



Solving Inverse Problems by Joint Posterior Maximization with Autoencoding Prior

Mario González, Andrés Almansa, Pauline Tan

► To cite this version:

Mario González, Andrés Almansa, Pauline Tan. Solving Inverse Problems by Joint Posterior Maximization with Autoencoding Prior. SIAM Journal on Imaging Sciences, 2022, 15 (2), pp.822-859. 10.1137/21M140225X . hal-03151455v4

HAL Id: hal-03151455

<https://hal.science/hal-03151455v4>

Submitted on 26 Apr 2022

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

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

Solving Inverse Problems by Joint Posterior Maximization with Autoencoding Prior*

Mario González[†], Andrés Almansa[‡], and Pauline Tan[§]

Abstract. In this work we address the problem of solving ill-posed inverse problems in imaging where the prior is a variational autoencoder (VAE). Specifically we consider the decoupled case where the prior is trained once and can be reused for many different log-concave degradation models without retraining. Whereas previous MAP-based approaches to this problem lead to highly non-convex optimization algorithms, our approach computes the joint (space-latent) MAP that naturally leads to alternate optimization algorithms and to the use of a stochastic encoder to accelerate computations. The resulting technique (JPMAP) performs Joint Posterior Maximization using an Autoencoding Prior. We show theoretical and experimental evidence that the proposed objective function is quite close to bi-convex. Indeed it satisfies a weak bi-convexity property which is sufficient to guarantee that our optimization scheme converges to a stationary point. We also highlight the importance of correctly training the VAE using a denoising criterion, in order to ensure that the encoder generalizes well to out-of-distribution images, without affecting the quality of the generative model. This simple modification is key to providing robustness to the whole procedure. Finally we show how our joint MAP methodology relates to more common MAP approaches, and we propose a continuation scheme that makes use of our JPMAP algorithm to provide more robust MAP estimates. Experimental results also show the higher quality of the solutions obtained by our JPMAP approach with respect to other non-convex MAP approaches which more often get stuck in spurious local optima.

Key words. Image Restoration, Inverse Problems, Bi-convex Optimization, Bayesian Statistics, Generative Models, Variational Auto-encoders

AMS subject classifications. 68U10, 65K10, 65D18, 68T05, 90C26, 90C25, 90C30,

1. Introduction. General inverse problems in imaging consist in estimating a clean image $\mathbf{x} \in \mathbb{R}^d$ from noisy, degraded measurements $\mathbf{y} \in \mathbb{R}^m$. In many cases the degradation model is known and its conditional density

$$p_{Y|X}(\mathbf{y} | \mathbf{x}) \propto e^{-F(\mathbf{x}, \mathbf{y})}$$

*Submitted to the editors on March 24th 2021. Revised version February 3rd 2022. The authors would like to sincerely thank Mauricio Delbracio, José Lezama and Pablo Musé for their help, their insightful comments, and their continuous support throughout this project.

Funding: This work was funded by ECOS Sud Project U17E04, by the French-Uruguayan Institute of Mathematics and Interactions (IFUMI), by CSIC I+D 2018-256 (Uruguay) and by ANII (Uruguay) under Grant 11 FCE_1.2017.1.135458 and by the French Research Agency through the PostProdLEAP project (ANR-19-CE23-0027-01). Computer experiments for this work ran on a Titan Xp GPU donated by NVIDIA, as well as on HPC resources from GENCI-IDRIS (Grants 2020-AD011011641 and 2021-AD011011641R1).

[†]DMEL, CenUR RN, Universidad de la República, Salto, Uruguay (mgonzalez@unorte.edu.uy, <http://dmel.interior.edu.uy/mario-gonzalez/>).

[‡]MAP5, CNRS & Université Paris Cité, France (andres.almansa@parisdescartes.fr).

[§]LJLL, Sorbonne Université, Paris, France (pauline.tan@sorbonne-universite.fr)

is log-concave with respect to $\mathbf{x}$. To illustrate this, let us consider the case where the negative log-conditional is quadratic with respect to $\mathbf{x}$

$$(1.1) \quad F(\mathbf{x}, \mathbf{y}) = \frac{1}{2\sigma^2} \|\mathbf{A}\mathbf{x} - \mathbf{y}\|^2.$$

This boils down to a linear degradation model that takes into account degradations such as, white Gaussian noise, blur, and missing pixels. When the degradation operator $\mathbf{A}$ is non-invertible or ill-conditioned, or when the noise level σ is high, obtaining a good estimate of $\mathbf{x}$ requires prior knowledge on the image, given by $p_X(\mathbf{x}) \propto e^{-\lambda U(\mathbf{x})}$. Variational and Bayesian methods in imaging are extensively used to derive MMSE or MAP estimators,

$$(1.2) \quad \hat{\mathbf{x}}_{\text{MAP}} = \arg \max_{\mathbf{x}} p_{X|Y}(\mathbf{x} | \mathbf{y}) = \arg \min_{\mathbf{x}} \{F(\mathbf{x}, \mathbf{y}) + \lambda U(\mathbf{x})\}$$

based on (a) explicit hand-crafted priors like Tikhonov regularization [67], total variation [59, 10, 43, 50] and its higher order [7] and non-local [25] variants, sparsity in a transformed domain [21], or in redundant representations like wavelet frames or patch dictionaries [22], or (b) learning-based priors like patch-based Gaussian mixture models [80, 74, 65].

Neural network regression. Since deep neural networks (NN) showed their superiority in image classification tasks [37] researchers started to look for ways to use this tool to solve inverse problems too. The most straightforward attempts employed neural networks as *regressors* to learn a risk minimizing mapping $\mathbf{y} \mapsto \mathbf{x}$ from many examples $(\mathbf{x}_i, \mathbf{y}_i)$ either agnostically [20, 75, 77, 24, 62, 23] or including the degradation model in the network architecture via unrolled optimization techniques [28, 12, 18, 26].

Implicitly decoupled priors. The main drawback of neural networks regression is that they require to retrain the neural network each time a single parameter of the degradation model changes. To avoid the need for retraining, another family of approaches seek to *decouple* the NN-based learned image prior from the degradation model. A popular approach within this methodology are *Plug & Play* (or PnP) methods. Instead of directly learning the log-prior $-\log p_X(\mathbf{x}) = U(\mathbf{x}) + C$, these methods seek to learn an approximation of its gradient ∇U [5, 4] or proximal operator prox_U [71, 45, 76, 11, 34, 60], by replacing it by a denoising NN. Then, these approximations are used in an iterative optimization algorithm to find the corresponding MAP estimator in Equation (1.2) or more generally some sort of consensus equilibrium among the data fitting term and the priors [9].

Taking an apparently different approach Romano *et al.* introduced the regularization by denoising (RED) algorithm [58] which uses a denoiser D_σ to construct an explicit regularizer $U(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T (\mathbf{x} - D_\sigma(\mathbf{x}))$. Under certain conditions (see below) its gradient $\nabla U = I - D_\sigma$ can be conveniently computed in terms of the denoiser, leading to a gradient descent scheme for the associated MAP estimator, which is very easy to implement.

Explicitly decoupled generative priors. In another series of works pioneered by Bora *et al.* [6] and followed by [63, 56, 46, 32, 30] the Plug & Play prior is explicitly provided by a generative model, most often a generative adversarial network G that maps a latent variable $\mathbf{z} \sim \mathcal{N}(0, I)$ to an image $\mathbf{x} = G(\mathbf{z})$ with the desired distribution p_X as represented by the learning dataset. More precisely these methods solve an optimization problem on the latent variable $\mathbf{z}$

$$(1.3) \quad \hat{\mathbf{z}}_{\text{MAP}} = \arg \min_{\mathbf{z}} \left\{ F(G(\mathbf{z}), \mathbf{y}) + \frac{1}{2} \alpha \|\mathbf{z}\|^2 \right\}$$

and the reconstructed image is provided by $\hat{\mathbf{x}}_{\text{MAP}} = \mathbf{G}(\hat{\mathbf{z}}_{\text{MAP}})$. As we show in the following sub-section and in appendix B.1, this corresponds (when $\alpha = 1$) to the Maximum A Posteriori (MAP) estimator with respect to the $\mathbf{z}$ variable. In this work we adopt this framework with some extensions that help avoid getting trapped in spurious critical points of the non-convex objective function.

Empirical success of Plug & Play and RED. Plug & Play and RED approaches became very popular because they allow to repurpose very powerful state of the art denoisers as regularizers of a large family of inverse problems in a quite straightforward manner. They have been successfully applied to many different problems in imaging and they have thus empirically proven their superiority (in terms of achievable reconstruction quality with respect to more classical regularization techniques), and opened the way for the solution of more difficult inverse problems in imaging.

Theoretical questions. The success of Plug & Play and RED approaches largely outpaced our understanding of why and when these techniques lead to algorithms that provably converge to well-posed statistical estimators with well known properties. This is not surprising because obtaining convergence guarantees for non-convex optimization problems under realistic conditions is quite challenging.

A notable exception where strong convergence results have been obtained is the particular case of compressed sensing, where the rows of the degradation operator (or sensing matrix) $\mathbf{A}$ are independent realizations of a zero-mean Gaussian distribution. For this problem, Hand *et al.* [30, 32] show that the optimization objective (1.3) has almost no spurious stationary points when the generator is assumed to be a random ReLU network with Gaussian weights. As a consequence, a minor modification to the gradient descent algorithm in [6] converges with high probability to the global optimum.

In this paper we are interested in more general inverse problems, where the sensing matrix $\mathbf{A}$ is not necessarily random but deterministic and highly structured most often dictated by our modeling of the acquisition device. In this more general setting the hypotheses of the CS results are not necessarily satisfied, and the kind of convergence guarantees that could be established for PnP algorithms with non-convex priors are much weaker (typically only convergence to a stationary point or fixed point is provided, not necessarily a global optimum), and most works concentrate in the implicit case, where the prior is not explicitly provided by a generative model, but implicit in a denoising algorithm.

In such a case the actual prior is unknown, the existence of a density whose gradient or proximal operator is well approximated by a neural denoiser is most often not guaranteed [57], and the convergence of the algorithm is not guaranteed either unless the denoiser is retrained with specific constraints like idempotence [29, 63], contractive residual [60] or exact, invertible, smooth MMSE denoisers [73].

The effect of such training constraints on the quality of the denoisers and the associated priors is yet to be explored in detail. But even when these constraints are satisfied, convergence conditions can be quite restrictive, either (a) requiring the data-fitting term F to be strongly convex [60] (thus excluding many important problems in computational imaging where $\mathbf{A}$ is not full rank like interpolation, super-resolution, deconvolution with a non-invertible kernel or compressive sensing), and/or (b) constraining the regularization parameter λ outside of its

useful range [60, 73].¹

Similarly, an early analysis of the RED approach [57] provides a convergence proof, but only under quite restrictive conditions (locally homogeneous denoisers with symmetric Jacobian) which exclude most state of the art denoisers like DnCNN, BM3D, NLMeans. A more recent analysis of a stochastic variant of the RED algorithm [39] (called PnP-SGD) significantly expands the family of denoisers that provide convergence guarantees, including DnCNN and the doubly-stochastic variant of NLM [64]. These guarantees come, however, at the expense of a very small descent step which leads to a very computationally expensive algorithm with slow convergence. In addition, the experiments with PnP-SGD show that this algorithm is extremely sensitive to the initial condition, and it can be easily get stuck on spurious local minima if not initialized very carefully.

Focus of this work. Very recent works focused on developing MAP estimation algorithms with convergence guarantees under more realistic conditions. The convergence analysis of the RED framework, and its RED-PRO variant was further refined under a demicontractive condition for the denoiser [15]. This condition is, however, difficult to verify according to Pesquet et al. [51] who provides an alternative convergence analysis framework based on firmly non-expansive denoisers for which explicit training procedures exist [66]. In this work we explore alternative new ways to bring theory and practice closer together, by proposing novel Plug & Play algorithms to compute the MAP estimator of an inverse problem with a neural regularizer. Unlike previous approaches which were based on implicit priors, or on GAN-based explicit priors, our approach is based on an explicit generative prior that has been trained as a Variational AutoEncoder (VAE). As we shall see later, the additional VAE structure provides: (i) powerful mechanisms to avoid getting stuck in spurious local minima of the associated non-convex functional, and (ii) convergence guarantees under much less restrictive conditions on the inverse problem F and regularization parameter λ .

The next Section 1.1 reviews previous work on similar approaches to compute a MAP estimator from a generative prior that was trained either as a VAE or a GAN. Section 1.2 briefly introduces our approach and how it relates to previous work. The section finishes with an overview of the rest of the paper.

1.1. Maximum a Posteriori meets Generative Models. Our approach focuses on PnP algorithms where the prior is provided by a generative model. For instance one could use a generative adversarial network (GAN) to learn a generative model for $X = G(Z)$ with $Z \sim N(0, I)$ a latent variable. The generative model induces a prior on X via the push-forward measure $p_X = G\#p_Z$, which following [48, section 5] can be developed as

$$p_X(\mathbf{x}) = \frac{p_Z(G^{-1}(\mathbf{x}))}{\sqrt{\det S(G^{-1}(\mathbf{x}))}} \delta_{\mathcal{M}}(\mathbf{x})$$

where $S = \left(\frac{\partial G}{\partial \mathbf{z}}\right)^T \left(\frac{\partial G}{\partial \mathbf{z}}\right)$ is the squared Jacobian and the manifold $\mathcal{M} = \{\mathbf{x} : \exists \mathbf{z}, \mathbf{x} = G(\mathbf{z})\}$ represents the image of the generator G . With such a prior p_X , the $\mathbf{x}$ -optimization (1.2)

¹In [39] the PnP-ADMM and PnP-FBS algorithms introduced in [60, 73] are reported to converge in practice quite far beyond the conditions of the theorem, but require (to obtain optimal performance) the regularization parameter λ to be tuned to values that are far outside the region where convergence is guaranteed. Also, the performance is significantly degraded if λ is constrained to the range where convergence is guaranteed.

required to obtain $\hat{\mathbf{x}}_{\text{MAP}}$ becomes intractable (in general), for various reasons:

- the computation of $\det S$,
- the inversion of $\mathbf{G}$, and
- the hard constraint $\mathbf{x} \in \mathcal{M}$.

These operations are all memory and/or computationally intensive, except when they are partially addressed by the use of a normalizing flow like in [31, 72].

