finally, we have shown in Corollary 3.3, that for almost every $\omega \in \Omega$, we have $L_{\mu_n}^{(\omega)} \circ C^T \xrightarrow[n \rightarrow +\infty]{e} L_\mu \circ C^T$. These conditions together are the necessary assumptions of [37, Theorem 12.40 (b)]. Hence we have convergence $\lim_{n \rightarrow +\infty} \nabla L_{\mu_n}^{(\omega)}(C^T z_n) = \nabla L_\mu(C^T \hat{z})$ for almost every $\omega \in \Omega$. ■

We now prove the main result.

Theorem 3.9. *There exists a set $\Xi \subseteq \Omega$ of probability one such that for each $\omega \in \Xi$ the following holds: Given $\varepsilon_n \searrow 0$, and $z_n(\omega)$ such that $\phi_{\mu_n}^{(\omega)}(z_n(\omega)) \leq \inf_z \phi_{\mu_n}^{(\omega)}(z) + \varepsilon_n$, define*

$$x_n(\omega) := \nabla L_{\mu_n}^{(\omega)}(C^T z_n).$$

If $z_{n_k}(\omega)$ is any convergent subsequence of $z_n(\omega)$ then $\lim_{k \rightarrow +\infty} x_{n_k}(\omega) = \bar{x}_\mu$, where $\bar{x}_\mu$ is the unique solution of (P). If (2.7) admits a unique solution $\bar{z}_\mu$, then in fact $x_n(\omega) \rightarrow \bar{x}_\mu$.

Proof. Let

$$\Xi = \{\omega \in \Omega \mid \phi_{\mu_n}^{(\omega)} \xrightarrow[n \rightarrow +\infty]{e} \phi_\mu\},$$

recalling that by Proposition 3.1, $\mathbb{P}(\Xi) = 1$. Fix $\omega \in \Xi$. By Lemma 3.7, for any convergent subsequence $z_{n_k}(\omega)$ with limit $\bar{z}(\omega)$, we have that $\bar{z}(\omega) \in \operatorname{argmin} \phi_\mu$. Furthermore, by Lemma 3.8

$$\lim_{k \rightarrow +\infty} x_{n_k}(\omega) = \lim_{k \rightarrow +\infty} \nabla L_{\mu_{n_k}}^{(\omega)}(C^T z_{n_k}) = \nabla L_\mu(C^T \bar{z}(\omega))$$

Using the primal-dual optimality conditions (2.4) we have that $\nabla L_\mu(C^T \bar{z}(\omega))$ solves the primal problem (P). As (P) admits a unique solution $\bar{x}_\mu$, necessarily $\lim_{k \rightarrow +\infty} x_{n_k}(\omega) = \bar{x}_\mu$. If additionally $\operatorname{argmin} \phi_\mu = \{\bar{z}_\mu\}$, then necessarily $z_n \rightarrow \bar{z}_\mu$ via Lemma 3.7, and the result follows from an identical application of Lemma 3.8 and (2.4). ■

The key novelty of this result is the almost sure convergence of minimizers, particularly as $\varepsilon \searrow 0$. The compact support of the measure μ is key for this result. In particular, this is stronger than the convergence in probability guaranteed by common statistical techniques such as m -estimation [46, Section 3.2].

⁷A sequence of functions $f_n : \mathbb{R}^d \rightarrow \overline{\mathbb{R}}$ is eventually level-bounded if for each α , the sequence of sets $\{f_n^{-1}([-\infty, \alpha])\}$ is eventually bounded, see [37, p. 266].

4. Convergence rates for quadratic fidelity.

When proving rates of convergence, we restrict ourselves to the case where $g = \frac{1}{2}\|b - (\cdot)\|_2^2$. Thus, the dual objective function reads

$$(4.1) \quad \phi_\mu(z) = \frac{1}{2\alpha}\|z\|^2 - \langle b, z \rangle + L_\mu(C^T z).$$

Clearly, ϕ_μ is finite valued and $(1/\alpha)$ -strongly convex, hence admits a unique minimizer $\bar{z}_\mu$. Recalling what was laid out in Subsection 2.2, as global Assumption 2.2 holds vacuously with $g = \frac{1}{2}\|b - (\cdot)\|_2^2$, the unique solution to the MEM primal problem (P) is given by $\bar{x}_\mu = \nabla L_\mu(C^T \bar{z}_\mu)$. Further by our global compactness assumption on $\mathcal{X}$, ϕ_μ is (infinitely many times) differentiable.

4.1. Epigraphical distances.

Our main tool to prove convergence rates are epigraphical distances. We mainly follow the presentation in Royset and Wets [40, Chapter 6.J], but one may find similar treatment in Rockafellar and Wets [37, Chapter 7]. For any norm $\|\cdot\|_*$ on $\mathbb{R}^d$, the distance (in said norm) between a point c and a set D is defined as $d_D(c) = \inf_{d \in D} \|c - d\|_*$. For C, D subsets of $\mathbb{R}^d$ we define the excess of C over D [40, p. 399] as

$$\text{exc}(C, D) := \begin{cases} \sup_{c \in C} d_D(c) & \text{if } C, D \neq \emptyset, \\ +\infty & \text{if } C \neq \emptyset, D = \emptyset, \\ 0 & \text{otherwise.} \end{cases}$$

We note that this excess explicitly depends on the choice of norm used to define d_D . For the specific case of the 2-norm, we denote the projection of a point $a \in \mathbb{R}^d$ onto a closed, convex set $B \subset \mathbb{R}^d$ as the unique point $\text{proj}_B(a) \in B$ which achieves the minimum $\|\text{proj}_B(a) - a\| = \min_{b \in B} \|b - a\|$. The truncated ρ -Hausdorff distance [40, p. 399] between two sets $C, D \subset \mathbb{R}^d$ is defined as

$$\hat{d}_\rho(C, D) := \max \{ \text{exc}(C \cap B_\rho, D), \text{exc}(D \cap B_\rho, C) \},$$

where $B_\rho = \{x \in \mathbb{R}^d : \|x\|_* \leq \rho\}$ is the closed ball of radius ρ in $\mathbb{R}^d$. When discussing distances on $\mathbb{R}^d$, we will consistently make the choice of $\|\cdot\|_* = \|\cdot\|$ the 2-norm. We note we can recover the usual Pompeiu-Hausdorff distance by taking $\rho \rightarrow +\infty$.

However, to extend the truncated ρ -distance to epigraphs of functions - which here are subsets of $\mathbb{R}^{d+1}$ - we equip $\mathbb{R}^{d+1}$ with a very particular norm. For any $z \in \mathbb{R}^{d+1}$, write $z = (x, a)$ for $x \in \mathbb{R}^d, a \in \mathbb{R}$. Then for any $z_1, z_2 \in \mathbb{R}^{d+1}$ we define the norm

$$\|z_1 - z_2\|_{*,d+1} = \|(x_1, a_1) - (x_2, a_2)\|_{*,d+1} := \max\{\|x_1 - x_2\|_2, |a_1 - a_2|\}.$$

With this norm, we can define an epi-distance as in [40, Equation 6.36]: for $f, h : \mathbb{R}^d \rightarrow \bar{\mathbb{R}}$ not identically $+\infty$, and $\rho > 0$ we define

$$(4.2) \quad \hat{d}_\rho(f, h) := \hat{d}_\rho(\text{epi } f, \text{epi } h),$$

where $\mathbb{R}^{d+1}$ has been equipped with the norm $\|\cdot\|_{*,d+1}$. This epi-distance quantifies epigraphical convergence in the following sense: [37, Theorem 7.58]⁸ if f is a proper function and f_n a sequence of proper functions, then for any constant $\rho_0 > 0$:

$$f_n \xrightarrow[n \rightarrow +\infty]{e} f \text{ if and only if } \hat{d}_\rho(f_n, f) \rightarrow 0 \text{ for all } \rho > \rho_0.$$

4.2. Convergence Rates.

We begin with a technical lemma which will prove expedient for future results.

⁸We remark that while at first glance the definition of epi-distance seen in [37, Theorem 7.58] differs from ours (which agrees with [40]), it is equivalent up to multiplication by a constant and rescaling in ρ . See [40, Proposition 6.58] and [37, Proposition 7.61] - the Kenmochi conditions - for details.

Lemma 4.1. Let $\rho > 0$ and $\nu \in \mathcal{P}(\mathcal{X})$. Then, for all $z \in B_\rho$ we have

$$M_\nu(C^T z) = \int_{\mathcal{X}} e^{\langle C^T z, \cdot \rangle} d\nu \in [\exp(-\rho \|C\| |\mathcal{X}|), \exp(\rho \|C\| |\mathcal{X}|)].$$

Proof. For all $x \in \mathcal{X}$, $z \in B_\rho$, we have, via Cauchy-Schwarz, that

$$\exp(-\rho \|C\| |\mathcal{X}|) \leq \exp(-\|z\| \|Cx\|) \leq \exp\langle C^T z, x \rangle.$$

In particular, $\exp(-\rho \|C\| |\mathcal{X}|) \leq \min_{x \in \mathcal{X}} \exp\langle C^T z, x \rangle$. On the other hand, we find that

$$\exp\langle C^T z, x \rangle \leq \exp(\|z\| \|Cx\|) \leq \exp(\rho \|C\| |\mathcal{X}|).$$

Thus, $\max_{x \in \mathcal{X}} \exp\langle C^T z, x \rangle \leq \exp(\rho \|C\| |\mathcal{X}|)$. Hence for any $\nu \in \mathcal{P}(\mathcal{X})$ we find

$$1 \cdot \exp(-\rho \|C\| |\mathcal{X}|) \leq \nu(\mathcal{X}) \min_{x \in \mathcal{X}} e^{\langle C^T z, x \rangle} \leq \int_{\mathcal{X}} e^{\langle C^T z, \cdot \rangle} d\nu \leq \nu(\mathcal{X}) \max_{x \in \mathcal{X}} e^{\langle C^T z, x \rangle} \leq 1 \cdot \exp(\rho \|C\| |\mathcal{X}|).$$

We now prove rates of convergence for arbitrary prior ν , and later specialize to the empirical case. To this end, we construct a key global constant ρ_0 induced by C, b, α in (4.1): We define

$$(4.3) \quad \rho_0 := \max \left\{ \hat{\rho}, \frac{\hat{\rho}^2}{2\alpha} + \|b\| \hat{\rho} + \hat{\rho} \|C\| |\mathcal{X}| \right\},$$

where $\hat{\rho} = 2\alpha(\|b\| + \|C\| |\mathcal{X}|)$ and $|\mathcal{X}| := \max_{x \in \mathcal{X}} \|x\|$. We emphasize that our running compactness assumption on $\mathcal{X}$ is essential for finiteness of ρ_0 . The main feature of this constant is the following.