Current attempts to use such a generative model as a prior, like the one proposed by Bora et al. [6] for GANs, circumvent these difficulties by performing an optimization on $\mathbf{z}$ (in the latent domain) instead of $\mathbf{x}$. Instead of solving Equation (1.2), they solve

$$(1.4) \quad \begin{aligned} \hat{\mathbf{z}}_{\text{MAP}} &= \arg \max_{\mathbf{z}} \{p_{Y|X}(\mathbf{y} | \mathbf{G}(\mathbf{z})) p_Z(\mathbf{z})\} \\ &= \arg \min_{\mathbf{z}} \left\{ F(\mathbf{G}(\mathbf{z}), \mathbf{y}) + \frac{1}{2} \|\mathbf{z}\|^2 \right\}, \end{aligned}$$

by assuming a standard Gaussian prior. This problem is much more tractable, and the corresponding $\mathbf{x}$ -estimate is obtained as

$$(1.5) \quad \hat{\mathbf{x}}_{\text{MAP}-\mathbf{z}} = \mathbf{G}(\hat{\mathbf{z}}_{\text{MAP}}).$$

As we show in appendix B.1, this new estimator does not necessarily coincide with $\hat{\mathbf{x}}_{\text{MAP}}$ but it does correspond to the MAP-estimator of $\mathbf{x}$ after the change of variable $\mathbf{x} = \mathbf{G}(\mathbf{z})$, namely

$$\hat{\mathbf{x}}_{\text{MAP}-\mathbf{z}} = \mathbf{G} \left(\arg \max_{\mathbf{z}} \{p_{Z|Y}(\mathbf{z} | \mathbf{y})\} \right).$$

Since $\mathbf{G}$ is non-linear, this problem (or its equivalent formulation (1.4)) is highly non-convex and difficult to solve with global optimality guarantees. Nevertheless, in the particular case where $\mathbf{A}$ is a random Gaussian matrix (compressed sensing case) or when F is strongly convex, recent work shows that the global optimum can be reached with linear convergence rates by a small modification of a gradient descent algorithm [32, 30], or by an ADMM algorithm with non-linear constraints [38, 2, 70]. To the best of our knowledge, these results do not extend, however, to the more general case we are interested in here, where $\mathbf{A}$ is deterministic and rank-deficient, and F is consequently not strongly convex. In this more general setting, convergence guarantees for this optimization problem remain extremely difficult to establish, as confirmed by experimental results presented in Section 3.

A common technique to solve difficult optimization problems like the one in Equation (1.4) is to use (Half Quadratic) splitting methods

$$(1.6) \quad \hat{\mathbf{x}}_{\beta} = \arg \min_{\mathbf{x}} \min_{\mathbf{z}} \underbrace{\left\{ F(\mathbf{x}, \mathbf{y}) + \frac{\beta}{2} \|\mathbf{x} - \mathbf{G}(\mathbf{z})\|^2 + \frac{1}{2} \|\mathbf{z}\|^2 \right\}}_{J_{1,\beta}(\mathbf{x}, \mathbf{z})}$$

combined with a continuation scheme, namely:

$$(1.7) \quad \hat{\mathbf{x}}_{\text{MAP}-\mathbf{z}} = \lim_{\beta \rightarrow \infty} \hat{\mathbf{x}}_{\beta}.$$

The convergence of the continuation scheme in the last line is a standard result in Γ -convergence (see [17] and appendix C). The corresponding splitting algorithm is presented in Algorithm 1.1.

Algorithm 1.1 MAP- $\mathbf{z}$ splitting

Require: Measurements $\mathbf{y}$, Initial condition $\mathbf{x}_0$, maxiter, $k_{\max}$, $\{\beta_0, \dots, \beta_{k_{\max}}\}$

Ensure: $\hat{\mathbf{x}} = \mathbf{G}(\arg \max_{\mathbf{z}} p_{Z|Y}(\mathbf{z} | \mathbf{y}))$

```

1: for  $k := 0$  to  $k_{\max}$  do
2:    $\beta := \beta_k$ 
3:   for  $n := 0$  to maxiter do
4:      $\mathbf{z}_{n+1} := \arg \min_{\mathbf{z}} J_{1,\beta}(\mathbf{x}_n, \mathbf{z})$  // Nonconvex
5:      $\mathbf{x}_{n+1} := \arg \min_{\mathbf{x}} J_{1,\beta}(\mathbf{x}, \mathbf{z}_{n+1})$  // Quadratic
6:   end for
7:    $\mathbf{x}_0 := \mathbf{x}_{n+1}$ 
8: end for
9: return  $\mathbf{x}_{n+1}$ 

```

However, unlike most cases of HQS which include a linear constraint between the two variables, this splitting algorithm still contains (line 4) a difficult non-convex optimization problem².

1.2. Proposed method: Joint MAP $_{\mathbf{x},\mathbf{z}}$. In this work we propose to address this challenge by substituting the difficult non-convex sub-problem by a local quadratic approximation provided by the encoder of a variational autoencoder.

Indeed, as we show in Section 2, a variational autoencoder allows to interpret the splitting Equation (1.6) as the negative logarithm of the joint posterior density $p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$. Therefore, solving Equation (1.6) amounts to compute a joint MAP $_{\mathbf{x},\mathbf{z}}$ estimator that we denote by $\hat{\mathbf{x}}_{\text{MAP}_{\mathbf{x},\mathbf{z}}}^\beta$. Moreover if the same joint conditional density $p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$ is decomposed in a different manner, it leads to an approximate expression that makes use of the encoder, and is quadratic in $\mathbf{z}$. If this approximation is good enough then the maximization of the joint log-posterior becomes a bi-concave optimization problem or approximately so. And in that case, an extension of standard bi-convex optimization results [27] shows that the algorithm converges to a stationary point.

We also highlight the importance of correctly training the VAE in such a way that the encoder generalizes well to noisy values of $\mathbf{x}$ outside of the support of $p_X(\mathbf{x})$. This can be achieved by training the VAE to reconstruct their clean inputs with noise injected at the input level, as proposed by Im et al. [33]. We observe that this modified training does not degrade the quality of the generative model, but makes our quasi-bi-convex optimization procedure much more robust.

Finally we show that a continuation scheme allows to obtain the MAP $_{\mathbf{z}}$ estimator as the limit of a series of joint MAP $_{\mathbf{x},\mathbf{z}}$ optimizations. This continuation scheme, in addition to the quasi-bi-convex optimization, and the initialisation heuristic provided by the denoising encoder

²In another context a primal-dual optimization algorithm was proposed to solve a similar optimization problem [2], but this approach was not explored in the context where $\mathbf{G}$ is a generative model.

leads to a much more robust non-convex optimization scheme which more often converges to the right critical point than a straightforward gradient descent of the $\text{MAP}_{\mathbf{z}}$ model.

The remainder of this paper is organized as follows. In Section 2 we derive a model for the joint conditional posterior distribution of space and latent variables $\mathbf{x}$ and $\mathbf{z}$, given the observation $\mathbf{y}$. This model makes use of a generative model, more precisely a VAE with Gaussian decoder. We then propose an alternate optimization scheme to maximize the joint posterior model, and state convergence guarantees. Section 3 presents first a set of experiments that illustrates the convergence properties of the optimization scheme. We then test our approach on classical image inverse problems, and compare its performance with state-of-the-art methods. Concluding remarks are presented in Section 4.

2. From Variational Autoencoders to Joint Posterior Maximization. Recently, some generative models based on neural networks have shown their capability to approximate the complex image distribution in a data-driven fashion. In particular, *Variational Autoencoders (VAE)* [36] combine variational inference to approximate unknown posterior distributions of latent variable models with the ability of neural networks to learn such approximations.

Consider a graphical model $\mathbf{z} \rightarrow \mathbf{x}$ in which we assume that a latent variable $\mathbf{z}$ is responsible of the observed image $\mathbf{x}$. For example, in an image of a handwritten digit we can imagine which digit is represented in the image, width, angle (and so on) as latent variables. We choose a generative model

$$p_{\theta}(\mathbf{x}, \mathbf{z}) = p_{\theta}(\mathbf{x}|\mathbf{z})p_{\mathbf{Z}}(\mathbf{z})$$

where $p_{\mathbf{Z}}(\mathbf{z})$ is some simple distribution (which we can easily sample from) and $p_{\theta}(\mathbf{x}|\mathbf{z})$ is the approximation of the probability distribution of $\mathbf{x}$ given $\mathbf{z}$ parameterized by a neural network (with weights θ) known as *stochastic decoder*.

The intractability of $p_{\theta}(\mathbf{x}) = \int p_{\theta}(\mathbf{x}|\mathbf{z})p_{\mathbf{Z}}(\mathbf{z}) d\mathbf{z}$ is related to the posterior distribution $p_{\theta}(\mathbf{z}|\mathbf{x})$ by

$$(2.1) \quad p_{\theta}(\mathbf{z}|\mathbf{x}) = \frac{p_{\theta}(\mathbf{x}|\mathbf{z})p_{\mathbf{Z}}(\mathbf{z})}{p_{\theta}(\mathbf{x})}.$$

The *variational inference* approach consists in approximating this posterior with another model $q_{\phi}(\mathbf{z}|\mathbf{x})$ which, in our case, is another neural network with parameters ϕ , called a *stochastic encoder*.

Following [36], we consider the *Evidence Lower BOund (ELBO)* as

$$(2.2) \quad \mathcal{L}_{\theta,\phi}(\mathbf{x}) := \log p_{\theta}(\mathbf{x}) - KL(q_{\phi}(\mathbf{z}|\mathbf{x}) || p_{\theta}(\mathbf{z}|\mathbf{x})) \leq \log p_{\theta}(\mathbf{x})$$

where KL is the Kullback-Leibler divergence. Thus, given a dataset $\mathcal{D} = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}$ of image samples, maximizing the averaged ELBO on $\mathcal{D}$ means maximizing $\log p_{\theta}(\mathcal{D})$ which is the maximum likelihood estimator of weights θ and minimizing $KL(q_{\phi}(\mathbf{z}|\mathbf{x}) || p_{\theta}(\mathbf{z}|\mathbf{x}))$ which enforces the approximated posterior $q_{\phi}(\mathbf{z}|\mathbf{x})$ to be similar to the true posterior $p_{\theta}(\mathbf{z}|\mathbf{x})$.

It can be shown [36] that the ELBO can be rewritten as

$$(2.3) \quad \mathcal{L}_{\theta,\phi}(\mathbf{x}) = \mathbb{E}_{q_{\phi}(\mathbf{z}|\mathbf{x})}[\log p_{\theta}(\mathbf{x}|\mathbf{z})] - KL(q_{\phi}(\mathbf{z}|\mathbf{x}) || p_{\mathbf{Z}}(\mathbf{z})).$$

The first term in (2.3) is a *reconstruction loss* similar to the one of plain autoencoders: it enforces that the code $\mathbf{z} \sim q_\phi(\cdot|\mathbf{x})$ generated by the encoder q_ϕ can be used by the decoder p_θ to reconstruct the original input $\mathbf{x}$. The second term is a *regularization term* that enforces the distribution $q_\phi(\mathbf{z}|\mathbf{x})$ of the latent code $\mathbf{z}$ (given $\mathbf{x}$) to be close to the prior distribution $p_Z(\mathbf{z})$. It is common to choose an isotropic Gaussian as the prior distribution of the latent code:

$$p_Z(\mathbf{z}) = \mathcal{N}(\mathbf{z} | 0, I) \propto e^{-\|\mathbf{z}\|^2/2}$$

and a Gaussian encoder $q_\phi(\mathbf{z}|\mathbf{x}) = \mathcal{N}(\mathbf{z} | \mu_\phi(\mathbf{x}), \Sigma_\phi(\mathbf{x}))$, so that the KL divergence in (2.3) is straightforward to compute. For the decoder $p_\theta(\mathbf{x}|\mathbf{z})$ a Gaussian decoder is the most common choice and as we will see we benefit from that.

2.1. Learning approximations vs. encoder approximations. In this work we construct an image prior using a Variational Autoencoder (VAE). Like any machine learning tool VAEs make different kinds of approximations. Let's distinguish two types of approximations that shall be important in the sequel:

Learning approximation: The ideal prior p_X^* can only be approximated by our VAE due to its architectural constraints, finite complexity, truncated optimization algorithms, finite amount of data and possible biases in the data. Due to all these approximations, after learning we have only access to an approximate prior $p_X \approx p_X^*$. VAEs give access to this approximate prior p_X via a generative model: taking samples of a latent variable Z with known distribution $\mathcal{N}(0, I)$ in $\mathbb{R}^l$ (with $l \ll d$), and feeding these samples through a learned decoder network, we obtain samples of $X \sim p_X$. The approximate prior itself

$$(2.4) \quad p_X(\mathbf{x}) = \int p_\theta(\mathbf{x}|\mathbf{z}) p_Z(\mathbf{z}) d\mathbf{z}$$

is intractable because it requires computing an integral over all possible latent codes $\mathbf{z}$. However the approximate joint distribution is readily accessible

$$p_{X,Z}(\mathbf{x}, \mathbf{z}) = p_\theta(\mathbf{x}|\mathbf{z}) p_Z(\mathbf{z})$$

thanks to $p_{X|Z}(\mathbf{x}|\mathbf{z}) = p_\theta(\mathbf{x}|\mathbf{z})$ which is provided by the decoder network.

Encoder approximation: In the previous item we considered the VAE as a generative model without making use of the encoder network. The encoder network

$$\tilde{p}_{Z|X}(\mathbf{z}|\mathbf{x}) := q_\phi(\mathbf{z}|\mathbf{x}) \approx p_{Z|X}(\mathbf{z}|\mathbf{x})$$

is introduced as an approximate way to solve the intractability of $p_{Z|X}(\mathbf{z}|\mathbf{x}) = p_\theta(\mathbf{z}|\mathbf{x})$ (which is related to the intractability of $p_\theta(\mathbf{x})$ as observed in equation (2.1)).

Using the encoder network we can provide an alternative approximation for the joint distribution

$$\tilde{p}_{X,Z}(\mathbf{x}, \mathbf{z}) := q_\phi(\mathbf{z}|\mathbf{x}) p_X(\mathbf{x}) \approx p_{X,Z}(\mathbf{x}, \mathbf{z})$$

which shall be useful in the sequel.

Put another way, the ideal joint distribution $p_{X,Z}^*$ is inaccessible, but can be approximated in two different ways:

The first expression denoted $p_{X,Z}(\mathbf{x}, \mathbf{z})$ only uses the decoder and is only affected by the *learning approximation*

$$p_{X,Z}^*(\mathbf{x}, \mathbf{z}) \approx p_{X,Z}(\mathbf{x}, \mathbf{z}) := p_\theta(\mathbf{x}|\mathbf{z}) p_Z(\mathbf{z}).$$

The second expression denoted $\tilde{p}_{X,Z}(\mathbf{x}, \mathbf{z})$ uses both encoder and decoder and is affected both by the *learning approximation* and by the *encoder approximation*

$$p_{X,Z}(\mathbf{x}, \mathbf{z}) \approx \tilde{p}_{X,Z}(\mathbf{x}, \mathbf{z}) := q_\phi(\mathbf{z}|\mathbf{x}) p_X(\mathbf{x})$$

In the following subsection we shall forget about the ideal prior p_X^* and joint distribution $p_{X,Z}^*$ which are both inaccessible. Instead we accept p_X (with its learning approximations) as our prior model which shall guide all our estimations. The approximation symbol shall be reserved to expressions that are affected by the encoder approximation *in addition to* the learning approximation.

2.2. Variational Autoencoders as Image Priors. To obtain the Maximum a Posteriori estimator (MAP), we could plug in the approximate prior p_X in equation (1.2), but this leads to a numerically difficult problem to solve due to the intractability of p_X . Instead, we propose to maximize the joint posterior $p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$ over $(\mathbf{x}, \mathbf{z})$ which is equivalent to minimizing

$$\begin{aligned} J_1(\mathbf{x}, \mathbf{z}) &:= -\log p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y}) \\ (2.5) \quad &= -\log p_{Y|X,Z}(\mathbf{y} | \mathbf{x}, \mathbf{z}) p_\theta(\mathbf{x} | \mathbf{z}) p_Z(\mathbf{z}) \\ &= F(\mathbf{x}, \mathbf{y}) + H_\theta(\mathbf{x}, \mathbf{z}) + \frac{1}{2} \|\mathbf{z}\|^2. \end{aligned}$$

Note that the first term is quadratic in $\mathbf{x}$ (assuming (1.1)), the third term is quadratic in $\mathbf{z}$ and all the difficulty lies in the coupling term $H_\theta(\mathbf{x}, \mathbf{z}) = -\log p_\theta(\mathbf{x} | \mathbf{z})$. For Gaussian decoders [36], the latter can be written as

$$\begin{aligned} (2.6) \quad H_\theta(\mathbf{x}, \mathbf{z}) &= \frac{1}{2} \left(d \log(2\pi) + \log \det \Sigma_\theta(\mathbf{z}) \right. \\ &\quad \left. + \|\Sigma_\theta^{-1/2}(\mathbf{z})(\mathbf{x} - \boldsymbol{\mu}_\theta(\mathbf{z}))\|^2 \right). \end{aligned}$$

which is also convex in $\mathbf{x}$. Hence, minimization with respect to $\mathbf{x}$ takes the convenient closed form:

$$\begin{aligned} (2.7) \quad \arg \min_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}) &= (\mathbf{A}^T \mathbf{A} + \sigma^2 \Sigma_\theta^{-1}(\mathbf{z}))^{-1} \\ &\quad \times (\mathbf{A}^T \mathbf{y} + \sigma^2 \Sigma_\theta^{-1}(\mathbf{z}) \boldsymbol{\mu}_\theta(\mathbf{z})). \end{aligned}$$

Unfortunately the coupling term H and hence J_1 is a priori non-convex in $\mathbf{z}$. As a consequence the $\mathbf{z}$ -minimization problem

$$(2.8) \quad \arg \min_{\mathbf{z}} J_1(\mathbf{x}, \mathbf{z})$$

is a priori more difficult. However, for Gaussian encoders, VAEs provide an approximate expression for this coupling term which is quadratic in $\mathbf{z}$. Indeed, given the equivalence

$$(2.9) \quad \begin{aligned} p_\theta(\mathbf{x} | \mathbf{z}) p_Z(\mathbf{z}) &= p_{X,Z}(\mathbf{x}, \mathbf{z}) \\ &= p_{Z|X}(\mathbf{z} | \mathbf{x}) p_X(\mathbf{x}) \\ &\approx q_\phi(\mathbf{z} | \mathbf{x}) p_X(\mathbf{x}) \end{aligned}$$

we have that

$$(2.10) \quad H_\theta(\mathbf{x}, \mathbf{z}) + \frac{1}{2} \|\mathbf{z}\|^2 \approx K_\phi(\mathbf{x}, \mathbf{z}) - \log p_X(\mathbf{x}).$$

where $K_\phi(\mathbf{x}, \mathbf{z}) = -\log q_\phi(\mathbf{z} | \mathbf{x})$. Therefore, this new coupling term becomes

$$\begin{aligned} K_\phi(\mathbf{x}, \mathbf{z}) &= -\log \mathcal{N}(\mathbf{z}; \boldsymbol{\mu}_\phi(\mathbf{x}), \boldsymbol{\Sigma}_\phi(\mathbf{x})) \\ &= \frac{1}{2} [l \log(2\pi) + \log \det \boldsymbol{\Sigma}_\phi(\mathbf{x}) \\ &\quad + \|\boldsymbol{\Sigma}_\phi^{-1/2}(\mathbf{x})(\mathbf{z} - \boldsymbol{\mu}_\phi(\mathbf{x}))\|^2], \end{aligned}$$

which is quadratic in $\mathbf{z}$. This provides an approximate expression for the energy (2.5) that we want to minimize, namely

$$(2.11) \quad J_2(\mathbf{x}, \mathbf{z}) := F(\mathbf{x}, \mathbf{y}) + K_\phi(\mathbf{x}, \mathbf{z}) - \log p_X(\mathbf{x}) \approx J_1(\mathbf{x}, \mathbf{z}).$$

This approximate functional is quadratic in $\mathbf{z}$, and minimization with respect to this variable yields

$$(2.12) \quad \arg \min_{\mathbf{z}} J_2(\mathbf{x}, \mathbf{z}) = \boldsymbol{\mu}_\phi(\mathbf{x}).$$

In the case of linear VAEs,

$$(2.13) \quad p_Z(\mathbf{z}) = \mathcal{N}(\mathbf{z}; 0, I)$$

$$(2.14) \quad p_{X|Z}(\mathbf{x} | \mathbf{z}) = \mathcal{N}(\mathbf{x}; V_\theta \mathbf{z} + v_\theta, \boldsymbol{\Sigma}_\theta).$$

It is easily shown that the posterior is also Gaussian [44], namely

$$(2.15) \quad p_{Z|X}(\mathbf{z} | \mathbf{x}) = \mathcal{N}(\mathbf{z}; M V_\theta^T (\mathbf{x} - v_\theta), \boldsymbol{\Sigma}_\theta M) \quad \text{where } M = (V_\theta^T V_\theta + \boldsymbol{\Sigma}_\theta)^{-1}.$$

Hence, the linear encoder $q_\phi(\mathbf{z} | \mathbf{x})$ that minimizes the ELBO is that of equation (2.15) so the approximation (2.9) is exact and then $J_1 = J_2$.

2.3. Alternate Joint Posterior Maximization. The previous observations suggest to adopt an alternate scheme to minimize $-\log p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$ in order to solve the inverse problem. We begin our presentation by a simple version of the proposed algorithm, which aims at managing the case where the approximation of J_1 by J_2 is exact (at least in the sense given in Assumption 1 below); then we propose an adaptation for the more realistic non-exact case and we explore its convergence properties.

When $J_1 = J_2$ is a bi-convex function, Algorithm 2.1 is known as *Alternate Convex Search*. Its behavior has been studied in [27, 1]. Here we shall consider the following (strong) assumption, which includes the strictly bi-convex case ($J_1 = J_2$):

Assumption 1. For any $\mathbf{x}$, if $\mathbf{z}^*$ is a global minimizer of $J_2(\mathbf{x}, \cdot)$, then $\mathbf{z}^*$ is a global minimizer of $J_1(\mathbf{x}, \cdot)$.

The proposed alternate minimization takes the simple and fast form depicted in Algorithm 2.1, which can be shown to converge to a stationary point of J_1 under Assumptions 1 and 2, as stated in Proposition 2.1 below. Note that the minimization in step 2 of Algorithm 2.1 does not require the knowledge of the unknown term $-\log p_X(\mathbf{x})$ in Equation (2.11) since it does not depend on $\mathbf{z}$.

Algorithm 2.1 Joint posterior maximization - exact case

Require: Measurements $\mathbf{y}$, Autoencoder parameters θ, ϕ , Initial condition $\mathbf{x}_0$

Ensure: $\hat{\mathbf{x}}, \hat{\mathbf{z}} = \arg \max_{\mathbf{x}, \mathbf{z}} p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$

```

1: for  $n := 0$  to maxiter do
2:    $\mathbf{z}_{n+1} := \arg \min_{\mathbf{z}} J_2(\mathbf{x}_n, \mathbf{z})$  // Quadratic approx
3:    $\mathbf{x}_{n+1} := \arg \min_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}_{n+1})$  // Quadratic
4: end for
5: return  $\mathbf{x}_{n+1}, \mathbf{z}_{n+1}$ 

```

The convergence analysis of the proposed schemes requires some general assumptions on the functions J_1 and J_2 :

Assumption 2. $J_1(\cdot, \mathbf{z})$ is convex and admits a unique minimizer for any $\mathbf{z}$. Moreover, J_1 is coercive and continuously differentiable.

The unicity of the minimizers of the partial function $J_1(\cdot, \mathbf{z})$ can be dropped. In this case, the proof of the convergence of Algorithm 2.1 has to be slightly adapted.

The convergence property of Algorithm 2.1 will be investigated in a wider framework below (Proposition 2.1). Note that all the properties required in Assumption 2 are satisfied if we use a differentiable activation function like the Exponential Linear Unit (ELU) [14] with $\alpha = 1$, instead of the more common ReLU activation function. More details can be found in Appendix A.

2.4. Approximate Alternate Joint Posterior Maximization. When the autoencoder approximation in (2.11) is not exact (Assumption 1), the energy we want to minimize in Algorithm 2.1, namely J_1 may not decrease. To ensure the decay, some additional steps can be added. Noting that the approximation provided by J_2 provides a fast and accurate heuristic to initialize the minimization of J_1 , an alternative scheme is proposed in Algorithm 2.2.

In Algorithm 2.2, GD is a gradient descent scheme such that for any starting point $\mathbf{z}_0$, the output $\mathbf{z}^+$ satisfies

$$\frac{\partial J_1}{\partial \mathbf{z}}(\mathbf{x}, \mathbf{z}^+) = 0 \quad \text{and} \quad J_1(\mathbf{x}, \mathbf{z}^+) \leq J_1(\mathbf{x}, \mathbf{z}_0)$$

Hence, one can consider for instance a gradient descent scheme which finds a local minimizer of $J_1(\mathbf{x}, \cdot)$ starting from $\mathbf{z}_0$.

Our experiments with Algorithm 2.2 (Section 3.5) show that during the first few iterations (where the approximation provided by J_2 is good enough) $\mathbf{z}^1$ and $\mathbf{z}^2$ reach convergence faster than $\mathbf{z}^3$. After a critical number of iterations the opposite is true (the initialization provided by the previous iteration is better than the J_2 approximation, and $\mathbf{z}^3$ converges faster).

Algorithm 2.2 Joint posterior maximization - approximate case**Require:** Measurements $\mathbf{y}$, Autoencoder parameters θ, ϕ , Initial conditions $\mathbf{x}_0, \mathbf{z}_0$ **Ensure:** $\hat{\mathbf{x}}, \hat{\mathbf{z}} = \arg \max_{\mathbf{x}, \mathbf{z}} p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$

```

1: for  $n := 0$  to maxiter do
2:    $\mathbf{z}^1 := \arg \min_{\mathbf{z}} J_2(\mathbf{x}_n, \mathbf{z})$  // Equation (2.12)
3:    $\mathbf{z}^2 := \text{GD}_{\mathbf{z}} J_1(\mathbf{x}_n, \mathbf{z})$ , starting from  $\mathbf{z} = \mathbf{z}^1$ 
4:    $\mathbf{z}^3 := \text{GD}_{\mathbf{z}} J_1(\mathbf{x}_n, \mathbf{z})$ , starting from  $\mathbf{z} = \mathbf{z}_n$ 
5:   for  $i := 1$  to 3 do
6:      $\mathbf{x}^i := \arg \min_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}^i)$  // Equation (2.7)
7:   end for
8:    $i^* := \arg \min_{i \in \{1,2,3\}} J_1(\mathbf{x}^i, \mathbf{z}^i)$ 
9:    $(\mathbf{x}_{n+1}, \mathbf{z}_{n+1}) := (\mathbf{x}^{i^*}, \mathbf{z}^{i^*})$ 
10: end for
11: return  $\mathbf{x}_{n+1}, \mathbf{z}_{n+1}$ 

```

These observations suggest that a faster execution, with the same convergence properties, can be achieved by the variant in Algorithm 2.3, which avoids the costly computation of $\mathbf{z}^2$ and $\mathbf{z}^3$ when unnecessary. Hence, in practice, we will use Algorithm 2.3 rather than Algorithm 2.2. However, Algorithm 2.2 provides a useful tool for diagnostics. Indeed, the comparison of the evaluation of $J_1(\mathbf{x}^i, \mathbf{z}^i)$ for $i = 1, 2, 3$ performed in step 8 permits to assess the evolution of the approximation of J_1 by J_2 .

Algorithm 2.3 is still quite fast when J_2 provides a sufficiently good approximation, since in that case the algorithm chooses $i^* = 1$, and avoids any call to the iterative gradient descent algorithm. Even if we cannot give a precise definition of what *sufficiently good* means, the sample comparison of K_ϕ and H_θ as functions of $\mathbf{z}$, displayed in Figure 3(a), shows that the approximation is fair enough in the sense that it preserves the global structure of J_1 . The same behavior was observed for a large number of random tests.