Lemma 4.2. For any $\nu \in \mathcal{P}(\mathcal{X})$, let ϕ_ν be the corresponding dual objective function as defined in (4.1), which has a unique minimizer $\bar{z}_\nu$. Then ρ_0 has the following two properties:

$$(a) \quad \phi_\nu(\bar{z}_\nu) \in [-\rho_0, \rho_0], \quad (b) \quad \|\bar{z}_\nu\| \leq \rho_0.$$

Proof. We first claim that $\|\bar{z}_\nu\| \leq \hat{\rho}$. Let $z \in \mathbb{R}^d$ be such that $\|z\| > \hat{\rho}$. Then,

$$(4.4) \quad \phi_\nu(z) \geq \frac{\|z\|^2}{2\alpha} - \|b\| \|z\| + \log \exp(-\|C\| \|x\| \|z\|) = \|z\| \left(\frac{\|z\|}{2\alpha} - \|b\| - \|C\| |\mathcal{X}| \right).$$

Where in the first inequality we have used Cauchy-Schwarz on $\langle b, z \rangle$, and Lemma 4.1 with $\rho = \|z\|$ to bound $M_\mu(C^T z) \geq \exp(-\|C\| \|x\| \|z\|)$. From (4.4) it is clear that $\|z\| > \hat{\rho}$ implies $\phi_\nu(z) > 0$. But observing that $\phi_\nu(0) = 0$, such z cannot be a minimizer. Hence necessarily $\|\bar{z}_\nu\| \leq \hat{\rho} \leq \rho_0$. Once more, via Cauchy-Schwarz and Lemma 4.1 we compute

$$\begin{aligned} |\phi_\nu(\bar{z}_\nu)| &= \left| \frac{\|\bar{z}_\nu\|^2}{2\alpha} - \langle b, \bar{z}_\nu \rangle + L_\nu(\bar{z}_\nu) \right| \leq \frac{\hat{\rho}^2}{2\alpha} + \hat{\rho} \|b\| + \log \exp(\|C\| |\mathcal{X}| \hat{\rho}) \\ &= \frac{\hat{\rho}^2}{2\alpha} + \hat{\rho} \|b\| + \hat{\rho} \|C\| |\mathcal{X}|. \end{aligned}$$

Lemma 4.3. Let ρ_0 be given by (4.3). Then for all $\rho > \rho_0$ and all $\mu, \nu \in \mathcal{P}(\mathcal{X})$, we have

$$\hat{d}_\rho(\phi_\mu, \phi_\nu) \leq \max_{z \in B_\rho} |L_\nu(C^T z) - L_\mu(C^T z)|.$$

Proof. Lemma 4.2 guarantees that for both measures $\mu, \nu \in \mathcal{P}(\mathcal{X})$, we have

$$(4.5) \quad \phi_\nu(\bar{z}_\nu), \phi_\mu(\bar{z}_\mu) \in [-\rho_0, \rho_0] \quad \text{and} \quad \|\bar{z}_\nu\|, \|\bar{z}_\mu\| \leq \rho_0.$$

These conditions imply, for any $\rho > \rho_0$, that the set $C_\rho := (\{z : \phi_\mu(z) \leq \rho\} \cup \{z : \phi_\nu(z) \leq \rho\}) \cap B_\rho$ is nonempty. This follows from (4.5) as for any $\rho > \rho_0$ the nonempty set $\{\bar{z}_\mu, \bar{z}_\nu\} \cap B_{\rho_0}$ is contained in $C_{\rho_0} \subset C_\rho$. As C_ρ is nonempty we may apply [40, Theorem 6.59] with $f = \phi_\mu$ and $g = \phi_\nu$ to obtain $\hat{d}_\rho(\phi_\mu, \phi_\nu) \leq \sup_{z \in C_\rho} |\phi_\nu(z) - \phi_\mu(z)|$. Then from the definition of ϕ_μ and ϕ_ν , we have

$$\begin{aligned} \sup_{z \in C_\rho} |\phi_\nu(z) - \phi_\mu(z)| &= \sup_{z \in C_\rho} |L_\nu(C^T z) - L_\mu(C^T z)| \\ &\leq \sup_{z \in B_\rho} |L_\nu(C^T z) - L_\mu(C^T z)| \\ &= \max_{z \in B_\rho} |L_\nu(C^T z) - L_\mu(C^T z)|, \end{aligned}$$

where in the penultimate line uses that $C_\rho \subseteq B_\rho$, and the final equality follows as the continuous function $L_\mu \circ C^T - L_\nu \circ C^T$ achieves a maximum over the compact set B_ρ . ■

For notational convenience, we will hereafter denote

$$D_\rho(\nu, \mu) := \max_{z \in B_\rho} |L_\mu(C^T z) - L_\nu(C^T z)|.$$

We also recall from Subsection 2.1 that $S_\varepsilon(\nu)$ denotes the set of ε -minimizers of ϕ_ν .

Lemma 4.4. *Let ρ_0 be given by (4.3). Then, for all $\mu, \nu \in \mathcal{P}(\mathcal{X})$, all $\rho > \rho_0$, and all $\varepsilon \in [0, \rho - \rho_0]$, the following holds: If*

$$\delta > \varepsilon + 2D_\rho(\nu, \mu),$$

then

$$|\phi_\nu(\bar{z}_\nu) - \phi_\mu(\bar{z}_\mu)| \leq D_\rho(\nu, \mu) \quad \text{and} \quad \text{exc}(S_\varepsilon(\nu) \cap B_\rho, S_\delta(\mu)) \leq D_\rho(\nu, \mu).$$

Proof. Let $\rho > \rho_0$ and $\varepsilon \in [0, \rho - \rho_0]$. By choice, we have $\varepsilon < 2\rho$. By Lemma 4.2(a) we have $\phi_\nu(\bar{z}_\nu), \phi_\mu(\bar{z}_\mu) \in [-\rho_0, \rho_0]$, and in turn by the choice of ρ, ε we have $[-\rho_0, \rho_0] \subseteq [-\rho, \rho] \subseteq [-\rho, \rho - \varepsilon]$. Also, as $\rho > \rho_0$, by Lemma 4.2(b) we have $\{z_\nu\} = \text{argmin } \phi_\nu \cap B_\rho$ and $\{z_\mu\} = \text{argmin } \phi_\mu \cap B_\rho$. These properties of ρ, ε are exactly the assumptions of [40, Theorem 6.56] for $f = \phi_\mu$ and $g = \phi_\nu$. This result yields that, if $\delta > \varepsilon + 2\hat{d}_\rho(\phi_\mu, \phi_\nu)$, then $|\phi_\nu(\bar{z}_\nu) - \phi_\mu(\bar{z}_\mu)| \leq \hat{d}_\rho(\phi_\nu, \phi_\mu)$ and $\text{exc}(S_\varepsilon(\nu) \cap B_\rho, S_\delta(\mu)) \leq \hat{d}_\rho(\phi_\nu, \phi_\mu)$. However as $\rho > \rho_0$, we may apply Lemma 4.3 to assert $\hat{d}_\rho(\phi_\mu, \phi_\nu) \leq D_\rho(\nu, \mu)$. Hence, for any $\delta > \varepsilon + 2D_\rho(\nu, \mu) \geq \varepsilon + 2\hat{d}_\rho(\phi_\mu, \phi_\nu)$ we obtain

$$\begin{aligned} |\phi_\nu(\bar{z}_\nu) - \phi_\mu(\bar{z}_\mu)| &\leq D_\rho(\nu, \mu), \\ \text{exc}(S_\varepsilon(\nu) \cap B_\rho, S_\delta(\mu)) &\leq D_\rho(\nu, \mu). \end{aligned} \quad \blacksquare$$

For the main results, Theorem 4.7 and Theorem 5.5, we require additional auxiliary results. With some additional computation, we can infer the following Lipschitz bound on ∇L_ν .

Corollary 4.5. *Let $\hat{\rho} > 0$ and $\nu \in \mathcal{P}(\mathcal{X})$. Then for all $x, y \in B_{\hat{\rho}} \subset \mathbb{R}^d$, we have that*

$$(4.6) \quad \|\nabla L_\nu(x) - \nabla L_\nu(y)\| \leq K\|x - y\|$$

for an explicit constant $K > 0$ which depends on $\hat{\rho}, d, |\mathcal{X}|$, but not on ν .

Proof. As discussed in Subsection 2.2, L_ν is twice continuously differentiable. Hence, using the fundamental theorem of calculus, we have $\nabla L_\nu(x) - \nabla L_\nu(y) = \int_0^1 \nabla^2 L_\nu(x + t(y - x)) \cdot (y - x) dt$.

Thus, as $x + t(y - x) \in B_{\hat{\rho}}$ for all $t \in [0, 1]$, we have

$$\begin{aligned} \|\nabla L_{\nu}(z) - \nabla L_{\nu}(y)\| &\leq \int_0^1 \|\nabla^2 L_{\nu}(x + t(y - x))\| \|y - x\| dt \\ &\leq \int_0^1 \max_{z \in B_{\hat{\rho}}} \|\nabla^2 L_{\nu}(z)\| \|y - x\| dt \\ &= \max_{z \in B_{\hat{\rho}}} \|\nabla^2 L_{\nu}(z)\| \|y - x\|. \end{aligned}$$

By convexity of L_{ν} , we observe that $\nabla^2 L_{\nu}(z)$ is (symmetric) positive semidefinite (for any z). Hence, $\max_{z \in B_{\hat{\rho}}} \|\nabla^2 L_{\nu}(z)\| \leq \max_{z \in B_{\hat{\rho}}} \text{Tr}(\nabla^2 L_{\nu}(z))$. Now, observe that

$$\frac{\partial}{\partial z_i} L_{\nu}(z) = \frac{1}{M_{\nu}(z)} \left[\int_{\mathcal{X}} x_i \exp\langle z, x \rangle d\nu(x) \right] = \frac{1}{\int_{\mathcal{X}} \exp\langle z, x \rangle d\nu(x)} \left[\int_{\mathcal{X}} x_i \exp\langle z, x \rangle d\nu(x) \right],$$

where the interchange of the derivative and integral is permitted by the Leibniz rule for finite measures, see e.g. [19, Theorem 2.27] or [27, Theorem 6.28]. Hence,

$$\frac{\partial^2}{\partial z_i^2} L_{\nu}(z) = \frac{-1}{(M_{\nu}(z))^2} \left[\int_{\mathcal{X}} x_i \exp\langle z, x \rangle d\nu(x) \right]^2 + \frac{1}{M_{\nu}(z)} \left[\int_{\mathcal{X}} x_i^2 \exp\langle z, x \rangle d\nu(x) \right].$$

Taking the absolute value in the last identity, we may bound $|x_i| \leq \|x\| \leq |\mathcal{X}|$, $\|z\| \leq \hat{\rho}$, and apply Lemma 4.1 to bound $M_{\nu}(z)$. This eventually yields

$$\left| \frac{\partial^2}{\partial z_i^2} L_{\nu}(z) \right| \leq \frac{|\mathcal{X}|}{\exp(-\hat{\rho}|\mathcal{X}|)^2} \exp(\hat{\rho}|\mathcal{X}|)^2 + \frac{|\mathcal{X}|^2}{\exp(-\hat{\rho}|\mathcal{X}|)} \exp(\hat{\rho}|\mathcal{X}|) =: \hat{K},$$

with $\hat{K} > 0$ which depends on $\hat{\rho}$ and $|\mathcal{X}|$. As this uniformly bounds every term in the trace, $K := d \cdot \hat{K}$ is the desired constant. ■

The key feature of the constant K is that it does not depend on the choice of measure ν . Hence we can uniformly apply this bound over a family of measures, the most pertinent example being $\{\mu_n^{(\omega)}\}$. We remark that our upper bound on K is a vast overestimate for practical examples, which can be observed numerically. Finally we state a useful property of the excess.

Lemma 4.6. *Let $A, B \subset \mathbb{R}^d$ be nonempty and let B be closed and convex. Then for $\bar{a} \in A$ and $\bar{b} = \text{proj}_B(\bar{a})$ we have*

$$\|\bar{a} - \bar{b}\| \leq \text{exc}(A; B).$$

We now have developed all the necessary tools to state and prove the main result for the case of $g = \frac{1}{2}\|(\cdot) - b\|$.