Note that Algorithm 2.1 is a particular instance of Algorithm 2.3 in the case where Assumption 1 holds, and $n_1 = n_2 = 0$ and if grad descent gives a global minimizer of the considered function (in this case, the computation of $\mathbf{z}^1, \mathbf{z}^2$, are skipped and only $\mathbf{z}^3$ is computed).

Proposition 2.1 (Convergence of Algorithm 2.3). *Let $\{(\mathbf{x}_n, \mathbf{z}_n)\}$ be a sequence generated by Algorithm 2.3. Under Assumption 2 we have that:*

1. *The sequence $\{J_1(\mathbf{x}_n, \mathbf{z}_n)\}$ converges monotonically when $n \rightarrow \infty$.*
2. *The sequence $\{(\mathbf{x}_n, \mathbf{z}_n)\}$ has at least one accumulation point.*
3. *All accumulation points of $\{(\mathbf{x}_n, \mathbf{z}_n)\}$ are stationary points of J_1 and they all have the same function value.*

Proof. Since we are interested in the behaviour for $n \rightarrow \infty$, we assume $n > n_2$ in Algorithm 2.3.

1. Since $n > n_2$ the algorithm chooses $i^* = 3$ and $\mathbf{z}_{n+1} = \mathbf{z}^3$. According to the definition of grad descent, one has

$$J_1(\mathbf{x}_n, \mathbf{z}_{n+1}) \leq J_1(\mathbf{x}_n, \mathbf{z}_n)$$

Algorithm 2.3 Joint posterior maximization - approximate case (faster version)

Require: Measurements $\mathbf{y}$, Autoencoder parameters θ, ϕ , Initial condition $\mathbf{x}_0$, iterations

```

 $n_1 \leq n_2 \leq n_{\max}$ 
Ensure:  $\hat{\mathbf{x}}, \hat{\mathbf{z}} = \arg \max_{\mathbf{x}, \mathbf{z}} p_{X,Z|Y}(\mathbf{x}, \mathbf{z} | \mathbf{y})$ 
1: for  $n := 0$  to  $n_{\max}$  do
2:   done := FALSE
3:   if  $n < n_1$  then
4:      $\mathbf{z}^1 := \arg \min_{\mathbf{z}} J_2(\mathbf{x}_n, \mathbf{z})$  // Equation (2.12)
5:      $\mathbf{x}^1 := \arg \min_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}^1)$  // Equation (2.7)
6:     if  $J_1(\mathbf{x}^1, \mathbf{z}^1) < J_1(\mathbf{x}_n, \mathbf{z}_n)$  then
7:        $i^* := 1$  //  $J_2$  is good enough
8:       done := TRUE
9:     end if
10:  end if
11:  if not done and  $n < n_2$  then
12:     $\mathbf{z}^1 := \arg \min_{\mathbf{z}} J_2(\mathbf{x}_n, \mathbf{z})$ 
13:     $\mathbf{z}^2 := \text{GD}_{\mathbf{z}} J_1(\mathbf{x}_n, \mathbf{z})$ , starting from  $\mathbf{z} = \mathbf{z}^1$ 
14:     $\mathbf{x}^2 := \arg \min_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}^2)$  // Equation (2.7)
15:    if  $J_1(\mathbf{x}^2, \mathbf{z}^2) < J_1(\mathbf{x}_n, \mathbf{z}_n)$  then
16:       $i^* := 2$  //  $J_2$  init is good enough
17:      done := TRUE
18:    end if
19:  end if
20:  if not done then
21:     $\mathbf{z}^3 := \text{GD}_{\mathbf{z}} J_1(\mathbf{x}_n, \mathbf{z})$ , starting from  $\mathbf{z} = \mathbf{z}_n$ 
22:     $\mathbf{x}^3 := \arg \min_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}^3)$  // Equation (2.7)
23:     $i^* := 3$ 
24:  end if
25:   $(\mathbf{x}_{n+1}, \mathbf{z}_{n+1}) := (\mathbf{x}^{i^*}, \mathbf{z}^{i^*})$ 
26: end for
27: return  $\mathbf{x}_{n+1}, \mathbf{z}_{n+1}$ 
    
```

and by optimality one has

$$J_1(\mathbf{x}_{n+1}, \mathbf{z}_{n+1}) \leq J_1(\mathbf{x}_n, \mathbf{z}_{n+1}).$$

Hence, since J_1 is coercive (thus, lowerbounded), Statement 1 is straightforward.

2. Thanks to the coercivity of J_1 , the sequences $\{(\mathbf{x}_n, \mathbf{z}_n)\}$ and $\{(\mathbf{x}_n, \mathbf{z}_{n+1})\}$ are bounded, thus admit an accumulation point.

3. Using Fermat's rule and the definition of grad descent, one has

$$\frac{\partial J_1}{\partial \mathbf{z}}(\mathbf{x}_n, \mathbf{z}_{n+1}) = 0 \quad \text{and} \quad \frac{\partial J_1}{\partial \mathbf{x}}(\mathbf{x}_{n+1}, \mathbf{z}_{n+1}) = 0.$$

Let $(\mathbf{x}^*, \mathbf{z}^*)$ be an accumulation point of $\{(\mathbf{x}_n, \mathbf{z}_n)\}$. By double extraction, one can find two

subsequences such that

$$(\mathbf{x}_{n_j+1}, \mathbf{z}_{n_j+1}) \rightarrow (\mathbf{x}^*, \mathbf{z}^*) \text{ and } (\mathbf{x}_{n_j}, \mathbf{z}_{n_j+1}) \rightarrow (\hat{\mathbf{x}}^*, \mathbf{z}^*)$$

By continuity of ∇J_1 , one gets that

$$\frac{\partial J_1}{\partial \mathbf{z}}(\hat{\mathbf{x}}^*, \mathbf{z}^*) = 0 \quad \text{and} \quad \frac{\partial J_1}{\partial \mathbf{x}}(\mathbf{x}^*, \mathbf{z}^*) = 0$$

In particular, the convexity of $J_1(\cdot, \mathbf{z}^*)$ and Assumption 2 ensure that $\mathbf{x}^*$ is a global minimizer of $J_1(\cdot, \mathbf{z}^*)$. Besides, the inequalities proved in Point 1 above show that

$$J_1(\mathbf{x}^*, \mathbf{z}^*) = J_1(\hat{\mathbf{x}}^*, \mathbf{z}^*) = \lim_{n \rightarrow \infty} J_1(\mathbf{x}_n, \mathbf{z}_n)$$

that is, $\hat{\mathbf{x}}^*$ is also a global minimizer of $J_1(\cdot, \mathbf{z}^*)$. Since $J_1(\cdot, \mathbf{z}^*)$ has a unique minimizer, one has $\hat{\mathbf{x}}^* = \mathbf{x}^*$, and

$$\frac{\partial J_1}{\partial \mathbf{z}}(\mathbf{x}^*, \mathbf{z}^*) = 0$$

namely $(\mathbf{x}^*, \mathbf{z}^*)$ is a stationary point of J_1 . Note that we have also proved that $\mathbf{x}_{n_j}$ and $\mathbf{x}_{n_j+1}$ have same limit. ■

Remark 2.2. Note that if $n_1 = n_2 = \infty$ we cannot assume that $i^* = 3$. In that case statements 1 and 2 are still valid but the third statement is not. The reason is that for $i^* \in \{1, 2\}$ we cannot guarantee the chain of inequalities

$$J_1(\mathbf{x}_{n+1}, \mathbf{z}_{n+1}) \leq J_1(\mathbf{x}_n, \mathbf{z}_{n+1}) \leq J_1(\mathbf{x}_n, \mathbf{z}_n)$$

but only

$$J_1(\mathbf{x}_{n+1}, \mathbf{z}_{n+1}) \leq J_1(\mathbf{x}_n, \mathbf{z}_n).$$

This is consistent with the design of the algorithm where iterations $n < n_2$ serve as an heuristic to guide the algorithm to a sensible critical point. However, convergence to a critical point is only guaranteed by the final iterations $n > n_2$.

2.5. MAP- $\mathbf{z}$ as the limit case for $\beta \rightarrow \infty$. If one wishes to compute the MAP- $\mathbf{z}$ estimator instead of the joint MAP- $\mathbf{x}$ - $\mathbf{z}$ from the previous section, one has two options:

1. Use your favorite gradient descent algorithm to solve equation (1.4).
2. Use Algorithm 2.3 to solve a series of joint MAP- $\mathbf{x}$ - $\mathbf{z}$ problems with increasing values of $\beta \rightarrow \infty$ as suggested in Algorithm 1.1.

In the experimental section we show that the second approach most often leads to a better optimum.

In practice, in order to provide a stopping criterion for Algorithm 1.1 and to make a sensible choice of β -values we reformulate Algorithm 1.1 as a constrained optimization problem

$$\arg \min_{\mathbf{x}, \mathbf{z} : \|\mathbf{G}(\mathbf{z}) - \mathbf{x}\|^2 \leq \varepsilon} F(\mathbf{x}, \mathbf{y}) + \frac{1}{2} \|\mathbf{z}\|^2.$$

The corresponding Lagrangian form is

$$(2.16) \quad \max_{\beta} \min_{\mathbf{x}, \mathbf{z}} F(\mathbf{x}, \mathbf{y}) + \frac{1}{2} \|\mathbf{z}\|^2 + \beta (\|\mathbf{G}(\mathbf{z}) - \mathbf{x}\|^2 - \varepsilon)^+$$

and we use the exponential multiplier method [68] to guide the search for the optimal value of β (see Algorithm 2.4)

Algorithm 2.4 MAP- $\mathbf{z}$ as the limit of joint MAP- $\mathbf{x}$ - $\mathbf{z}$.

Require: Measurements $\mathbf{y}$, Tolerance ε , Rate $\rho > 0$, Initial β_0 , Initial $\mathbf{x}_0$, Iterations $0 \leq$

$$n_1 \leq n_2 \leq n_{\max}$$

Ensure: $\arg \min_{\mathbf{x}, \mathbf{z} : \|\mathbf{G}(\mathbf{z}) - \mathbf{x}\|^2 \leq \varepsilon} F(\mathbf{x}, \mathbf{y}) + \frac{1}{2} \|\mathbf{z}\|^2$.

```

1:  $\beta := \beta_0$ 
2:  $\mathbf{x}^0, \mathbf{z}^0 :=$  Algorithm 2.3 starting from  $\mathbf{x} = \mathbf{x}_0$  with  $\beta, n_1, n_2, n_{\max}$ .
3: converged := FALSE
4:  $k := 0$ 
5: while not converged do
6:    $\mathbf{x}^{k+1}, \mathbf{z}^{k+1} :=$  Algorithm 2.3 starting from  $\mathbf{x} = \mathbf{x}^k$  with  $\beta$  and  $n_1 = n_2 = 0$ 
7:    $C = \|\mathbf{G}(\mathbf{z}^{k+1}) - \mathbf{x}^{k+1}\|^2 - \varepsilon$ 
8:    $\beta := \beta \exp(\rho C)$ 
9:   converged :=  $(C \leq 0)$ 
10:   $k := k + 1$ 
11: end while
12: return  $\mathbf{x}^k, \mathbf{z}^k$ 
```

3. Experimental results.

3.1. Baseline algorithms. To validate our approach, we perform comparisons on several inverse problems with the following algorithms:

- CSGM (Bora et al. [6]) directly computes the $\mathbf{z}$ – MAP estimator as defined in Equation (1.4) using gradient descent. We run CSGM using the decoder of a VAE as generator $\mathbf{G}$ starting at random $\mathbf{z}_0$. In addition, as Bora et al. note that random restarts are important for good performance, we also compute the best result (as measured by (1.4)) among $m = 10$ different random initializations $\mathbf{z}_0$ and refer to this variant as mCSGM.
- PULSE [46] is very similar to CSGM but restricts the search of the latent code $\mathbf{z}$ to the sphere of radius $\sqrt{l}$, arguing that it concentrates most of the probability mass of a Gaussian distribution $\mathcal{N}(0, I)$ on a high-dimensional space $\mathbb{R}^l$.
- PGD-GAN [63] performs a projected gradient descent of $F(\mathbf{x}, \mathbf{y})$ wrt $\mathbf{x}$:

$$(3.1) \quad \begin{cases} \mathbf{w}_k = \mathbf{x}_k - \eta \mathbf{A}^T (\mathbf{A} \mathbf{x}_k - \mathbf{y}) \\ \mathbf{x}_{k+1} = \mathbf{G}(\arg \min_{\mathbf{z}} \|\mathbf{w}_k - \mathbf{G}(\mathbf{z})\|). \end{cases}$$

- In addition, we implement the splitting method of Algorithm 1.1 which is a simple continuation scheme for the $\mathbf{z}$ – MAP estimator of Equation (1.4).

For a fair comparison we run all algorithms on the same prior, *i.e.* the same generator network $\mathbf{G} = \boldsymbol{\mu}_\theta$ where $\boldsymbol{\mu}_\theta$ is the decoder mean from the VAE model that we trained for JPMAP.

3.2. Inverse problems. Here, we briefly describe the inverse problems $\mathbf{y} = \mathbf{A}\mathbf{x} + \eta$, $\eta \sim \mathcal{N}(0, \sigma^2 I)$ to be considered for validating our approach:

- *Denoising:* $\mathbf{A} = I$ and σ large.
- *Compressed Sensing:* the sensing matrix $\mathbf{A} \in \mathbb{R}^{q \times d}$ has Gaussian random entries $\mathbf{A}_{ij} \sim \mathcal{N}(0, 1/q)$, where $q \ll d$ is the number of measurements.

- *Interpolation*: $\mathbf{A}$ is a diagonal matrix with random binary entries, so masking a percentage p of the image pixels.
- *(Non-blind) Deblurring*: $\mathbf{A}\mathbf{x} = h * \mathbf{x}$ where h is a known convolution kernel.
- *Super-resolution*: $\mathbf{A}$ is a downsampling/decimation operator of scaling factor s .

3.3. AutoEncoder and dataset. In order to test our joint prior maximization model we first train a Variational Autoencoder like in [36] on the training data of MNIST handwritten digits [40].

The *stochastic encoder* takes as input an image $\mathbf{x}$ of $28 \times 28 = 784$ pixels and produces as an output the mean and (diagonal) covariance matrix of the Gaussian distribution $q_\phi(\mathbf{z}|\mathbf{x})$, where the latent variable $\mathbf{z}$ has dimension 8. The architecture of the encoder is composed of 3 fully connected layers with ELU activations (to preserve continuous differentiability). The sizes of the layers are as follows: $784 \rightarrow 500 \rightarrow 500 \rightarrow (8 + 8)$. Note that the output is of size $8 + 8$ in order to encode the mean and diagonal covariance matrix, both of size 8.

The *stochastic decoder* takes as an input the latent variable $\mathbf{z}$ and outputs the mean and covariance matrix of the Gaussian distribution $p_\theta(\mathbf{x}|\mathbf{z})$. Following [16] we chose here an isotropic covariance $\Sigma_\theta(\mathbf{z}) = \gamma^2 I$ where $\gamma > 0$ is trained, but independent of $\mathbf{z}$. This choice simplifies the minimization problem (2.8), because the term $\det \Sigma_\theta(\mathbf{z})$ (being constant) has no effect on the $\mathbf{z}$ -minimization. The architecture of the decoder is also composed of 3 fully connected layers with ELU activations (to preserve continuous differentiability). The sizes of the layers are as follows: $8 \rightarrow 500 \rightarrow 500 \rightarrow 784$. Note that the covariance matrix is constant, so it does not augment the size of the output layer which is still $784 = 28 \times 28$ pixels.

We also trained a VAE on CelebA [42] images cropped to $64 \times 64 \times 3$, with latent dimension ranging from 64 to 512. We choose a DCGAN-like [54] CNN architecture as encoder and a symmetrical one as decoder with ELU activations, batch normalization and isotropic covariance as before. For more details, see the code³.

We train these architectures using PyTorch [49] with batch size 128 and Adam algorithm for 200 epochs with learning rate 0.0001 and rest of the parameters as default.

3.4. Need to train the VAE with a denoising criterion. It should be noted that when training our Variational Autoencoder we should be more careful than usual. Indeed in the most widespread applications of VAEs they are only used as a generative model or as a way to interpolate between images that are close to $\mathcal{M}$, *i.e.* the image of the generator μ_θ . For such applications it is sufficient to train the encoder μ_ϕ, Σ_ϕ on a training set that is restricted to $\mathcal{M}$.

In our case however, we need the encoder to provide sensible values even when its input $\mathbf{x}$ is quite far away from $\mathcal{M}$: the encoder has to actually fulfill two functions at the same time:

1. (Approximately) project $\mathbf{x}$ to its closest point in $\mathcal{M}$, and
2. compute the encoding of this projected value (which should be the same as the encoding of the original $\mathbf{x}$).

³Code available at <https://github.com/mago876/JPMAP>.

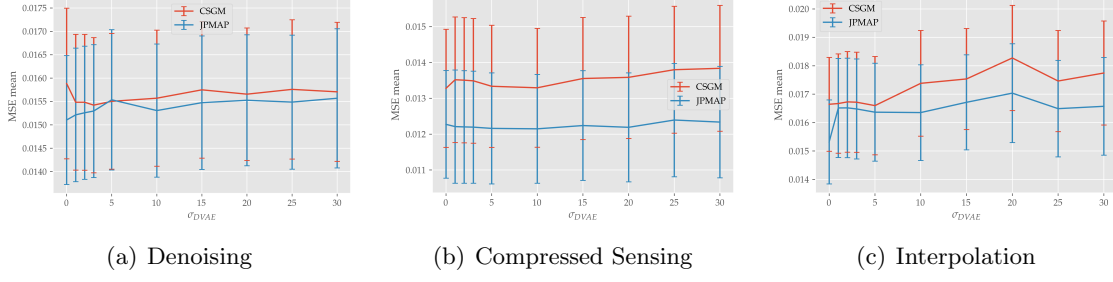


Figure 1. Evaluating the quality of the generative model as a function of σ_{DVAE} . On (a) Denoising (Gaussian noise $\sigma = 150$), (b) Compressed Sensing ($\sim 10.2\%$ measurements, noise $\sigma = 10$) and (c) Interpolation (80% of missing pixels, noise $\sigma = 10$). Results of both algorithms are computed on a batch of 50 images and initializing on ground truth $\mathbf{x}^*$ (for CSGM we use $\mathbf{z}_0 = \mu_\phi(\mathbf{x}^*)$).

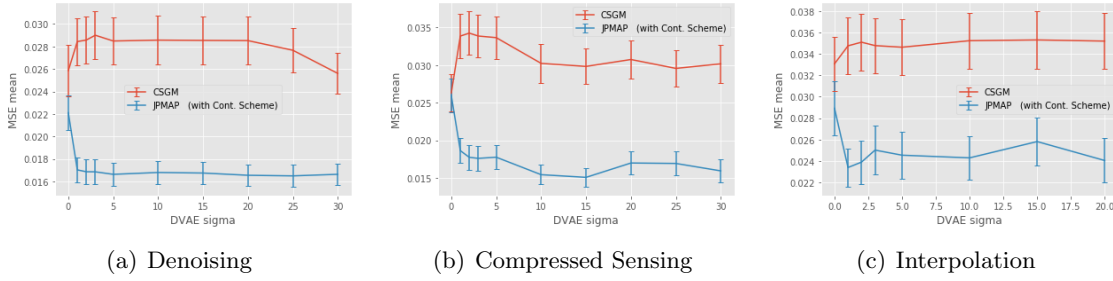


Figure 2. Evaluating the effectiveness of JPMAP vs CSGM as a function of σ_{DVAE} (same setup of Figure 1). Without a denoising criterion $\sigma_{\text{DVAE}} = 0$ the JPMAP algorithm may provide wrong guesses $\mathbf{z}^1$ when applying the encoder in step 2 of Algorithm 2.2. For $\sigma_{\text{DVAE}} > 0$ however, the alternating minimization algorithm can benefit from the robust initialization heuristics provided by the encoder, and it consistently converges to a better local optimum than the simple gradient descent in CSGM.

Traditional VAE training procedures do not ensure that the encoder generalizes well to $\mathbf{x} \notin \mathcal{M}$. In order to ensure this generalization ability we adopt the training procedure of the DVAE (Denoising VAE) proposed by Im et al. [33], which consists in adding various realizations of zero-mean Gaussian noise of variance σ_{DVAE}^2 to the samples $\mathbf{x}$ presented to the encoder, while still requiring the decoder to match the noiseless value, *i.e.* we optimize the parameters in such a way that

$$(3.2) \quad \mu_\theta(\mu_\phi(\tilde{\mathbf{x}})) \approx \mathbf{x}$$

where $\tilde{\mathbf{x}} = \mathbf{x} + \sigma_{\text{DVAE}}\varepsilon$ and $\varepsilon \sim \mathcal{N}(0, I)$ for all $\mathbf{x}$ in the training set and for many realizations of ε .

More specifically, if we take a corruption model $p(\tilde{\mathbf{x}}|\mathbf{x})$ like above, it can be shown [33] that

$$(3.3) \quad \tilde{\mathcal{L}}_{\theta,\phi}(\mathbf{x}) = \mathbb{E}_{p(\tilde{\mathbf{x}}|\mathbf{x})} \left[\mathbb{E}_{q_\phi(\mathbf{z}|\tilde{\mathbf{x}})} [\log p_\theta(\mathbf{x}|\mathbf{z})] - KL(q_\phi(\mathbf{z}|\tilde{\mathbf{x}}) \parallel p_Z(\mathbf{z})) \right]$$

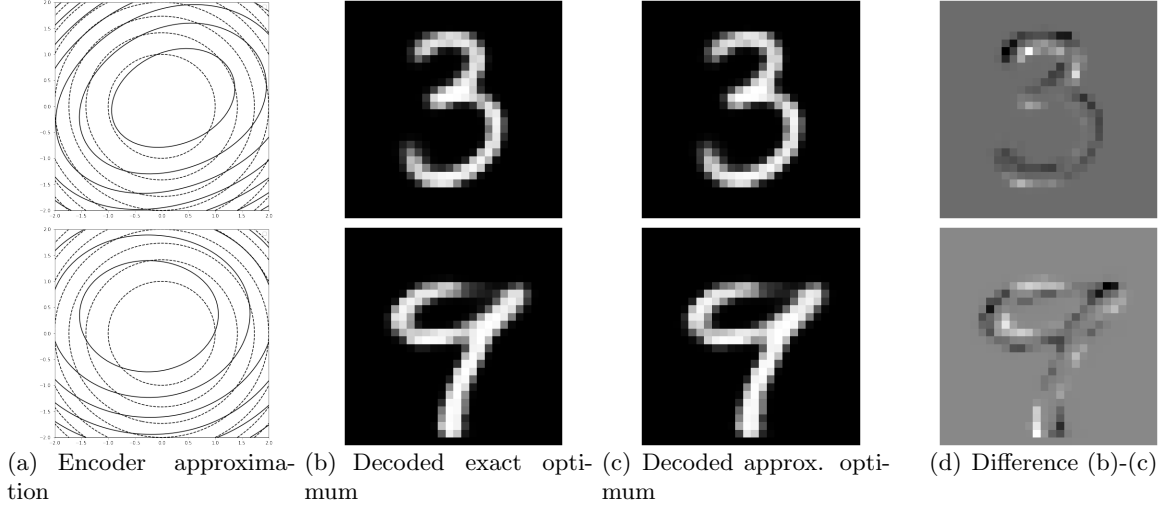


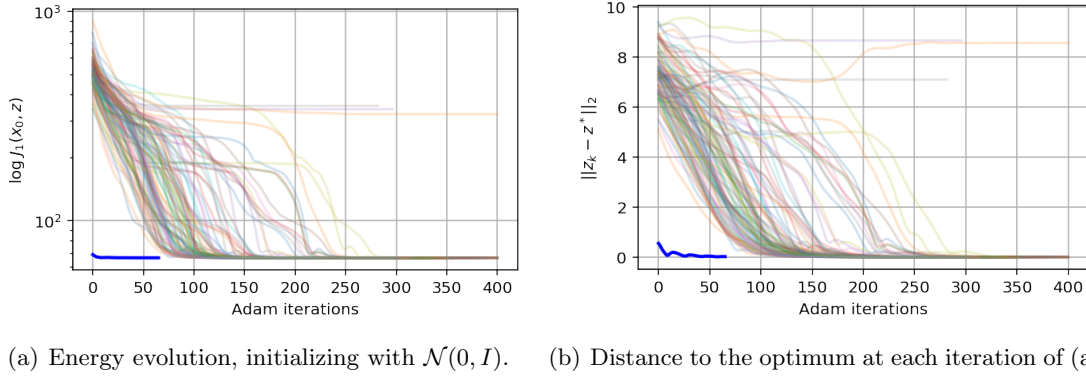
Figure 3. Encoder approximation: (a) Contour plots of $-\log p_\theta(\mathbf{x}|\mathbf{z}) + \frac{1}{2}\|\mathbf{z}\|^2$ and $-\log q_\phi(\mathbf{z}|\mathbf{x})$ for a fixed $\mathbf{x}$ and for a random 2D subspace in the $\mathbf{z}$ domain (the plot shows $\pm 2\Sigma_\phi^{1/2}$ around μ_ϕ). Observe the relatively small gap between the true posterior $p_\theta(\mathbf{z}|\mathbf{x})$ and its variational approximation $q_\phi(\mathbf{z}|\mathbf{x})$. This figure shows some evidence of partial $\mathbf{z}$ -convexity of J_1 around the minimum of J_2 , but it does not show how far is $\mathbf{z}^1$ from $\mathbf{z}^2$. (b) Decoded exact optimum $\mathbf{x}_1 = \mu_\theta(\arg \max_{\mathbf{z}} p_\theta(\mathbf{x}|\mathbf{z})e^{\frac{1}{2}\|\mathbf{z}\|^2})$. (c) Decoded approximate optimum $\mathbf{x}_2 = \mu_\theta(\arg \max_{\mathbf{z}} q_\phi(\mathbf{z}|\mathbf{x}))$. (d) Difference between (b) and (c).

is an alternative ELBO of (2.3). In practice, using Monte Carlo for estimating the expectation $\mathbb{E}_{p(\bar{\mathbf{x}}|\mathbf{x})}$ in (3.3), we only need to add noise to $\mathbf{x}$ before passing it to the encoder q_ϕ during training, as mentioned in (3.2).

Our experiments with this denoising criterion confirm the observation by Im et al. [33] that it does not degrade the quality of the generative model, as long as σ_{DVAE} is not too large (see Figure 1). As a side benefit, however, we obtain a more robust encoder that generalizes well for values of $\mathbf{x}$ that are not in $\mathcal{M}$ but within a neighbourhood of size $\approx \sigma_{\text{DVAE}}$ around $\mathcal{M}$. This side benefit, which was not the original intention of the DVAE training algorithm in [33] is nevertheless crucial for the success of our algorithm as demonstrated in Figure 2. The same figure shows that as long as $\sigma_{\text{DVAE}} \geq 5$ its value does not significantly affect the performance. In the sequel we use $\sigma_{\text{DVAE}} = 15$.

3.5. Effectiveness of the encoder as a fast approximate minimizer. Proposition 2.1 shows that the proposed alternate minimization scheme in Algorithm 2.3 converges to a stationary point of J_1 . And so does the gradient descent scheme in [6]. Since both algorithms have to deal with non-convex energies, they both risk converging to spurious local minima. Also both algorithms solve essentially the same model when the variance γ of the coupling term tends to zero.

If our algorithm shows better performance (see next subsection), it is mainly because it relies on a previously trained VAE in two fundamental ways: (i) to avoid getting trapped in spurious local minima and (ii) to accelerate performance during the initial iterations ($n < n_{\min}$).



(a) Energy evolution, initializing with $\mathcal{N}(0, I)$. (b) Distance to the optimum at each iteration of (a).

Figure 4. Effectiveness of the encoder approximation: We take $\mathbf{x}_0$ from the test set of MNIST and minimize $J_1(\mathbf{x}_0, \mathbf{z})$ with respect to $\mathbf{z}$ using gradient descent from random Gaussian initializations $\mathbf{z}_0$. The blue thick curve represents the trajectory if we initialize at the encoder approximation $\mathbf{z}^1 = \arg \min_{\mathbf{z}} J_2(\mathbf{x}_0, \mathbf{z}) = \boldsymbol{\mu}_\phi(\mathbf{x}_0)$. (a): Plots of the energy iterates $J_1(\mathbf{x}_0, \mathbf{z}_k)$. (b): ℓ^2 distances of each trajectory with respect to the global optimum $\mathbf{z}^*$. Conclusion: Observe that the encoder initialization allows much faster convergence both in energy and in $\mathbf{z}$, and it avoids the few random initializations that lead to a wrong stationary point different from the unique global minimizer.