Theorem 4.7. *Let ρ_0 be given by (4.3), and suppose $\text{rank}(C) = d$. Then for all $\mu, \nu \in \mathcal{P}(\mathcal{X})$, all $\rho > \rho_0$ and all $\varepsilon \in [0, \rho - \rho_0]$, we have the following: If $\bar{z}_{\nu, \varepsilon}$ is an ε -minimizer of ϕ_{ν} as defined in (4.1), then*

$$\bar{x}_{\nu, \varepsilon} := \nabla L_{\nu}(C^T \bar{z}_{\nu, \varepsilon})$$

satisfies the error bound

$$\|\bar{x}_{\nu, \varepsilon} - \bar{x}_{\mu}\| \leq \frac{1}{\alpha \sigma_{\min}(C)} D_{\rho}(\nu, \mu) + \frac{2\sqrt{2}}{\sqrt{\alpha} \sigma_{\min}(C)} \sqrt{D_{\rho}(\nu, \mu)} + \left(K \|C\| \sqrt{2\alpha} + \frac{2}{\sqrt{\alpha} \sigma_{\min}(C)} \right) \sqrt{\varepsilon},$$

where $\bar{x}_{\mu}$ is the unique solution to the MEM primal problem (P) for μ and $K > 0$ is a constant which does not depend on μ, ν .

522 *Proof.* Let $\rho > \rho_0$, $\nu, \mu \in \mathcal{P}(\mathcal{X})$ and $\varepsilon \in [0, \rho - \rho_0]$. Let $\bar{z}_{\nu, \varepsilon}$ be a ε -minimizer of ϕ_ν , and denote
 523 the unique minimizers of ϕ_μ and ϕ_μ as $\bar{z}_\nu$ and $\bar{z}_\mu$, respectively. Then

$$\begin{aligned} 524 \quad \|\bar{x}_{\nu, \varepsilon} - \bar{x}_\mu\| &= \|\nabla L_\nu(C^T \bar{z}_{\nu, \varepsilon}) - \nabla L_\mu(C^T \bar{z}_\mu)\| \\ 525 \quad &= \|\nabla L_\nu(C^T \bar{z}_{\nu, \varepsilon}) - \nabla L_\nu(C^T \bar{z}_\nu) + \nabla L_\nu(C^T \bar{z}_\nu) - \nabla L_\mu(C^T \bar{z}_\mu)\|, \end{aligned}$$

526 and so

$$527 \quad (4.7) \quad \|\bar{x}_{\nu, \varepsilon} - \bar{x}_\mu\| \leq \|\nabla L_\nu(C^T \bar{z}_{\nu, \varepsilon}) - \nabla L_\nu(C^T \bar{z}_\nu)\| + \|\nabla L_\nu(C^T \bar{z}_\nu) - \nabla L_\mu(C^T \bar{z}_\mu)\|.$$

528 To estimate the first term on the right hand side of (4.7), we require an auxiliary bound. Observe
 529 that, as ϕ_ν is strongly $1/\alpha$ -convex with $\nabla \phi_\nu(\bar{z}_\nu) = 0$, we have

$$530 \quad (4.8) \quad \|\bar{z}_\nu - \bar{z}_{\nu, \varepsilon}\| \leq \sqrt{2\alpha} |\phi_\nu(\bar{z}_\nu) - \phi_\nu(\bar{z}_{\nu, \varepsilon})|^{1/2} \leq \sqrt{2\alpha\varepsilon}.$$

531 Here the first inequality uses (2.1), while the second follows from the definition of $\bar{z}_\nu$ and $\bar{z}_{\nu, \varepsilon}$, as
 532 $|\phi_\nu(\bar{z}_\nu) - \phi_\nu(\bar{z}_{\nu, \varepsilon})| = \phi_\nu(\bar{z}_{\nu, \varepsilon}) - \phi_\nu(\bar{z}_\nu) \leq \varepsilon$.

533 From Lemma 4.2(b), we find that $\|\bar{z}_\nu\| \leq \rho_0$. Thus, (4.8) yields $\|\bar{z}_{\nu, \varepsilon}\| \leq \rho_0 + \sqrt{2\alpha\varepsilon}$. This
 534 implies $\|C^T \bar{z}_\nu\|, \|C^T \bar{z}_{\nu, \varepsilon}\| \leq \|C\|(\rho_0 + \sqrt{2\alpha\varepsilon})$. Hence, Corollary 4.5 with $\hat{\rho} = \|C\|(\rho_0 + \sqrt{2\alpha\varepsilon})$
 535 yields

$$536 \quad \|\nabla L_\nu(C^T \bar{z}_{\nu, \varepsilon}) - \nabla L_\nu(C^T \bar{z}_\nu)\| \leq K \|C^T \bar{z}_\nu - C^T \bar{z}_{\nu, \varepsilon}\|,$$

537 where K depends on $\hat{\rho}, |\mathcal{X}|, d$ and therefore on $|\mathcal{X}|, \|C\|, b, \varepsilon, \alpha, d$. The right-hand side in the last
 538 inequality can be further estimated with (4.8) to find

$$539 \quad (4.9) \quad \|C^T \bar{z}_\nu - C^T \bar{z}_{\nu, \varepsilon}\| \leq \|C\| \|\bar{z}_\nu - \bar{z}_{\nu, \varepsilon}\| \leq \|C\| \sqrt{2\alpha\varepsilon}.$$

540 We now turn to the second term on the right-hand side of (4.7). First order optimality condi-
 541 tions give

$$542 \quad 0 = -\frac{\bar{z}_\nu}{\alpha} + b + C \nabla L_\nu(C^T \bar{z}_\nu), \quad 0 = -\frac{\bar{z}_\mu}{\alpha} + b + C \nabla L_\mu(C^T \bar{z}_\mu),$$

543 and therefore $\|C(\nabla L_\nu(C^T \bar{z}_\nu) - \nabla L_\mu(C^T \bar{z}_\mu))\| = \frac{1}{\alpha} \|\bar{z}_\nu - \bar{z}_\mu\|$. Furthermore, as $\text{rank}(C) = d$ we
 544 have $\sigma_{\min}(C) > 0$. We also have for, any $x \in \mathbb{R}^d$, that $\|Cx\| \geq \sigma_{\min}(C)\|x\|$, and hence

$$545 \quad \|\nabla L_\nu(C^T \bar{z}_\nu) - \nabla L_\mu(C^T \bar{z}_\mu)\| \leq \frac{1}{\sigma_{\min}(C)} \|C(\nabla L_\nu(C^T \bar{z}_\nu) - \nabla L_\mu(C^T \bar{z}_\mu))\| = \frac{1}{\alpha \sigma_{\min}(C)} \|\bar{z}_\nu - \bar{z}_\mu\|.$$

546 In order to bound $\|\bar{z}_\nu - \bar{z}_\mu\|$ from above, we define $\delta := 2(\varepsilon + 2D_\rho(\nu, \mu))$. Denoting as usual $S_\delta(\mu)$
 547 as the set of δ -minimizers of ϕ_μ , which is a closed, convex set by the continuity and convexity of
 548 ϕ_μ respectively, define $y = \text{proj}_{S_\delta(\mu)}(\bar{z}_\nu)$. The triangle inequality gives

$$549 \quad (4.10) \quad \|\bar{z}_\nu - \bar{z}_\mu\| \leq \|\bar{z}_\nu - y\| + \|y - \bar{z}_\mu\|.$$

550 By the choice of $\rho > \rho_0$, we have by Lemma 4.2(b) that $\bar{z}_\nu \in S_\varepsilon(\nu) \cap B_\rho$. Therefore applying
 551 Lemma 4.6 with $A = S_\varepsilon(\nu) \cap B_\rho$, $B = S_\delta(\mu)$ we can bound the first term on the right hand side of
 552 (4.10) as

$$553 \quad (4.11) \quad \|\bar{z}_\nu - y\| \leq \text{exc}(S_\varepsilon(\nu) \cap B_\rho; S_\delta(\mu)).$$

554 For the remaining term of the right hand side of (4.10), we use the characterization (2.1) of the
 555 $\frac{1}{\alpha}$ -strong convexity in the differentiable case for ϕ_μ , noting $\nabla \phi_\mu(\bar{z}_\mu) = 0$. Hence

$$556 \quad (4.12) \quad \|y - \bar{z}_\mu\| \leq \sqrt{2\alpha} |\phi_\mu(y) - \phi_\mu(\bar{z}_\mu)|^{1/2} \leq \sqrt{2\alpha\delta},$$

557 where $y \in S_\delta(\mu)$ for the second inequality. Combining (4.9)–(4.12) with (4.7), we find that

$$558 \quad \|\bar{x}_{\nu,\varepsilon} - \bar{x}_\mu\| \leq K\|C\|\sqrt{2\alpha\varepsilon} + \frac{1}{\alpha\sigma_{\min}(C)}\text{exc}(S_\varepsilon(\nu) \cap B_\rho; S_\delta(\mu)) + \frac{1}{\alpha\sigma_{\min}(C)}\sqrt{2\alpha\delta}.$$

559 By the choice of $\delta = 2(\varepsilon + 2D_\rho(\nu, \mu))$, Lemma 4.4 asserts $\text{exc}(S_\varepsilon(\nu) \cap B_\rho; S_\delta(\mu)) \leq D_\rho(\nu, \mu)$.
 560 Therefore

$$\begin{aligned} 561 \quad \|\bar{x}_{\nu,\varepsilon} - \bar{x}_\mu\| &\leq \frac{1}{\alpha\sigma_{\min}(C)}\text{exc}(S_\varepsilon(\nu) \cap B_\rho; S_\delta(\mu)) + \frac{1}{\alpha\sigma_{\min}(C)}\sqrt{2\alpha\delta} + K\|C\|\sqrt{2\alpha\varepsilon} \\ 562 \quad &\leq \frac{1}{\alpha\sigma_{\min}(C)}D_\rho(\nu, \mu) + \frac{1}{\sigma_{\min}(C)}\sqrt{\frac{4\varepsilon}{\alpha}} + \frac{1}{\sigma_{\min}(C)}\sqrt{\frac{8D_\rho(\nu, \mu)}{\alpha}} + K\|C\|\sqrt{2\alpha\varepsilon} \\ 563 \quad &= \frac{1}{\alpha\sigma_{\min}(C)}D_\rho(\nu, \mu) + \frac{2\sqrt{2}}{\sqrt{\alpha}\sigma_{\min}(C)}\sqrt{D_\rho(\nu, \mu)} + \left(K\|C\|\sqrt{2\alpha} + \frac{2}{\sqrt{\alpha}\sigma_{\min}(C)}\right)\sqrt{\varepsilon} \end{aligned}$$

564 where in the second line we have used the definition of δ and the concavity of $\sqrt{x+y} \leq \sqrt{x} + \sqrt{y}$. ■

565 Note that we may set $\varepsilon = 0$ for a corollary on exact minimizers. However, the error bound still
 566 has the same scaling in terms of $D_\rho(\nu, \mu)$. This theorem is the first of its type to link MEM solu-
 567 tions with epigraphical distances. This result has the benefit applying uniformly in the choice η ,
 568 so this theorem can be applied for *any* choice of η which provides a bound on $D_\rho(\mu, \eta)$ - and we
 569 demonstrate the particular case of $\nu = \mu_n^{(\omega)}$ in the next section. While this result has the stringent
 570 assumption that $\text{rank}(C) = d$, we emphasize that this does not appear in the qualitative Theo-
 571 rem 3.9 which guarantees the almost sure convergence of solutions. We conjecture that weakening
 572 this assumption may be possible, but will require alternative techniques. Finally we remark that
 573 the scaling $\sqrt{D_\rho(\mu, \eta)}$ arises as a direct consequence of the smoothness and strong convexity of
 574 $g = \frac{1}{2}\|b - (\cdot)\|^2$, see e.g. [40, Theorem 4.2] and the following discussion - and hence it may be
 575 feasible to prove improved rates for other specialized choices of g .

576 **5. A statistical dependence on n .** This section is devoted to making the dependence on n
 577 explicit in Theorem 4.7 for the special case $\nu = \mu_n^{(\omega)}$. We briefly recall the empirical setting
 578 developed in Subsection 2.3. Given i.i.d. random vectors $\{X_1, X_2, \dots, X_n, \dots\}$ on $(\Omega, \mathcal{F}, \mathbb{P})$ with
 579 shared law $\mu = \mathbb{P} X_1^{-1}$, we define $\mu_n^{(\omega)} = \sum_{i=1}^n \delta_{X_i(\omega)}$. For this measure, the dual objective reads

$$580 \quad \phi_{\mu_n^{(\omega)}}(z) = \frac{1}{2\alpha}\|z\|^2 - \langle b, z \rangle + \log \frac{1}{n} \sum_{i=1}^n e^{\langle C^T z, X_i(\omega) \rangle}.$$

581 Given $\bar{z}_{n,\varepsilon}(\omega)$, an ε -minimizer of $\phi_{\mu_n^{(\omega)}}(z)$, define

$$582 \quad \bar{x}_{n,\varepsilon}(\omega) := \nabla L_{\mu_n^{(\omega)}}(C^T \bar{z}_{n,\varepsilon}(\omega)) = \frac{\sum_{i=1}^n C X_i(\omega) e^{\langle C^T \bar{z}_{n,\varepsilon}(\omega), X_i(\omega) \rangle}}{\sum_{i=1}^n e^{\langle C^T \bar{z}_{n,\varepsilon}(\omega), X_i(\omega) \rangle}}.$$

583 We begin with a simplifying lemma, recalling the notation developed in Section 4 of the moment
 584 generating function M_μ of μ and empirical moment generating function $M_n(\cdot, \omega)$ of $\mu_n^{(\omega)}$.

585 **Lemma 5.1.** *Let $\rho > 0$, $n \in \mathbb{N}$ and $\omega \in \Omega$ and set $K := \exp(\rho\|C\|\|\mathcal{X}\|)$. Then*

$$586 \quad D_\rho(\mu, \mu_n^{(\omega)}) \leq K \max_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|.$$

587 **Proof.** Applying Lemma 4.1 to the particular probability measures μ and $\mu_n^{(\omega)}$ gives

$$588 \quad (5.1) \quad M_\mu(C^T z), M_n(C^T z, \omega) \in [\exp(-\rho\|C\|\|\mathcal{X}\|), \exp(\rho\|C\|\|\mathcal{X}\|)] =: [c, d]$$

589 where $0 < c < d$. Furthermore, for any $s, t \in [c, d]$ we have $|\log(s) - \log(t)| \leq \frac{1}{c}|s - t|$, and hence

$$590 \quad D_\rho(\mu, \mu_n^{(\omega)}) = \max_{z \in B_\rho} |L_{\mu_n^{(\omega)}}(C^T z) - L_\mu(C^T z)| \leq \exp(\rho\|C\|\|\mathcal{X}\|) \max_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|. \quad \blacksquare$$

591 **Lemma 5.2.** Let ρ_0 be as defined in (4.3) and suppose $\text{rank}(C) = d$. Then, for all $\rho > \rho_0$, all
 592 $\varepsilon \in [0, \rho - \rho_0]$, and for all $n \in \mathbb{N}$ we have: if $\bar{z}_{n,\varepsilon}(\omega) \in S_\varepsilon(\mu_n^{(\omega)})$, then $\bar{x}_{n,\varepsilon}(\omega) = \nabla L_{\mu_n^{(\omega)}}(C^T \bar{z}_{n,\varepsilon})$
 593 satisfies

$$594 \quad \|\bar{x}_{n,\varepsilon}(\omega) - \bar{x}_\mu\| \leq \frac{K_1}{\alpha \sigma_{\min}(C)} \max_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|$$

$$595 \quad + \frac{2\sqrt{2}K_1}{\sqrt{\alpha} \sigma_{\min}(C)} \sqrt{\max_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|} + \left(K_2 \|C\| \sqrt{2\alpha} + \frac{2}{\sqrt{\alpha} \sigma_{\min}(C)} \right) \sqrt{\varepsilon}$$

596 where K_1 is a constant which depends on $\rho, |\mathcal{X}|, \|C\|$, and K_2 on $|\mathcal{X}|, \|C\|, b, \varepsilon, d, \alpha$.

597 **Proof.** As $\mu_n^{(\omega)} \in \mathcal{P}(\mathcal{X})$, Theorem 4.7 yields for all n , for ρ_0 as defined in (4.3), for all $\rho > \rho_0$,
 598 and all $\varepsilon \in [0, \rho - \rho_0]$, if $\bar{z}_{n,\varepsilon}(\omega) \in S_\varepsilon(\mu_n^{(\omega)})$, then $\bar{x}_{n,\varepsilon}(\omega) = \nabla L_\nu(C^T \bar{z}_{n,\varepsilon})$ satisfies

$$599 \quad \|\bar{x}_{n,\varepsilon}(\omega) - \bar{x}_\mu\| \leq \frac{1}{\alpha \sigma_{\min}(C)} D_\rho(\mu, \mu_n^{(\omega)}) + \frac{2\sqrt{2}}{\sqrt{\alpha} \sigma_{\min}(C)} \sqrt{D_\rho(\mu, \mu_n^{(\omega)})}$$

$$600 \quad + \left(K_2 \|C\| \sqrt{2\alpha} + \frac{2}{\sqrt{\alpha} \sigma_{\min}(C)} \right) \sqrt{\varepsilon},$$

601 where we stress that the constant K_2 depends on $|\mathcal{X}|, \|C\|, b, \varepsilon, \alpha, d$, but does not depend on n . Ap-
 602 plying Lemma 5.1 to bound $D_\rho(\mu, \mu_n^{(\omega)}) \leq K_1 \sup_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|$ gives the result. ■

603 In order to construct a final bound which depends explicitly on n it remains to estimate the term
 604 $\max_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|$. This fits into the language of empirical process theory where
 605 this type of convergence is well studied. The main reference of interest here is van der Vaart [45].

606 For compact $\mathcal{X} \subset \mathbb{R}^d$, let $f : \mathcal{X} \rightarrow \mathbb{R}$ be a function and $\beta = (\beta_1, \dots, \beta_d)$ be a multi-index, i.e.
 607 a vector of d nonnegative integers. We call $|\beta| = \sum_i \beta_i$ the order of β , and define the differential
 608 operator $D^\beta = \frac{\partial^{\beta_1}}{\partial x_1^{\beta_1}} \cdots \frac{\partial^{\beta_d}}{\partial x_d^{\beta_d}}$. For integer k , we denote by $\mathcal{C}^k(\mathcal{X})$ as the space of k -smooth (also
 609 known as k -Hölder continuous) functions on $\mathcal{X}$, namely, those f which satisfy [45, p. 2131]

$$610 \quad \|f\|_{\mathcal{C}^k(\mathcal{X})} := \max_{|\beta| \leq k} \sup_{x \in \text{int}(\mathcal{X})} \|D^\beta f(x)\| + \max_{|\beta|=k} \sup_{\substack{x, y \in \text{int}(\mathcal{X}) \\ x \neq y}} \left| \frac{D^\beta f(x) - D^\beta f(y)}{\|x - y\|} \right| < +\infty.$$

611 Moreover, let $\mathcal{C}_R^k(\mathcal{X})$ denote the ball of radius R in $\mathcal{C}^k(\mathcal{X})$. With this notation developed we can state
 612 the classical results of van der Vaart [45]. In the notation therein, we apply the machinery of Sections
 613 1 and 2 to the measure space $(\mathcal{X}_1, \mathcal{A}_1) = (\mathcal{X}, \mathcal{B}_\mathcal{X})$, equipped with probability measure μ . Taking
 614 $\mathbb{G}_n = \sqrt{n}(\mu_n^{(\omega)} - \mu)$ and $\mathcal{F}_1 = \mathcal{F} = \mathcal{C}_R^k(\mathcal{X})$, this induces the norm $\|\mathbb{G}_n\|_{\mathcal{F}} = \sup_{f \in \mathcal{C}_R^k(\mathcal{X})} \{|\int_{\mathcal{X}} f d\mathbb{G}_n|\}$,
 615 and hence the results of [45, p. 2131] give

616 **Theorem 5.3.** Let $\mu \in \mathcal{P}(\mathcal{X})$. If $k > d/2$, then for any $R > 0$,

$$617 \quad \mathbb{E}_{\mathbb{P}}^* \left[\sup_{f \in \mathcal{C}_R^k(\mathcal{X})} \sqrt{n} \left| \int_{\mathcal{X}} f d\mu_n^{(\cdot)} - \int_{\mathcal{X}} f d\mu \right| \right] \leq D$$

618 where D is a constant depending (polynomially) on $k, d, |\mathcal{X}|, R$, and $\mathbb{E}_{\mathbb{P}}^*$ is the outer expectation to
 619 avoid concerns of measurability (see e.g. [46, Section 1.2]).

620 Here, the outer expectation is defined for $f : \Omega \rightarrow \mathbb{R}$ as

$$621 \quad \mathbb{E}_{\mathbb{P}}^*(f) := \inf_{\substack{h \geq f \text{ pointwise} \\ h \text{ measurable}}} \mathbb{E}_{\mathbb{P}}(h)$$

622 which coincides with the usual expectation for measurable functions. We remark that outer expec-
 623 tation is also known as the outer integral [37, Chapter 14.F]. A self-contained proof of Theorem 5.3

is non-trivial, requiring the development of entropy and bracketing numbers of function spaces which is beyond the scope of this article.⁹ We simply take this result as given. However, we show the following corollary.

Corollary 5.4. *For all $n \in \mathbb{N}$, we have*

$$\mathbb{E}_{\mathbb{P}} \left[\max_{z \in B_{\rho}} |M_{\mu}(C^T z) - M_n(C^T z, \cdot)| \right] \leq \frac{D}{\sqrt{n}},$$

where D is a constant depending on $d, |\mathcal{X}|, \|C\|, \rho$.