These two features are only possible if the autoencoder approximation is good enough and if the encoder is able to provide good initializations for the non-convex $\mathbf{z}$ - optimization sub-problem in line 13 of Algorithm 2.3.

Figures 3 and 4 illustrate these two properties of our VAE. We do so by selecting a random $\mathbf{x}_0$ from MNIST test set and by computing $\mathbf{z}^*(\mathbf{z}_0) := \text{GD}_{\mathbf{z}} J_1(\mathbf{x}_0, \mathbf{z})$ with different initial values $\mathbf{z}_0$. These experiments were performed using the ADAM minimization algorithm with learning rate equal to 0.01. Figure 4(a) shows that $\mathbf{z}^*(\mathbf{z}_0)$ reaches the global optimum for most (but not all) initializations $\mathbf{z}_0$. Indeed from 200 random initializations $\mathbf{z}_0 \sim \mathcal{N}(0, I)$, 195 reach the same global minimum, whereas 5 get stuck at a higher energy value. However these 5 initial values have energy values $J_1(\mathbf{x}_0, \mathbf{z}_0) \gg J_1(\mathbf{x}_0, \mathbf{z}^1)$ far larger than those of the encoder initialization $\mathbf{z}^1 = \boldsymbol{\mu}_\phi(\mathbf{x}_0)$, and are thus never chosen by Algorithm 2.3. The encoder initialization $\mathbf{z}^1$ on the other hand provides much faster convergence to the global optimum.

In addition, this experiment shows that we cannot assume $\mathbf{z}$ -convexity: The presence of plateaux in the trajectories of many random initializations as well as the fact that a few initializations do not lead to the global minimum indicates that J_1 may not be everywhere convex with respect to $\mathbf{z}$. However, in contrast to classical works on alternate convex search, our approach adopts weaker assumptions and does not require convexity on $\mathbf{z}$ to prove convergence in Proposition 2.1.

In Figure 4(b) we display the distances of each trajectory to the global optimum $\mathbf{z}^*$ (taken as the median over all initializations $\mathbf{z}_0$ of the final iterates $\mathbf{z}^*(\mathbf{z}_0)$); note that this optimum is always reached, which suggests that $\mathbf{z} \mapsto J_1(\mathbf{x}_0, \mathbf{z})$ has a unique global minimizer in this case. Finally, Figure 3 shows that the encoder approximation is quite good both in the latent

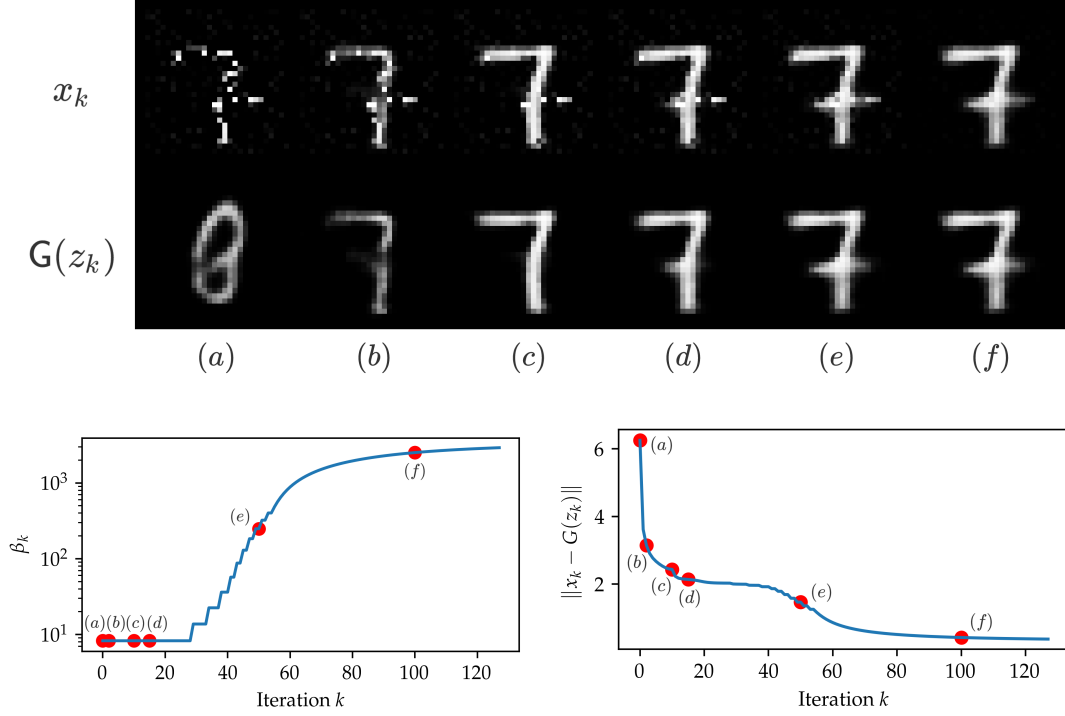


Figure 5. Evolution of Algorithm 2.4. In this interpolation example, JPMAP starts with the initialization in (a). During first iterations (b) – (d) where β_k is small, $\mathbf{x}_k$ and $\mathbf{G}(\mathbf{z}_k)$ start loosely approaching each other at a coarse scale, and $\mathbf{x}_k$ only fills missing pixels with the ones of $\mathbf{G}(\mathbf{z}_k)$ (in particular the noise of $\mathbf{y}$ is still present). By increasing β_k in (e) – (f) we enforce $\|\mathbf{G}(\mathbf{z}_k) - \mathbf{x}_k\|^2 \leq \epsilon$. Here we set $\epsilon = (\frac{3}{255})^2 d$, that is, MSE of 3 gray levels.

space (Figure 3(a)) and in image space (Figures 3(b) and 3(c)). It also shows that the true posterior $p_\theta(\mathbf{z}|\mathbf{x})$ is pretty close to log-concave near the maximum of $q_\phi(\mathbf{z}|\mathbf{x})$.

3.6. Image restoration experiments. *Choice of $\mathbf{x}_0$:* In the previous section, our validation experiments used a random $\mathbf{x}_0$ from the data set as initialization. When dealing with an image restoration problem, Algorithms 2.2 and 2.3 require an initial value of $\mathbf{x}_0$ to be chosen. In all experiments we choose this initial value as $\mathbf{A}^T \mathbf{y}$.

Choice of n_1 and n_2 : After a few runs of Algorithm 2.2 we find that in most cases, during the first 10 or 20 iterations $\mathbf{z}^1$ decreases the energy with respect to the previous iteration, and this value depends on the inverse problems (for example, for denoising is smaller than for compressed sensing). But after at most 150 iterations the autoencoder approximation is no longer good enough and we need to perform gradient descent on $\mathbf{z}_n$ in order to further decrease the energy. Based on these findings we set $n_1 = 25$ and $n_2 = 150$ in Algorithm 2.3 for all experiments. Note that we could also choose $n_1 = n_2 = n_{\max}$, since in all our experiments we observed that the algorithm auto-regulates itself, evolving from $i^* = 1$ in the first few dozen iterations to $i^* = 3$ when it is close to convergence. Choosing a finite value for n_1 and

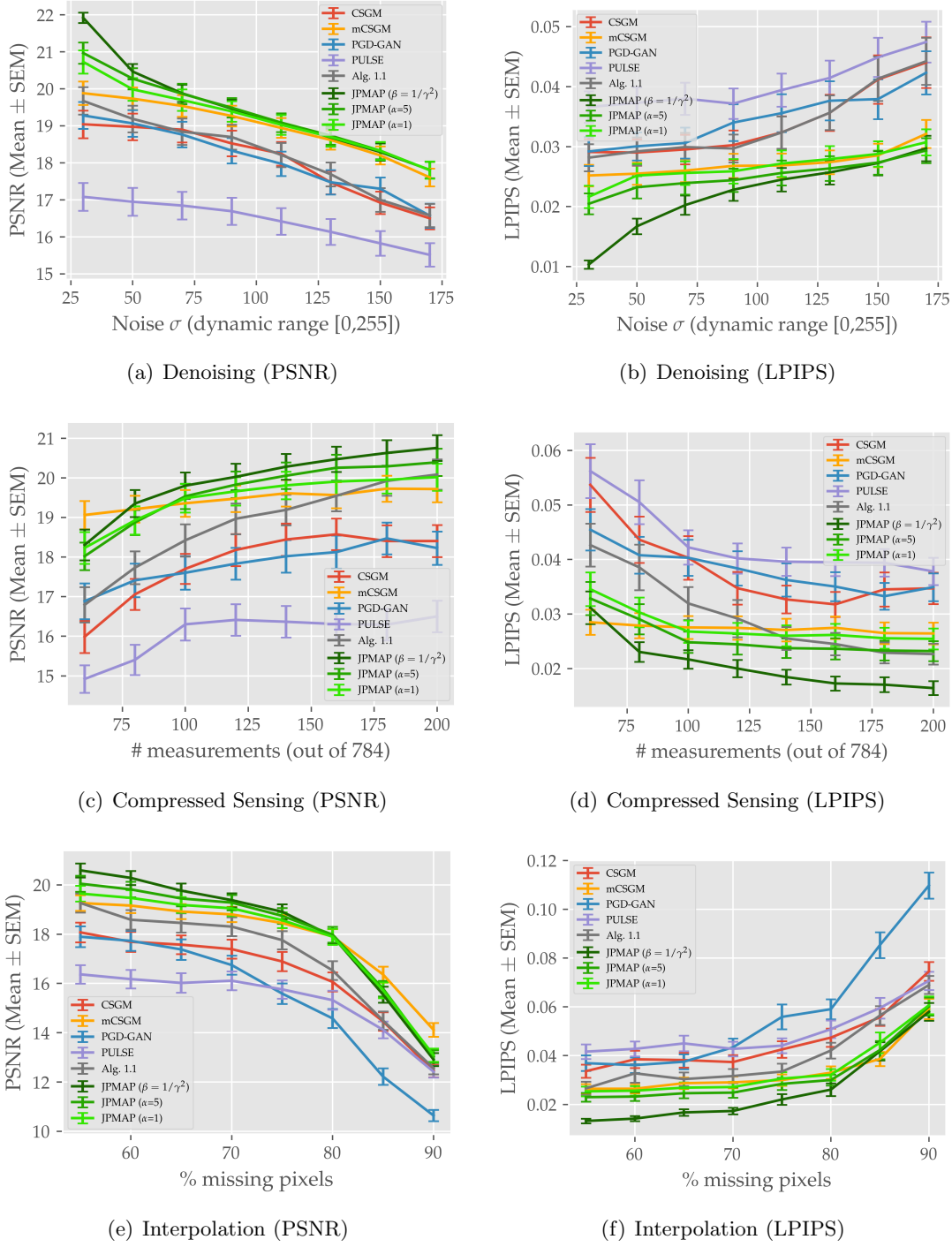


Figure 6. Denoising, Compressed Sensing and Interpolation: *Evaluating the effectiveness of Algorithm 2.3 (fixed β) and Algorithm 2.4 for different values of $\epsilon = (\frac{\alpha}{255})^2 n$, with $\sigma_{\text{DVAE}} = 15$ (metrics were computed on a batch of 100 test images). For PSNR, higher is better and for LPIPS, lower is better. For comparison we provide the results of the baselines introduced in Section 3.1 (namely, Algorithm 1.1, CSGM [6], mCSGM (CSGM with restarts), PGD-GAN [63] and PULSE [46].)*

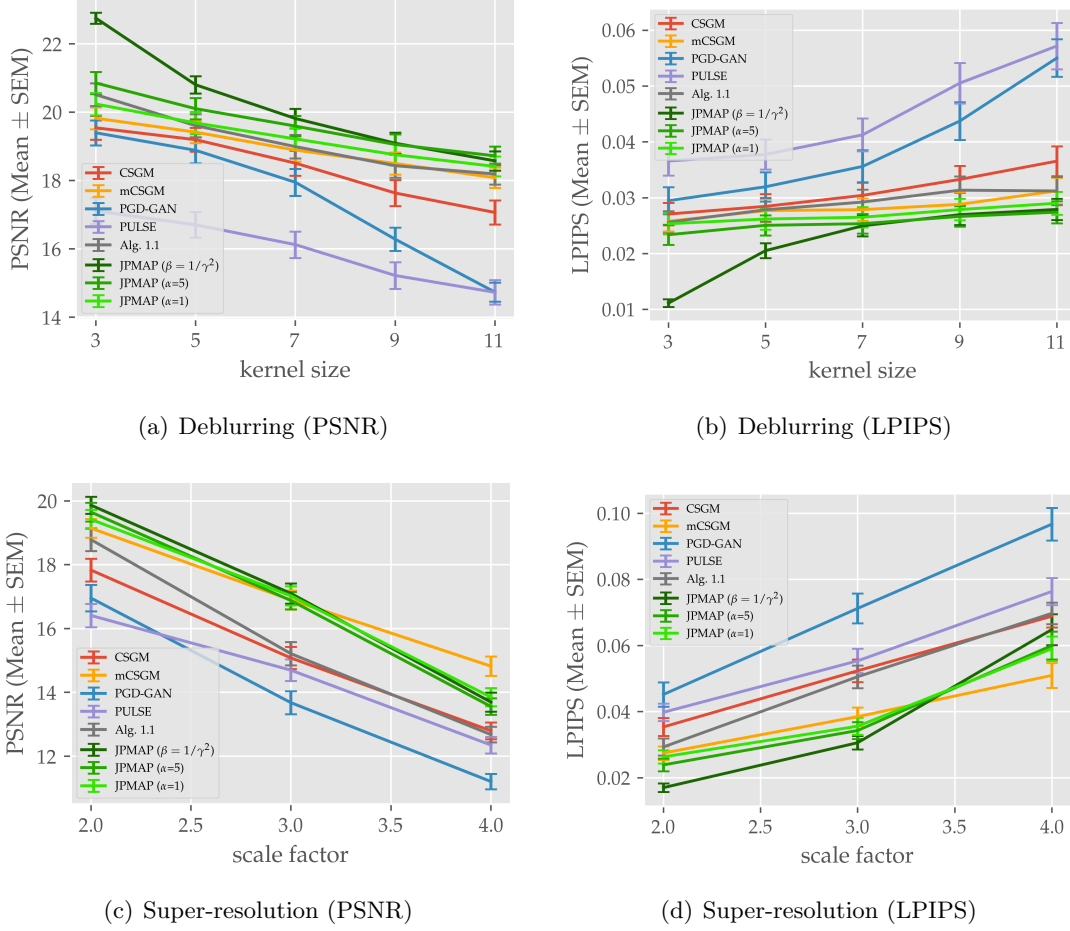


Figure 7. Deblurring and Super-resolution: *Evaluating the effectiveness of Algorithm 2.3 (fixed β) and Algorithm 2.4 for different values of $\epsilon = (\frac{\alpha}{255})^2 n$, with $\sigma_{\text{DVAE}} = 15$ (metrics were computed on a batch of 100 test images). For PSNR, higher is better and for LPIPS, lower is better. For comparison we provide the results of the baselines introduced in Section 3.1 (namely, Algorithm 1.1, CSGM [6], mCSGM (CSGM with restarts), PGD-GAN [63] and PULSE [46].)*

n_2 is only needed to ensure that $i^* = 3$ when $n \rightarrow \infty$, which is a necessary condition to prove statement 3 of Proposition 2.1.