Proof. Observe that for each z , the function $f_z(x) = \exp\langle C^T z, \cdot \rangle$ is an infinitely differentiable function on the compact set $\mathcal{X}$, and thus has bounded derivatives of all orders, in particular, of order $k = d > d/2$. Hence, for $\rho > 0$, the set of functions f_z parameterized by $z \in B_{\rho}$ satisfies

$$\left\{ f_z(x) = \exp\langle C^T z, x \rangle : z \in B_{\rho} \right\} \subset \mathcal{C}_{R_d}^d(\mathcal{X}),$$

where R_d is a constant which depends on $d, \rho, \|C\|$, and $|\mathcal{X}|$. Furthermore, as $\mathbb{Q}^m \cap B_{\rho} \subset B_{\rho}$ is a countable dense subset, $\max_{z \in B_{\rho}} |M_{\mu}(C^T z) - M_n(C^T z, \cdot)| = \sup_{z \in \mathbb{Q}^m \cap B_{\rho}} |M_{\mu}(C^T z) - M_n(C^T z, \cdot)|$ is a supremum of countably many $\mathbb{P}$ -measurable functions and is hence $\mathbb{P}$ -measurable. In particular the usual expectation agrees with the outer expectation. Hence applying [Theorem 5.3](#) we may assert

$$\begin{aligned} \mathbb{E}_{\mathbb{P}} \left[\max_{z \in B_{\rho}} |M_{\mu}(C^T z) - M_n(C^T z, \cdot)| \right] &= \mathbb{E}_{\mathbb{P}}^* \left[\sup_{z \in B_{\rho}} \left| \int_{\mathcal{X}} f_z d\mu_n^{(\cdot)} - \int_{\mathcal{X}} f_z d\mu \right| \right] \\ &\leq \mathbb{E}_{\mathbb{P}}^* \left[\sup_{f \in \mathcal{C}_{R_d}^d(\mathcal{X})} \left| \int_{\mathcal{X}} f d\mu_n^{(\cdot)} - \int_{\mathcal{X}} f d\mu \right| \right] \\ &\leq \frac{D}{\sqrt{n}}, \end{aligned}$$

for a constant D which depends on $d, |\mathcal{X}|$ and R_d . We remark that the choice of $k = d$ in the above was aesthetic, to remove the dependence of D on k . ■

The final result now follows as a simple consequence of [Corollary 5.4](#) and [Lemma 5.2](#):

Theorem 5.5. *Suppose $\text{rank}(C) = d$. For all $n \in \mathbb{N}$, and all $\bar{z}_{n,\varepsilon}(\omega) \in S_{\varepsilon}(\mu_n^{(\omega)})$, the associated $\bar{x}_{n,\varepsilon}(\omega) = \nabla L_{\mu_n^{(\omega)}}(C^T \bar{z}_{n,\varepsilon}(\omega))$ satisfies*

$$\begin{aligned} \mathbb{E}_{\mathbb{P}}^* \|\bar{x}_{n,\varepsilon}(\cdot) - \bar{x}_{\mu}\| &\leq \frac{DK_1}{\alpha \sigma_{\min}(C)} \frac{1}{\sqrt{n}} + \frac{2D\sqrt{2K_1}}{\sqrt{\alpha} \sigma_{\min}(C)} \sqrt{\frac{1}{\sqrt{n}}} + \left(K_2 \|C\| \sqrt{2\alpha} + \frac{2}{\sqrt{\alpha} \sigma_{\min}(C)} \right) \sqrt{\varepsilon} \\ &= O\left(\frac{1}{n^{1/4}} + \sqrt{\varepsilon} \right), \end{aligned}$$

where the leading constants K_1, K_2, D depend on $|\mathcal{X}|, \|C\|, b, \varepsilon, \alpha, d$. In particular, for the case $\varepsilon = 0$, the unique minimizer $\bar{x}_n(\omega)$ is always measurable and hence the outer expectation is the usual expectation.

⁹Bounds of this ‘‘Donsker’’ type have previously been applied to empirical approximations of stochastic optimization problems, to derive large deviation-style results for specific problems, see [\[39, Section 5.5\]](#) for a detailed exposition and discussion, in particular [\[39, Theorem 5.2\]](#). In principle this machinery could be used here to derive similar large deviation results.

652 *Proof.* Take $\rho = 2\rho_0$ in [Lemma 5.2](#). Then for any n , if $\bar{z}_{n,\varepsilon}(\omega) \in S_\varepsilon(\mu_n^{(\omega)})$, we have for all ω

$$653 \quad \|\bar{x}_{n,\varepsilon}(\omega) - \bar{x}_\mu\| \leq \frac{K_1}{\alpha\sigma_{\min}(C)} \sup_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|$$

$$654 \quad + \frac{2\sqrt{2}\sqrt{K_1}}{\sqrt{\alpha}\sigma_{\min}(C)} \sqrt{\sup_{z \in B_\rho} |M_\mu(C^T z) - M_n(C^T z, \omega)|} + \left(K_2 \|C\| \sqrt{2\alpha} + \frac{2}{\sqrt{\alpha}\sigma_{\min}(C)} \right) \sqrt{\varepsilon},$$

655 holds with constants K_1, K_2 that depend only on $|\mathcal{X}|, \|C\|, b, \varepsilon, \alpha, d$. Taking the outer expectation
 656 on both sides, and applying [Corollary 5.4](#) to the measurable right hand side gives the result.
 657 For the latter assertion, by [37, Theorem 14.37] there always exists a measurable selection $\omega \rightarrow$
 658 $\operatorname{argmin}_{\mu_n^{(\omega)}} \phi_{\mu_n^{(\omega)}} = \{\bar{z}_n(\omega)\}$. In particular, $\bar{z}_n(\omega)$ is unique by the $\frac{1}{\alpha}$ -strong convexity of $\phi_{\mu_n^{(\omega)}}$, and thus
 659 there is only one possible selection which is immediately measurable. Hence the function $\omega \rightarrow \bar{x}_n(\omega)$
 660 is the composition of the continuous function $\nabla L_{\mu_n^{(\omega)}}$ and the measurable function $C^T \bar{z}_n(\omega)$ - and is
 661 hence measurable. Remarking also that $\bar{x}_n(\omega)$ is unique, the left hand side is always $\mathbb{P}$ -measurable
 662 and the outer expectation agrees with the usual expectation. ■

663 More generally we note that for all $\varepsilon > 0$, there always exists a measurable selection $z_{n,\varepsilon}(\omega) \in$
 664 $S_\varepsilon(\mu_n^{(\omega)})$ via [37, Theorem 14.6, Proposition 14.33]. Furthermore any algorithm $\mathcal{A}$ which performs
 665 a composition of operations which preserve Borel measurability - such summations, products, differ-
 666 entiations, and evaluations of measurable functions (see e.g. [19, Section 2.1]) - defines a selection
 667 $\omega \rightarrow \mathcal{A}(\phi_{\mu_n^{(\omega)}}) = \bar{z}_{n,\varepsilon}(\omega)$ which is a composition of Borel-measurable functions and hence mea-
 668 surable. With this in mind, issues of measurability are a technical note rather than of practical
 669 concern.

670 Finally, we remark that this result is the first to give a parametric rate for convergence of
 671 approximate minimizers of the empirical MEM problem. We conjecture however that this result
 672 is not sharp, and that a sharper convergence rate of $n^{1/2}$ may be proven using different analytical
 673 techniques. This conjecture will be examined numerically in the next section.

674 **6. Numerical experiments.** We now shift to a numerical examination of the convergence $\bar{x}_{n,\varepsilon} \rightarrow$
 675 $\bar{x}_\mu$. We focus entirely on the most recent setting of [Section 5](#), the MEM problem with an empirical
 676 prior $\mu_n^{(\omega)}$ and the fidelity term $g = \frac{1}{2}\|b - (\cdot)\|^2$. As discussed throughout, the empirical dual
 677 objective function $\phi_{\mu_n^{(\omega)}}$ is smooth and strongly convex, with easily computable derivatives:¹⁰

$$678 \quad \nabla \phi_{\mu_n^{(\omega)}}(z) = \frac{1}{2\alpha} z - b + \frac{\sum_{i=1}^n C X_i(\omega) e^{\langle C^T z, X_i(\omega) \rangle}}{\sum_{i=1}^n e^{\langle C^T z, X_i(\omega) \rangle}}.$$

679 Hence we may solve this problem with any standard optimization package, and here we choose
 680 to use limited memory BFGS [9]. As a companion to this paper, we include a jupyter notebook,
 681 <https://github.com/mattkingros/MEM-Denoising-and-Deblurring.git>, which can be used to repro-
 682 duce and extend all computations performed hereafter. We emphasize that, as this work is the first
 683 to propose using empirical data in the MEM framework, there are many numerical considerations
 684 that we will not address. More sophisticated and higher order optimization routines would natu-
 685 rally be of interest for moderate n , as are questions of how to efficiently solve this problem when n
 686 grows prohibitively large or when C is exceedingly close to singular.

687 For our numerical experiments we will focus purely on denoising, namely the case $C = I$.
 688 We will use two datasets and two distributions of noise as proof of concept. For noise, we use
 689 additive Gaussian noise and “salt-and-pepper” corruption noise where each pixel has an independent
 690 probability of being set to purely black or white. For datasets, the first is the MNIST digits dataset
 691 [15] which contains 60000 28x28 grayscale pixel images of hand-drawn digits 0-9, and the second is

¹⁰The derivatives are easily available analytically, but a naïve implementation may run into stability issues stem-
 ming from overflow when computing large sums of exponential terms. This can be easily addressed, see e.g. [25].

the more expressive dataset of Fashion-MNIST [47], which is once again 60000 28x28 grayscale pixel images but of various garments such as shoes, sneakers, bags, pants, and so on. We remark these choices of noise, dataset, and matrix are far from real world problems, and serve as a lightweight implementation of data into the MEM framework.

In all experiments we *always* include enough noise so that the nearest neighbour to b in the dataset belongs to a different class than the ground truth, ensuring that recovery is non-trivial. We include this in all figures, captioned “closest point”. Furthermore, we take the ground truth to be a new data-point hand drawn by the authors; in particular, it is *not present* in either of the original MNIST or MNIST fashion datasets.

We begin with a baseline examination of the error in n for the MNIST digit dataset, for various choices of b . To generate $\mu_n^{(\omega)}$ practically, we sample n datapoints uniformly at random without replacement. As a remark on this methodology, the target best possible approximation is the one which uses all possible data, i.e., $\mu_D := \mu_{60,000}^{(\omega)}$. Hence, for error plots we compare to $\bar{x}_{\mu_D}$, as we do not have access to the full image distribution to construct $\bar{x}_\mu$.

Given a noisy image b , the experimental setup is as follows: for 20 values of n , spaced linearly between 10000 and 60000, we perform 15 random samples of size n . For each random sample, we compute an approximation $\bar{x}_{n,\epsilon}$ and the relative approximation error $\frac{\|\bar{x}_{n,\epsilon} - \bar{x}_{\mu_D}\|}{\|\bar{x}_{\mu_D}\|}$, which is then averaged over the trials. Superimposed is the upper bound $K_1 n^{1/4}$ convergence rate, as well as the conjectured $K_2 n^{1/2}$ rate, for moderate constants K_1, K_2 which changes between figures. We also visually exhibit several reconstructions, the nearest neighbour, and postprocessed images for comparison. This methodology is used to create figures [Figures 1, 3, 6, 8, 10 and 12](#)

A remark on postprocessing: As alluded to in the introduction (see [Remark 1.1](#)), an advantage of the dual approach is the ability to reconstruct the optimal measure Q_n which solves the measure-valued primal problem as seen in [34]. In particular as $Q_n \ll \mu_n^{(\omega)}$, we have $\bar{x}_n = \mathbb{E}_{Q_n}$ is a particular weighted linear combination of the input data. This allows for two types of natural postprocessing: at the level of the linear combination, i.e., setting all weights below some given threshold to zero, or at the level of the pixels: i.e. setting all pixels above $1-\gamma$ to 1 and below γ to 0. These postprocessing steps are motivated by the observation that solutions are compressible both in the linear combination and the pixel-intensity level. Hence our figures also include the final measure Q_n with all entries below 0.01 set to zero, the corresponding linear combination of the remaining datapoints (bottom right image), and a further masking at the pixel level (top right image).