Figure 5 shows the evolution of $\mathbf{x}_k$ and $\mathbf{G}(\mathbf{z}_k)$ from Algorithm 2.4 in an interpolation example. Here we can see how the exponential multiplier method in Equation (2.16) updates the values of β_k to ensure $\|\mathbf{G}(\mathbf{z}_k) - \mathbf{x}_k\|^2 \leq \epsilon$.

Figures 6 and 7 show the results of denoising, interpolation, compressed sensing, deblurring and super-resolution experiments on MNIST for different degradation levels using the proposed algorithm (JPMAP) and the baseline algorithms introduced in section 3.1. The

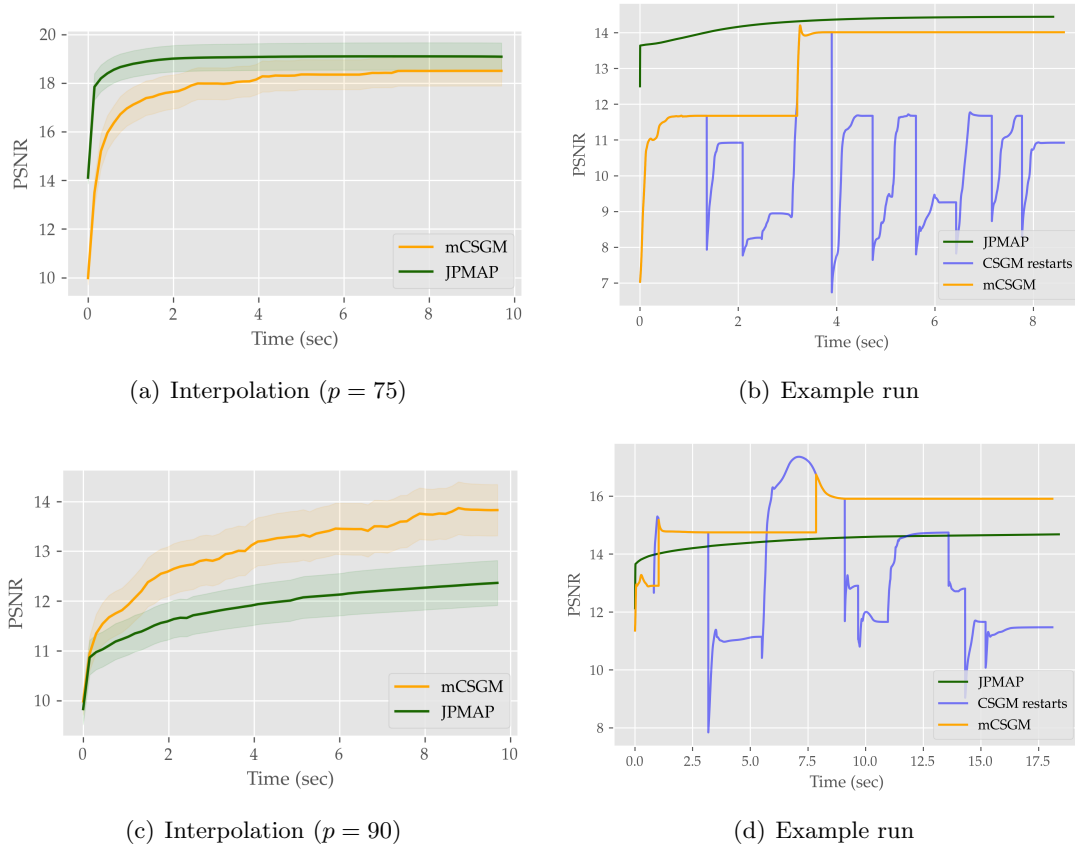
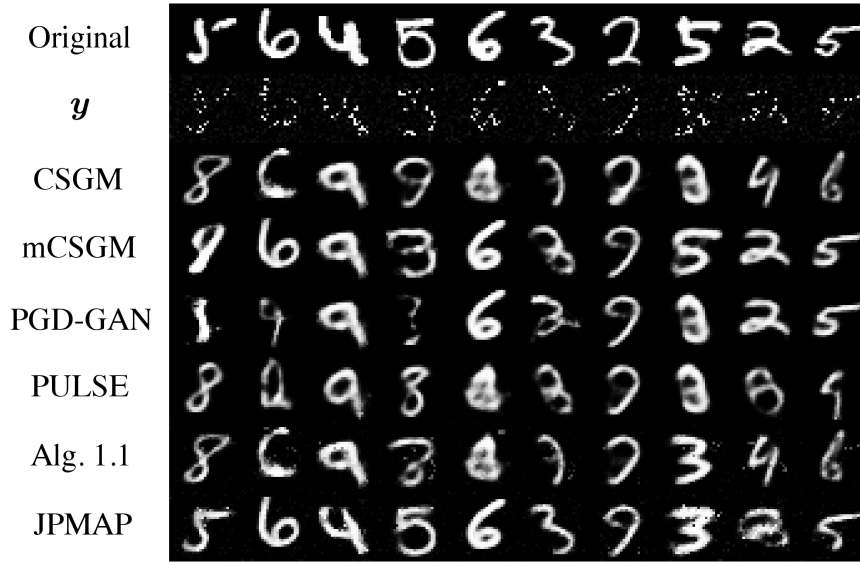


Figure 8. Time/PSNR comparison between mCSGM and JPMAP. Left: *Confidence intervals (for a batch of 100 random experiments) for PSNR vs computing time for both algorithms on the interpolation problem with $p\%$ of missing pixels with noise std $\sigma = 10/255$.* Right: *Detailed view of one of the 100 random experiments on the left. Blue lines represent $m = 10$ random restarts of CSGM and the orange line is the PSNR of the best $\mathbf{z}_k$ at iteration k of mCSGM as measured by (1.4).*

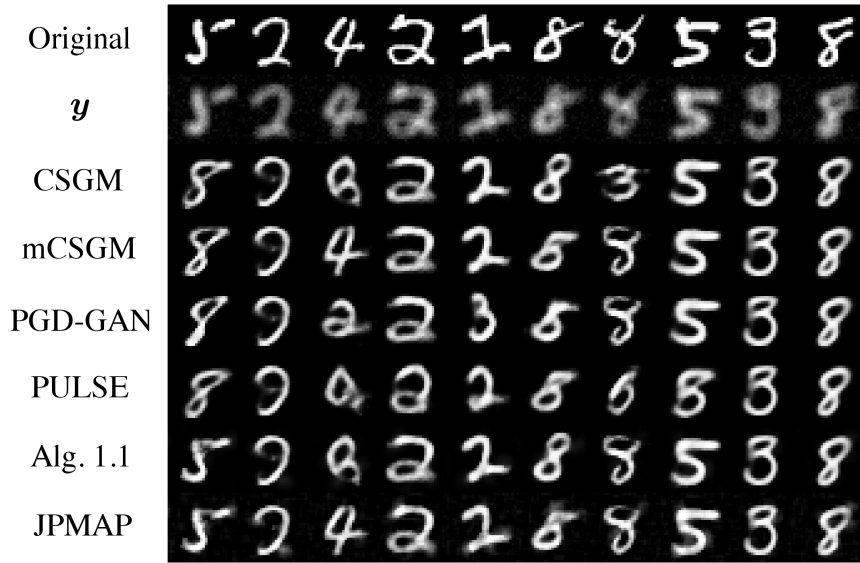
metrics used are PSNR and LPIPS⁴ [78] mean $\pm$ its standard error computed over 100 random experiments for each problem. Figure 9 displays the images of 10 representative interpolation and deblurring experiments from the hundreds of experiments summarized in figures 6 and 7.

These results show that JPMAP outperforms all other baseline algorithms in terms of PSNR and LPIPS when random restarts are not allowed. When 10 random restarts are allowed for CSGM, but not for JPMAP, then both algorithms (JPMAP and mCSGM) show a similar global performance: JPMAP tends to provide a slightly better result than mCSGM except for the most extremely ill-posed interpolation, super-resolution and compressed sensing experiments (when available measurements are less than 10% the number of pixels). In that case mCSGM outperforms JPMAP by an equally small margin. The latter case can be explained by the fact that the encoder (which is used by JPMAP but not by CSGM) struggles

⁴MNIST images were zero-padded to 32×32 because LPIPS does not accept 28×28 images.



(a) Results on interpolation



(b) Results on deblurring.

Figure 9. Experimental results on MNIST. Comparison of JPMAP with the baseline algorithms described in Section 3.1. (a) Some selected results from the interpolation experiment with 80% of missing pixels and Gaussian noise with $\sigma = 10/255$. From top to bottom: original image x^* , corrupted image y , and the results computed by CSGM, mCSGM, PGD-GAN, PULSE, Algorithm 1.1 and JPMAP. (b) Same as (a) from the deblurring experiment with kernel size 3×3 and Gaussian noise with $\sigma = 10/255$. Conclusion: Our algorithm performs generally better than the baseline algorithms, although in some cases it falls behind mCSGM.

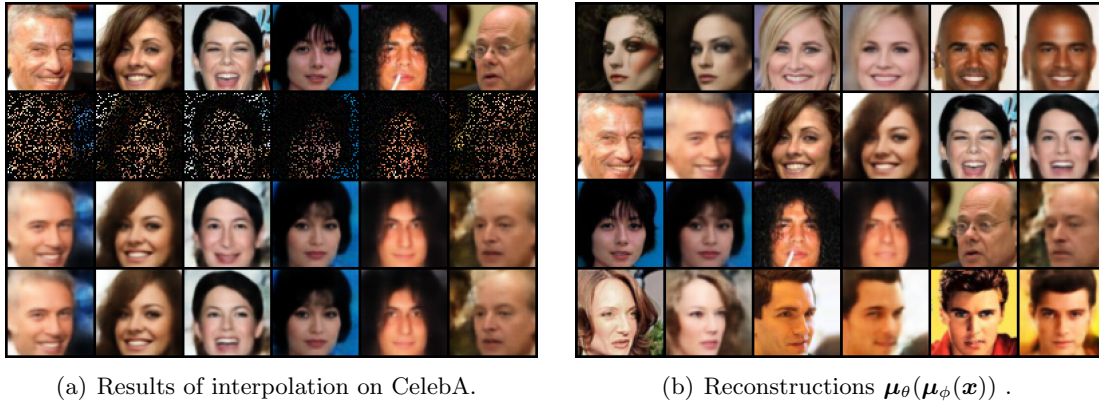


Figure 10. (a) Some preliminary results on CelebA: 80% of missing pixels, noise std $\sigma = 10/255$. From top to bottom: original image $\mathbf{x}^*$, corrupted image $\tilde{\mathbf{x}}$, restored by CSGM [6], restored image $\hat{\mathbf{x}}$ by our framework. (b) Reconstructions $\mu_\theta(\mu_\phi(\mathbf{x}))$ (even columns) for some test samples $\mathbf{x}$ (odd columns), showing the over-regularization of data manifold imposed by the trained vanilla VAE. As a consequence, $-\log p_{Z|Y}(\mathbf{z}|\mathbf{y})$ does not have as many local minima and then a simple gradient descent yields almost the same result as JPMAP (except on third column of (a)).

to generalize to images $\mathbf{x}$ which are very far away from $\mathcal{M}$ (the range of the generator). Indeed, in Section 3.4 we trained the VAE’s encoder to generalize to $\mathcal{M} + n$ where $n \sim \mathcal{N}(0, \sigma_{\text{DVAE}}^2 I)$ and $\sigma_{\text{DVAE}} = 15/255$. This value is optimal for moderately ill-posed problems, but more extreme problems may require larger values of σ_{DVAE} or a coarse to fine scheme, where a coarse VAE (with large σ_{DVAE}) is used during the first few iterations and a finer VAE (with smaller σ_{DVAE}) is used later until convergence. Finally one may consider using random restarts for both JPMAP and CSGM for a more fair comparison.

Figure 8 performs a more detailed comparison between JPMAP and mCSGM, which also considers running times of both algorithms. For the stopping criteria used in our experiments, one run of JPMAP requires roughly as much time as mCSGM (with $m = 10$ restarts). In addition for moderately ill-posed problems (like interpolation of 75% missing pixels see subfigures (a) and (b)) where JPMAP’s performance beats mCSGM, we can observe that JPMAP also converges much faster to that solution. For more extremely ill-posed problems (like interpolation of 90% missing pixels, see subfigures (c) and (d)) the opposite is true.

In the case of CelebA, we did not observe as much difference between JPMAP and CSGM as on MNIST. In Figure 10(a) the restorations on an interpolation problem (80% of missing pixels) are very similar to each other, but blurry. Also, although this problem is very ill-posed, both algorithms impressively find a solution $\mathbf{z}^*$ very close to the code $\mu_\phi(\mathbf{x})$ of the ground truth image $\mathbf{x}$, except for the third column where CSGM converges to a local minimum.