With this methodology developed, we present the results. For the first figures, the ground truth is a hand-drawn 8, which is approximated by sampling the MNIST digit dataset. For [Figure 1](#) we use additive Gaussian noise of variance $\sigma = 0.10\|x\|$. For [Figure 3](#) we use salt-and-pepper corruptions with an equal probability of 0.2 of any given pixel being set to 1 or 0.

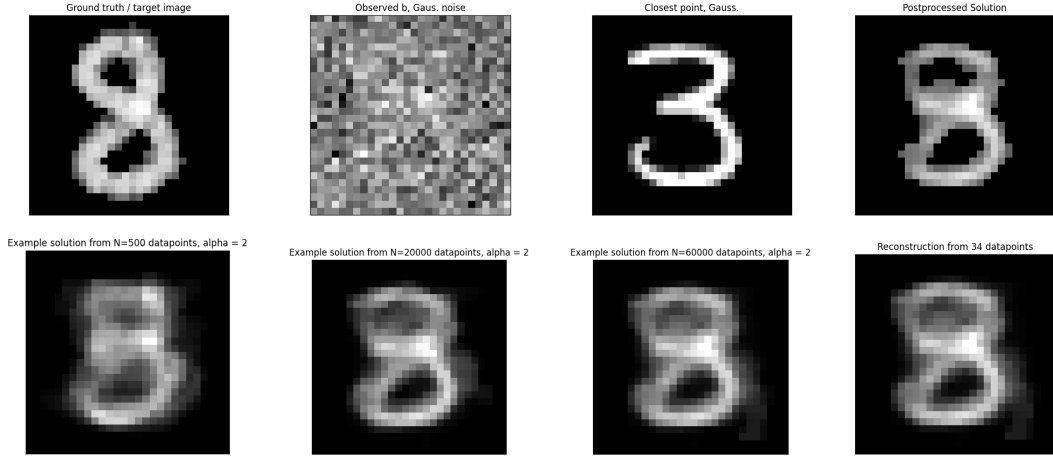


Figure 1: Recovery of an 8 with additive Gaussian noise

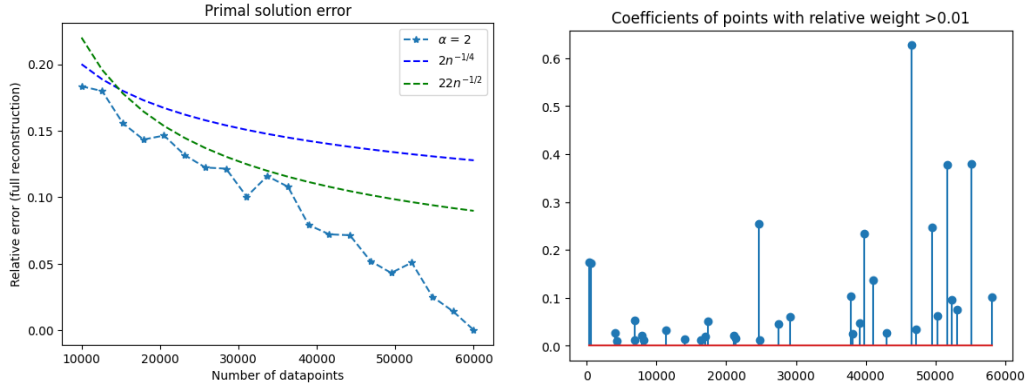


Figure 2: Rates and thresholding of optimal measure Q_n , for Figure 1

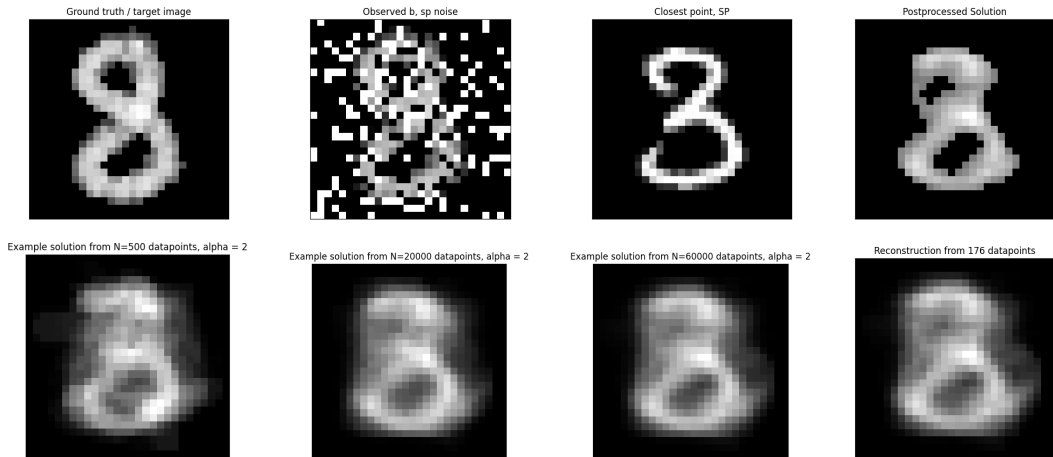


Figure 3: Recovery of an 8 with salt-and-pepper corruption noise

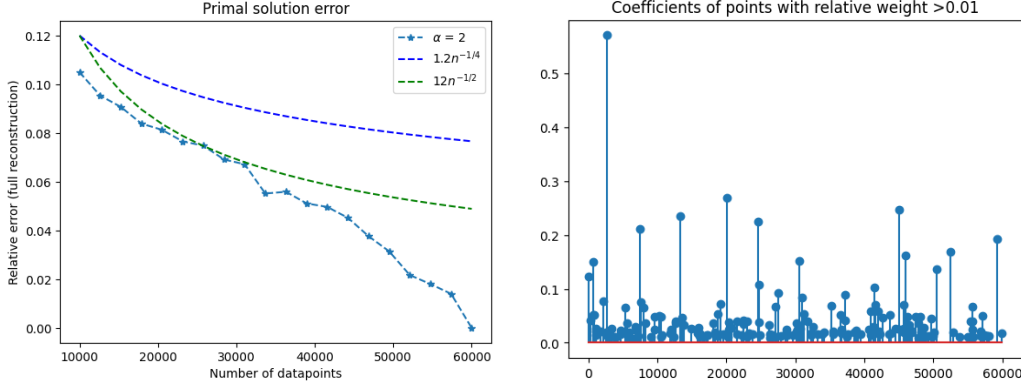


Figure 4: Rates and thresholding of optimal measure Q_n , for Figure 3

Examining Figures 1 and 3, we can make several observations which will persist globally in all numerics. As seen in Figures 2 and 4, the theoretical rate $n^{1/4}$ agrees with practical results, and appears to be tight for moderate n on the order of $n < 30000$. While the leading constants given by theory are quite large, growing polynomially with d, ρ e.g., in practice these are much more moderate - here $O(1)$. We also note the linear combination forming the final solution is quite compressible, here having only 157 datapoints having dual measure above 0.01, and after thresholding being visually indistinguishable from the full solution. Furthermore the *visual* convergence is fast, in the sense that there is not an appreciable visual difference between the solution given 20,000 and 60,000 datapoints. As they are formed as linear combinations of observed data, all solutions suffer from artefacting in the form of blurred edges, which can be solved by a final mask at the pixel level.

Before shifting away from the MNIST dataset, we also include a cautionary experiment where, for small random samples, the method is visually “confidently incorrect”, which can be seen in Figure 5. In the previous examples Figures 1 and 3 for small n we simply have blurry images. This type of failure is generally representative of how the method performs when n is too small. However there is another failure case which distinctly highlights the risk of taking n too small, which can lead to a biased sample which does not well approximate μ - leading to correspondingly biased solutions. In Figure 5 we see the recovered image after masking is clearly a 3 for one random samples of size $n = 1000$, but a 5 for other random samples of size $n = 800, 2000$.

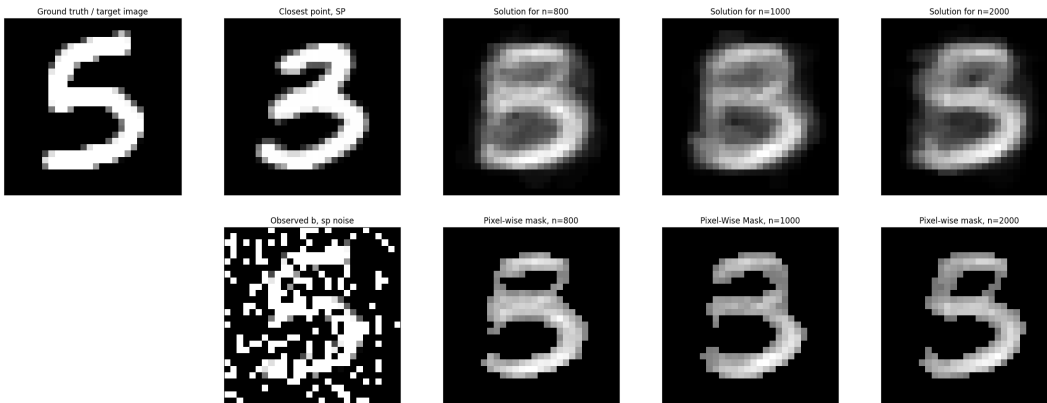


Figure 5: A specific type of failure case for small n .

We now move to the more expressive MNIST fashion dataset. We once again use a hand drawn target, which is not originally found in the dataset. To emphasize differences compared to handwritten digits, the fashion dataset is immediately more challenging, especially in regards to

748 fine details. To illustrate, there are few examples of garments with spots, stripes, or other detailed
749 patterns - and as such are much more difficult to learn from uniform random samples. Similarly,
750 if the target ground truth image contains details or patterns which are not present in the fashion
751 dataset, there is little hope of constructing a reasonable linear combination to approximate the
752 ground truth. On the other hand, some classes - such as heels or sandals - are extremely easy to
753 learn, as there are many near-identical examples in the dataset, and are visually quite distinct from
754 many other classes. Disregarding the change in dataset, our methodology for remains the same as
755 [Figure 1](#).

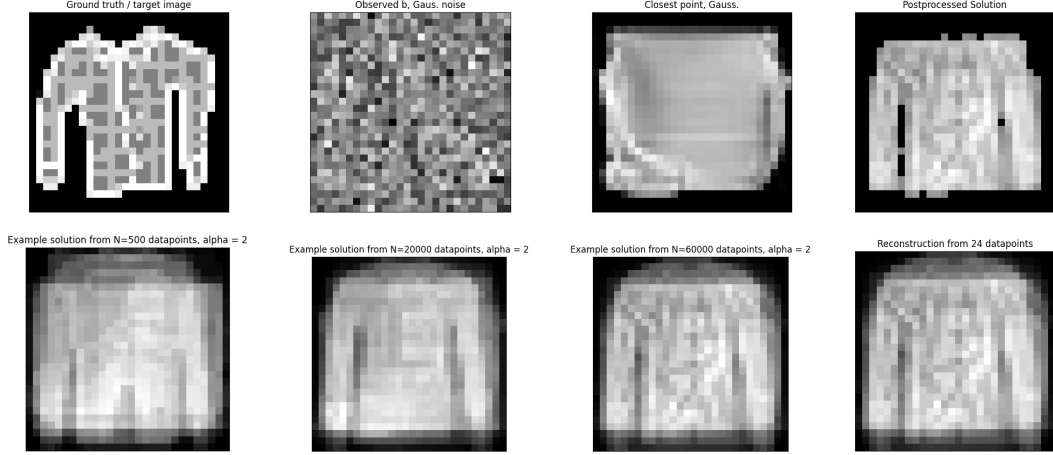


Figure 6: Recovery of a hand drawn shirt with additive Gaussian noise

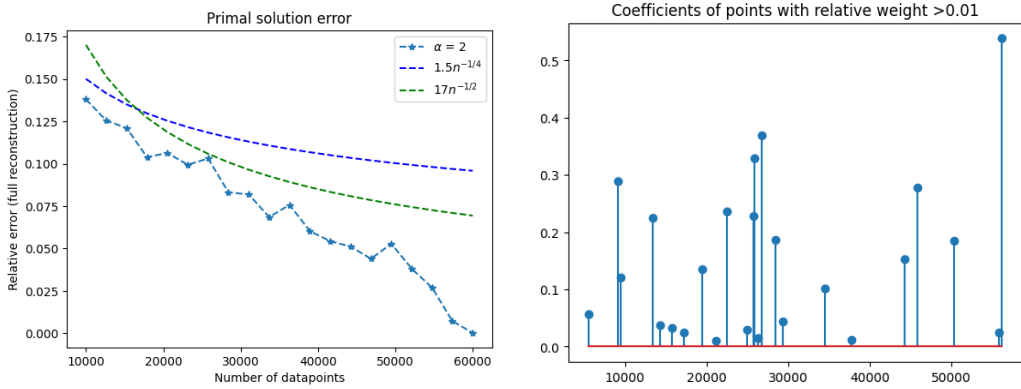


Figure 7: Rates and thresholding of optimal measure Q_n , for [Figure 6](#)

756 For the experiment with salt-and-pepper noise [Figure 8](#), while MEM denoising clearly recovers
757 a shirt, many of the finer details are lost or washed out. In contrast, with Gaussian noise [Figure 6](#),
758 there is some remnant of the shirts' pattern visible in the final reconstruction. In both cases the
759 nearest neighbor is a bag or purse, and not visually close to the ground truth. While the constant
760 leading constant is larger, once again we are firmly below the theoretically expected convergence
761 rate of $n^{1/4}$, and solutions are compressible with respect to the optimal measure Q_n .

762 Finally, we conclude with an experiment on MNIST fashion with a hand drawn target of a
763 heel, where we observe once again recovery well within the expected convergence rate, and visually
764 recovers (after postprocessing) a reasonable approximation to the ground truth.

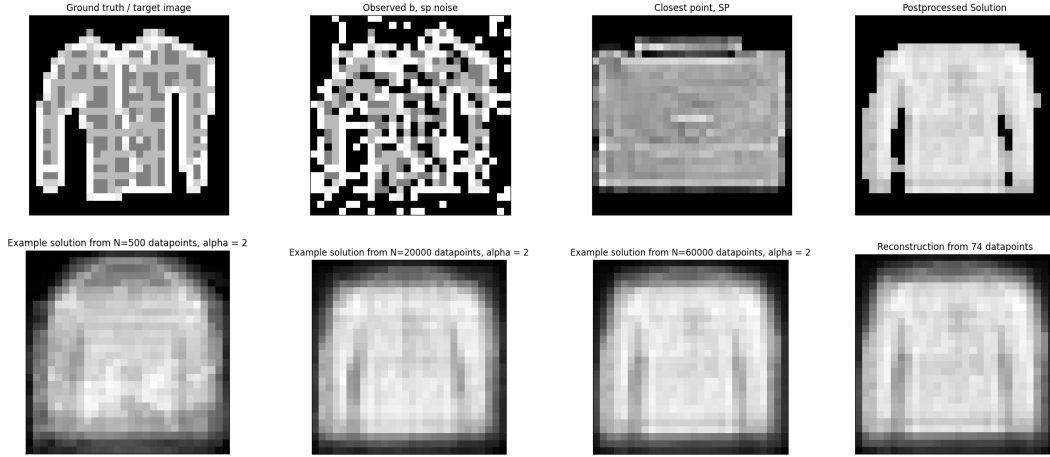


Figure 8: Recovery of a hand drawn shirt with salt and pepper corruption noise.

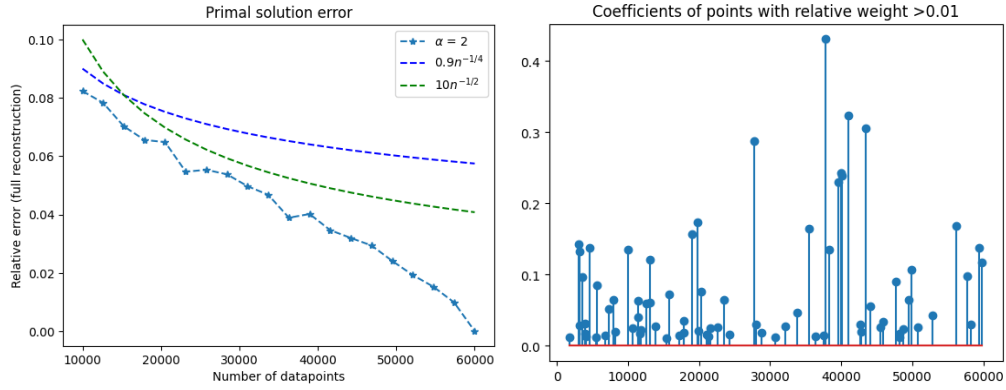


Figure 9: Rates and thresholding of optimal measure Q_n , for Figure 8

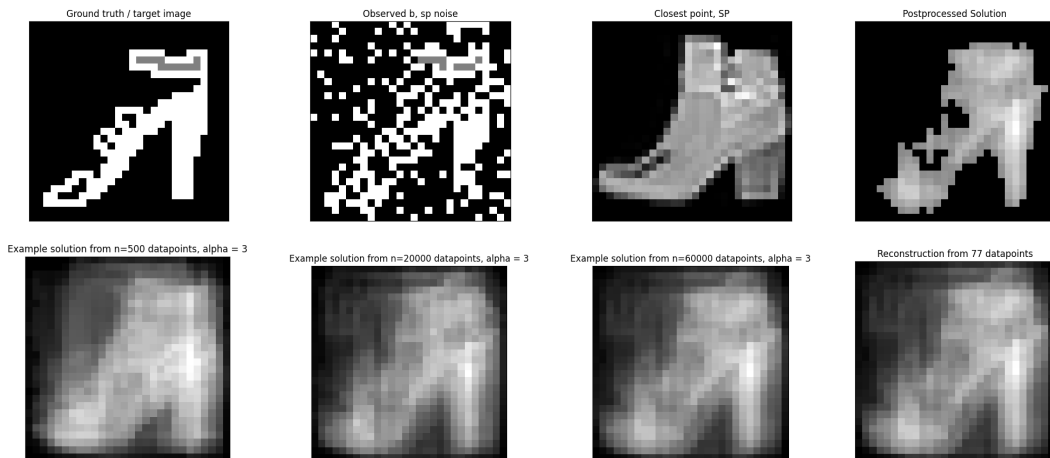


Figure 12: Recovery of hand drawn heel with Salt and Pepper corruption.

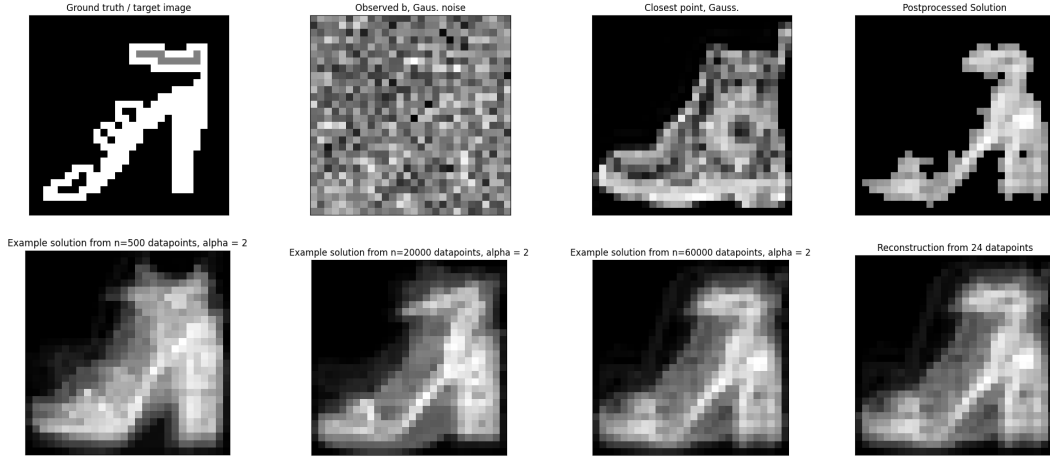


Figure 10: Recovery hand drawn heel with additive Gaussian noise.

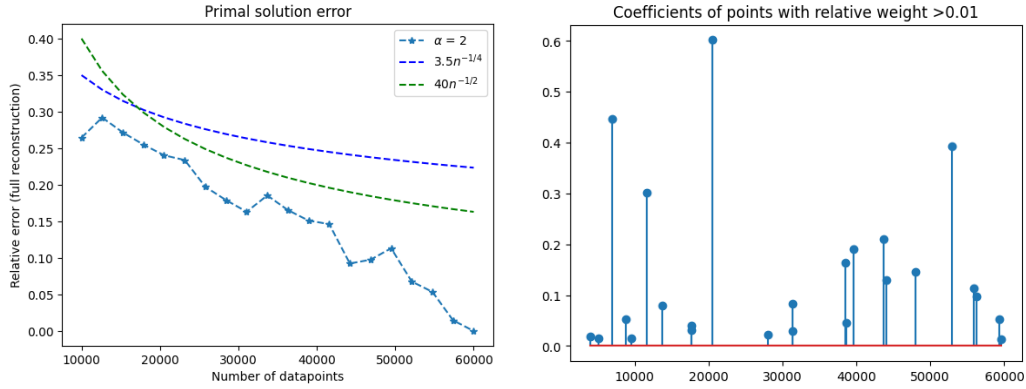


Figure 11: Rates and thresholding of optimal measure Q_n , for Figure 10

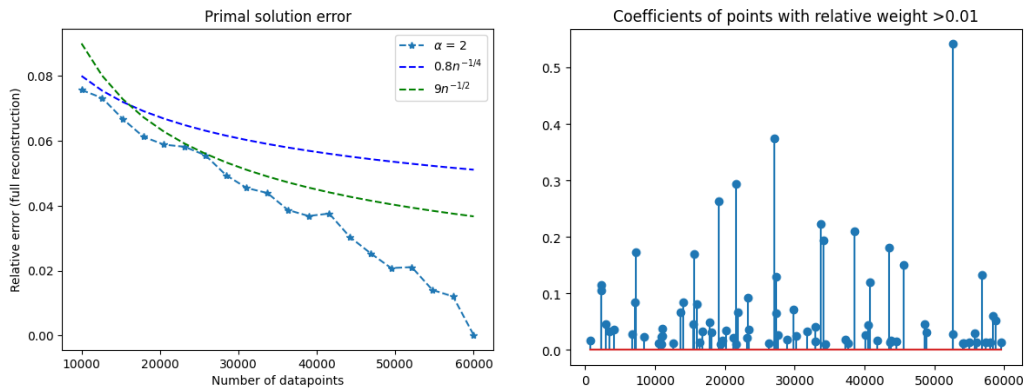


Figure 13: Rates and thresholding of optimal measure Q_n , for Figure 12

765 **7. Conclusion and Future Work.** Using the MEM framework, we have proposed using empirical
766 cal priors $\mu_n^{(\omega)}$ derived from data to approximate the unknown solution $\bar{x}_\mu$. We have shown that this
767 method has desirable theoretical properties, with Theorem 3.9 proving an almost sure convergence
768 of $\bar{x}_{n,\varepsilon} \rightarrow \bar{x}_\mu$, while Theorem 4.7 and Theorem 5.5 give upper bounds on the error $\|x_{\nu,\varepsilon} - x_\mu\|$ for an

769 arbitrary prior ν in terms of epigraphical distances and a parametric rate $\|x_{n,\varepsilon} - x_\mu\| = O\left(\frac{1}{n^{1/4}}\right)$.
770 One advantage of our approach here is that most of our results are valid not only for minimizers
771 but also for ϵ -minimizers. We conjecture, however, that at least for exact minimizers ($\epsilon = 0$) the rate
772 result is not sharp and can be improved to $O(\frac{1}{\sqrt{n}})$ - as one might suspect from a function-valued
773 central limit theorem.

774 As proof of concept, we numerically demonstrated the success of our method for denoising
775 ($C = I$) based upon certain standard test data sets. These experiments were performed with
776 modest computing resources and off-the-shelf optimization routines. In future work, our aim is to
777 design specialized optimization routines that would give us access to moderate and large regimes in
778 n and hence open the door to larger data sets. To this end, it would also be useful to reformulate
779 our dual problem so that one can take advantage of stochastic gradient descent. In addition to
780 denoising, it would be natural to present experiments for tasks such as deconvolution and inpainting.

781 Although we focused on the empirical approximation $\mu_n^{(\omega)}$, it would be natural to investigate
782 more sophisticated approximations of μ . This is particularly true given that the approach taken
783 here has been via epigraphical distances between the respective LMGFs.

784 Finally, in a certain sense this paper aims to learn a *regularizer* based upon a data set. It would
785 also be interesting to approach the *fidelity* term, which can be understood as the MEM estimator
786 based on a noise distribution (see [44]), in a similar vein. That is, can one learn the fidelity norm
787 based upon samples of noise?

- [1] C. AMBLARD, E. LAPALME, AND J.-M. LINA, *Biomagnetic source detection by maximum entropy and graphical models*, IEEE Trans. Biomed. Eng., 51 (2004), pp. 427–442.
- [2] H. ATTOUCH, *Convergence de fonctions convexes, des sous-différentiels et semi-groupes associés*, C.R. Acad. Sci. Paris, 284 (1977), pp. 539–542.
- [3] H. ATTOUCH, *Variational Convergence for Functions and Operators*, Pitman Advanced Pub. Program, Boston, 1984.
- [4] A. AUSLENDER AND M. TEBoulLE, *Interior gradient and proximal methods for convex and conic optimization*, SIAM J. Optim., 16 (2006), pp. 697–725.
- [5] H. BAUSCHKE AND P. COMBETTES, *Convex Analysis and Monotone Operator Theory in Hilbert Spaces.*, Springer, New York, 2019.
- [6] P. BILLINGSLEY, *Probability and Measure*, John Wiley & Sons, Hoboken, NJ, 2017.
- [7] L. D. BROWN, *Fundamentals of Statistical Exponential Families: with Applications in Statistical Decision Theory*, Institute of Mathematical Statistics, Hayward, California, 1986.
- [8] J. BURKE AND T. HOHEISEL, *A note on epi-convergence of sums under the inf-addition rule*, Optimization Online, (2015).
- [9] R. H. BYRD, J. NOCEDAL, AND R. B. SCHNABEL, *Representations of quasi-newton matrices and their use in limited memory methods*, Math. Program., 63 (1994), pp. 129–156.
- [10] Z. CAI, A. MACHADO, R. A. CHOWDHURY, A. SPILKIN, T. VINCENT, Ü. AYDIN, G. PELLEGRINO, J.-M. LINA, AND C. GROVA, *Diffuse optical reconstructions of functional near infrared spectroscopy data using maximum entropy on the mean*, Scientific reports, 12 (2022). 2316.
- [11] R. A. CHOWDHURY, J. M. LINA, E. KOBAYASHI, AND C. GROVA, *MEG source localization of spatially extended generators of epileptic activity: comparing entropic and hierarchical bayesian approaches*, PloS one, 8 (2013). E55969.
- [12] S. CSÖRGÖ, *Kernel-transformed empirical processes*, J. Multivariate Anal., 13 (1983), pp. 517–533.
- [13] D. DACUNHA-CASTELLE AND F. GAMBOA, *Maximum d’entropie et problème des moments*, Ann. Inst. Henri Poincaré, Probab. Stat., 26 (1990), pp. 567–596.
- [14] A. DEN DEKKER AND J. SIJBERS, *Data distributions in magnetic resonance images: A review*, Phys. Med., 30 (2014), pp. 725–741.
- [15] L. DENG, *The MNIST database of handwritten digit images for machine learning research [best of the web]*, IEEE Signal Process. Mag., 29 (2012), pp. 141–142.
- [16] M. D. DONSKER AND S. S. VARADHAN, *Asymptotic evaluation of certain Markov process expectations for large time. III*, Comm. Pure. Appl. Math., 29 (1976), pp. 389–461.
- [17] A. FERMIN, J.-M. LOUBES, AND C. LUDENA, *Bayesian methods for a particular inverse problem seismic tomography*, Int. J. Tomogr. Stat, 4 (2006), pp. 1–19.
- [18] A. FEUERVERGER, *On the empirical saddlepoint approximation*, Biometrika, 76 (1989), pp. 457–464.
- [19] G. FOLLAND, *Real Analysis: Modern Techniques and their Applications*, vol. 40, John Wiley & Sons, Hoboken, NJ, 1999.
- [20] F. GAMBOA, *Méthode du Maximum d’Entropie sur la Moyenne et Applications*, PhD thesis, Université Paris-Sud, 1989.
- [21] C. GROVA, J. DAUNIZEAU, J.-M. LINA, C. G. BÉNAR, H. BENALI, AND J. GOTMAN, *Evaluation of EEG localization methods using realistic simulations of interictal spikes*, Neuroimage, 29 (2006), pp. 734–753.
- [22] M. HEERS, R. A. CHOWDHURY, T. HEDRICH, F. DUBEAU, J. A. HALL, J.-M. LINA, C. GROVA, AND E. KOBAYASHI, *Localization accuracy of distributed inverse solutions for electric and magnetic source imaging of interictal epileptic discharges in patients with focal epilepsy*, Brain Topography, 29 (2016), pp. 162–181.
- [23] E. T. JAYNES, *Information theory and statistical mechanics*, Phys. Rev., 106 (1957), pp. 620–630.
- [24] E. T. JAYNES, *Information theory and statistical mechanics. II*, Phys. Rev., 108 (1957), pp. 171–190.
- [25] N. KANTAS ET AL., *On Particle Methods for Parameter Estimation in State-Space Models*, Statist. Sci., 30 (2015), pp. 328 – 351.
- [26] A. J. KING AND R. WETS, *Epi-consistency of convex stochastic programs*, Stoch. and Stoch. Rep., 34 (1991), pp. 83–92.
- [27] A. KLENKE, *Probability Theory: a Comprehensive Course*, Springer, Cham, Switzerland, 2013.
- [28] S. KULLBACK AND R. LEIBLER, *On information and sufficiency*, Ann. Math. Stat., 22 (1951), pp. 79–86.
- [29] G. LE BESNERAIS, J.-F. BERCHER, AND G. DEMOMENT, *A new look at entropy for solving linear inverse problems*, IEEE Trans. Inform. Theory, 45 (1999), pp. 1565–1578.
- [30] P. MARÉCHAL AND A. LANNES, *Unification of some deterministic and probabilistic methods for the solution of linear inverse problems via the principle of maximum entropy on the mean*, Inverse Problems, 13 (1997).
- [31] J. NAVAZA, *On the maximum-entropy estimate of the electron density function*, Acta Crystallogr. Sect. A, 41 (1985), pp. 232–244.
- [32] J. NAVAZA, *The use of non-local constraints in maximum-entropy electron density reconstruction*, Acta Crystallogr. Sect. A, 42 (1986), pp. 212–223.
- [33] E. RIETSCH ET AL., *The maximum entropy approach to inverse problems-spectral analysis of short data records*

- 850 *and density structure of the earth*, J. Geophys., 42 (1976), pp. 489–506.
- 851 [34] G. RIOUX, R. CHOKSI, T. HOHEISEL, P. MARÉCHAL, AND C. SCARVELIS, *The maximum entropy on the mean*
852 *method for image deblurring*, Inverse Problems, 37 (2020). 015011.
- 853 [35] G. RIOUX, C. SCARVELIS, R. CHOKSI, T. HOHEISEL, AND P. MARÉCHAL, *Blind deblurring of barcodes via*
854 *Kullback-Leibler divergence*, IEEE Trans. on Pattern Anal. Mach. Intell., 43 (2021), pp. 77–88.
- 855 [36] R. T. ROCKAFELLAR, *Convex Analysis*, Princeton University Press, Princeton, NJ, 1997.
- 856 [37] R. T. ROCKAFELLAR AND R. WETS, *Variational Analysis*, Springer, Berlin, 1998.
- 857 [38] R. T. ROCKAFELLAR AND R. WETS, *Variational systems, an introduction*, in Multifunctions and Integrands:
858 Stochastic Analysis, Approximation and Optimization Proceedings of a Conference held in Catania, Italy,
859 June 7–16, 1983, Springer, 2006, pp. 1–54.
- 860 [39] W. RÖMISCH AND R. WETS, *Stability of ε -approximate solutions to convex stochastic programs*, SIAM J. Optim,
861 18 (2007), pp. 961–979.
- 862 [40] J. ROYSET AND R. WETS, *An Optimization Primer*, Springer, Cham, Switzerland, 2021.
- 863 [41] T. SEVERINI, *Elements of Distribution Theory*, Cambridge University Press, Cambridge, UK, 2005.
- 864 [42] A. TORRALBA AND A. OLIVA, *Statistics of natural image categories*, Network: Computation in Neural Systems,
865 14 (2003). 391.
- 866 [43] B. URBAN, *Retrieval of atmospheric thermodynamical parameters using satellite measurements with a maximum*
867 *entropy method*, Inverse Problems, 12 (1996). 779.
- 868 [44] Y. VAISBOURD, R. CHOKSI, A. GOODWIN, T. HOHEISEL, AND C.-B. SCHÖNLIEB, *Maximum entropy on the*
869 *mean and the cramér rate function in statistical estimation and inverse problems: properties, models, and*
870 *algorithms*, Math. Program, in press, (2025).
- 871 [45] A. VAN DER VAART, *New donsker classes*, Ann. Probab., 24 (1996), pp. 2128–2140.
- 872 [46] A. VAN DER VAART AND J. WELLNER, *Weak Convergence*, Springer, 1996.
- 873 [47] H. XIAO, K. RASUL, AND R. VOLLGRAF, *Fashion-MNIST: a Novel Image Dataset for Benchmarking Machine*
874 *Learning Algorithms*. 2017, <https://arxiv.org/abs/cs.LG/1708.07747>.
- 875 [48] C. ZĂLINESCU, *Convex Analysis in General Vector Spaces*, World Scientific, Singapore, 2002.