We hypothesize that, as CelebA is a substantially more complex dataset than MNIST, a simple model like vanilla VAE is over-regularizing the manifold of samples (underfitting problem). In particular, because of the spectral bias [55] the learned manifold perhaps only contains low-frequency approximations of the true images as we can see in the reconstructions $\mu_\theta(\mu_\phi(\mathbf{x}))$ of test samples (see Figure 10(b)). This may cause the posterior $p_{Z|Y}(\mathbf{z}|\mathbf{y})$ to

have fewer local minima. With more realistic generative models such as VDVAE [13] or NVAE [69], which better represent the true data manifold, we expect the objective function $-\log p_{Z|Y}(\mathbf{z}|\mathbf{y})$ to exhibit a much larger number of local minima, thus making it more difficult to optimize by a simple gradient descent scheme. In that situation the proposed JPMAP method would more clearly show its advantages.

4. Conclusions and Future work. In this work we presented a new framework to solve inverse problems with a convex data-fitting term and a non-convex regularizer learned in the latent space via variational autoencoders. Unlike similar approaches like CSGM [6], PULSE [46] and PGD-GAN [63] which learn the prior based on generative models, our approach is based on a generalization of alternate convex search to quasi-bi-convex functionals. This quasi-bi-convexity is the result of considering the joint posterior distribution of latent and image spaces. As a result, the proposed approach provides convergence guarantees that extend to a larger family of inverse problems. Experiments on denoising, interpolation, deconvolution, super-resolution and compressed sensing confirm this, since our approach gets stuck much less often in spurious local minima than CSGM, PGD-GAN or PULSE, which are simply based on gradient descent of a highly non-convex functional. This leads to restored images which are significantly better in terms of PSNR and LPIPS.

JPMAP vs related Plug & Play approaches. When compared to other decoupled *plug & play* approaches that solve inverse problems using NN-based priors, our approach is constrained in different ways:

(a) In a certain sense our approach is *less constrained* than existing decoupled approaches since we do not require to retrain the NN-based denoiser to enforce any particular property to ensure convergence: Ryu et al. [60] requires the denoiser’s residual operator to be non-expansive, and Gupta et al. [29], Shah and Hegde [63] require the denoiser to act as a projector. The effect of these modifications to the denoiser on the quality of the underlying image prior has never been studied in detail and chances are that such constraints degrade it. Our method only requires a variational autoencoder without any further constraints, and the quality and expressiveness of this prior can be easily checked by sampling and reconstruction experiments. Checking the quality of the prior is a much more difficult task for Ryu et al. [60], Gupta et al. [29], Shah and Hegde [63] which rely on an implicit prior, and do not provide a generative model.

(b) Unlike [60] which requires the data-fitting term $F(\mathbf{x})$ to be *strongly convex* to ensure convergence, our method admits weakly convex and ill-posed data-fitting terms like missing pixels, compressed sensing and non-invertible blurring for instance.

(c) On the other hand our method is *more constrained* in the sense that it relies on a generative model of a *fixed size*. Even if the generator and encoder are both convolutional neural networks, training and testing the same model on images of different sizes is a priori not possible because the latent space has a fixed dimension and a fixed distribution. As a future work we plan to explore different ways to address this limitation. The most straightforward way is to use our model to learn a prior of image patches of a fixed size and stitch this model via aggregation schemes like in EPLL [80] to obtain a global prior model for images of any size. Alternatively we can use hierarchical generative models like in [35, 69] or resizable ones like in [3, 72], and adapt our framework accordingly.

MAP- x or MAP- z or joint MAP- x - z . In this work we explored and clarified the tight relationships between joint MAP- x - z estimation, splitting and continuation schemes and the more common MAP- z estimator in the context of inverse problems with a generative prior. On the other hand MAP- x estimators (which are otherwise standard in bayesian imaging) remained largely unexplored in the context of generative priors, due to the optimization challenges they impose, until the recent work of Helminger et al. [31], Whang et al. [72] showed that a normalizing flow-based generative model allows to overcome those challenges and deems this problem tractable. Similarly Oberlin and Verm [47] use Glow (an invertible normalizing flow) to compare synthesis-based and analysis-based reconstructions. Yet an extensive comparison of the advantages and weaknesses of these three families of estimators under the same prior model is still missing, and so is the link between these MAP estimators and the analysis/synthesis-based estimators in [47]. This will be the subject of future work.

Extension to higher dimensional problems. The present paper provides a first proof of concept of our framework, on a very simple dataset (MNIST) with a very simple VAE. More experiments are needed to verify that the framework preserves its qualitative advantages on more high-dimensional datasets (like CelebA, FFHQ, etc.), and a larger selection of inverse problems.

Generalizing our proposed method to much higher dimensional problems implies training much more complex generative models which can match the finer details and higher complexity of such data. We can still use over-simplified generative models in those cases, but our preliminary experiments suggest that in that situation, not only do we obtain relatively poor reconstructions, but the objective function associated to the MAP- z problem presents less spurious local-minima: as a consequence our proposed joint MAP- x - z is overkill in that configuration, and does not present such a great competitive advantage.

The big challenge of generalizing our proposed method to much higher dimensional problems is then to train sufficiently detailed and complex generative models. And in this area VAEs traditionally lagged behind GANs in terms of quality of the generated samples, the former producing in general more blurred samples. Nevertheless some studies [61] show that VAEs and Normalizing Flows produce more accurate representations of the probability distribution. In the medium term our work should be able to benefit from recent advances in VAE architectures [13, 69, 16], and adversarial training for VAEs [52, 53, 79] that reach GAN-quality samples with the additional benefits of VAEs. These extensions are however non-trivial, since these VAEs have a huge number of parameters and they need to be retrained or fine-tuned using a denoising criterion (see section 3.4 and [33]) for our method to work properly. In addition, the latent space of the most competitive VAEs is much larger than the image space, which may reduce its regularization capabilities.

As an alternative, GAN-based generative models can be augmented with a denoising encoder network [19], and Normalizing Flows can also act as projectors or denoising VAEs if we split the latent space to separate the data manifold from its complement, as suggested in [8, 41]. In combination with relaxation techniques, such augmented GANs or specially tailored Flows may provide SOTA priors that fit our quasi-bi-convex optimization framework.

Towards stronger convergence guarantees under weaker conditions.. The proposed Algorithm 2.4 bears strong similarities with ADMM with non-linear constraints as introduced by Valkonen et al. [70, 2] and analyzed by Latorre-Gómez et al. [38]. Latorre-Gómez result

provides very strong convergence guarantees (linear convergence rates to a global optimum), but requires the data fitting term to be strongly convex or to satisfy a restricted strong convexity property. Our result, on the other hand, provides much weaker convergence guarantees (convergence to a stationary point), but does not require strong convexity. Further exploring these connections might hopefully lead to something closer to the best of both worlds.

Acknowledgments. We would like to sincerely thank Mauricio Delbracio, José Lezama and Pablo Musé for their help, their insightful comments, and their continuous support throughout this project.

References.

- [1] Cecilia Aguerrebere, Andres Almansa, Julie Delon, Yann Gousseau, and Pablo Muse. A Bayesian Hyperprior Approach for Joint Image Denoising and Interpolation, With an Application to HDR Imaging. *IEEE Transactions on Computational Imaging*, 3(4):633–646, dec 2017. ISSN 2333-9403. doi: 10.1109/TCI.2017.2704439. URL https://nounsse.github.io/HBE_project/.
- [2] Martin Benning, Florian Knoll, Carola Bibiane Schönlieb, and Tuomo Valkonen. Pre-conditioned ADMM with nonlinear operator constraint. *IFIP Advances in Information and Communication Technology*, 494:117–126, 2016. ISSN 18684238. doi: 10.1007/978-3-319-55795-3_10.
- [3] Urs Bergmann, Nikolay Jetchev, and Roland Vollgraf. Learning Texture Manifolds with the Periodic Spatial GAN. *(ICML) International Conference on Machine Learning*, 1: 722–730, may 2017.
- [4] Siavash Arjomand Bigdeli and Matthias Zwicker. Image Restoration using Autoencoding Priors. Technical report, 2017.
- [5] Siavash Arjomand Bigdeli, Meiguang Jin, Paolo Favaro, and Matthias Zwicker. Deep Mean-Shift Priors for Image Restoration. In *(NIPS) Advances in Neural Information Processing Systems 30*, pages 763–772, sep 2017. URL <http://papers.nips.cc/paper/6678-deep-mean-shift-priors-for-image-restoration>.
- [6] Ashish Bora, Ajil Jalal, Eric Price, and Alexandros G Dimakis. Compressed sensing using generative models. In *(ICML) International Conference on Machine Learning*, volume 2, pages 537–546. JMLR. org, 2017. ISBN 9781510855144.
- [7] Kristian Bredies, Karl Kunisch, and Thomas Pock. Total generalized variation. *SIAM Journal on Imaging Sciences*, 3(3):492–526, 2010. ISSN 19364954. doi: 10.1137/090769521.
- [8] Johann Brehmer and Kyle Cranmer. Flows for simultaneous manifold learning and density estimation. mar 2020. URL <http://arxiv.org/abs/2003.13913>.
- [9] Gregory T Buzzard, Stanley H Chan, Suhas Sreehari, and Charles A Bouman. Plug-and-play unplugged: Optimization-free reconstruction using consensus equilibrium. *SIAM Journal on Imaging Sciences*, 11(3):2001–2020, 2018.
- [10] A Chambolle. An algorithm for total variation minimization and applications. *Journal of Mathematical Imaging and Vision*, 20:89–97, 2004. doi: 10.1023/B:JMIV.0000011325.36760.1e.
- [11] S. H. Chan, X. Wang, and O. A. Elgendy. Plug-and-play admm for image restoration:

- Fixed-point convergence and applications. *IEEE Transactions on Computational Imaging*, 3(1):84–98, March 2017. ISSN 2333-9403. doi: 10.1109/TCI.2016.2629286.
- [12] Yunjin Chen and Thomas Pock. Trainable Nonlinear Reaction Diffusion: A Flexible Framework for Fast and Effective Image Restoration. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 39(6):1256–1272, 2017. ISSN 01628828. doi: 10.1109/TPAMI.2016.2596743.
 - [13] Rewon Child. Very Deep VAEs Generalize Autoregressive Models and Can Outperform Them on Images. In *(ICLR) International Conference on Learning Representations*, pages 1–17, 2020. URL <https://openreview.net/forum?id=RLRXCV6DbEJ>.
 - [14] Djork-Arné Clevert, Thomas Unterthiner, and Sepp Hochreiter. Fast and Accurate Deep Network Learning by Exponential Linear Units (ELUs). In *(ICLR) International Conference on Learning Representations*, nov 2016.
 - [15] Regev Cohen, Michael Elad, and Peyman Milanfar. Regularization by Denoising via Fixed-Point Projection (RED-PRO). aug 2020. ISSN 23318422. URL <http://arxiv.org/abs/2008.00226>.
 - [16] Bin Dai and David Wipf. Diagnosing and Enhancing VAE Models. *ICLR*, pages 1–42, 2019. URL <https://openreview.net/forum?id=B1e0X3C9tQ>.
 - [17] Gianni Dal Maso. *An Introduction to Γ -Convergence*. Birkhäuser Boston, Boston, MA, 1993. ISBN 978-1-4612-6709-6. doi: 10.1007/978-1-4612-0327-8. URL <http://link.springer.com/10.1007/978-1-4612-0327-8>.
 - [18] Steven Diamond, Vincent Sitzmann, Felix Heide, and Gordon Wetzstein. Unrolled optimization with deep priors. 2017.
 - [19] Jeff Donahue and Karen Simonyan. Large Scale Adversarial Representation Learning. 2019. URL <http://arxiv.org/abs/1907.02544>.
 - [20] Chao Dong, Chen Change Loy, Kaiming He, and Xiaoou Tang. Learning a deep convolutional network for image super-resolution. In *European conference on computer vision*, pages 184–199. Springer, 2014.
 - [21] D.L. Donoho. De-noising by soft-thresholding. *IEEE Transactions on Information Theory*, 41(3):613–627, may 1995. ISSN 00189448. doi: 10.1109/18.382009.
 - [22] Michael Elad. *Sparse and Redundant Representations: From Theory to Applications in Signal and Image Processing*. Springer New York, New York, NY, 2010. ISBN 978-1-4419-7011-4. doi: 10.1007/978-1-4419-7011-4.
 - [23] Hongyun Gao, Xin Tao, Xiaoyong Shen, and Jiaya Jia. Dynamic scene deblurring with parameter selective sharing and nested skip connections. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3848–3856, 2019.
 - [24] Michaël Gharbi, Gaurav Chaurasia, Sylvain Paris, and Frédo Durand. Deep joint demosaicking and denoising. *ACM Transactions on Graphics (TOG)*, 35(6):191, 2016.
 - [25] Guy Gilboa and Stanley Osher. Nonlocal operators with applications to image processing. *Multiscale Modeling and Simulation*, 7(3):1005–1028, 2008. ISSN 15403459. doi: 10.1137/070698592.
 - [26] Davis Gilton, Greg Ongie, and Rebecca Willett. Neumann networks for inverse problems in imaging. 2019.
 - [27] Jochen Gorski, Frank Pfeuffer, and Kathrin Klamroth. Biconvex sets and optimization with biconvex functions: a survey and extensions. *Mathematical Methods of Operations*

- Research*, 66(3):373–407, nov 2007. ISSN 1432-2994. doi: 10.1007/s00186-007-0161-1.
- [28] Karol Gregor and Yann LeCun. Learning fast approximations of sparse coding. In *Proceedings of the 27th International Conference on International Conference on Machine Learning*, pages 399–406. Omnipress, 2010.
 - [29] Harshit Gupta, Kyong Hwan Jin, Ha Q Nguyen, Michael T McCann, and Michael Unser. Cnn-based projected gradient descent for consistent ct image reconstruction. *IEEE transactions on medical imaging*, 37(6):1440–1453, 2018. doi: 10.1109/TMI.2018.2832656.
 - [30] Paul Hand and Vladislav Voroninski. Global Guarantees for Enforcing Deep Generative Priors by Empirical Risk. *IEEE Transactions on Information Theory*, 66(1):401–418, 2020. ISSN 15579654. doi: 10.1109/TIT.2019.2935447.
 - [31] Leonhard Helminger, Michael Bernasconi, Abdelaziz Djelouah, Markus Gross, and Christopher Schroers. Blind Image Restoration with Flow Based Priors. Technical report, sep 2020. URL <http://arxiv.org/abs/2009.04583>.
 - [32] Wen Huang, Paul Hand, Reinhard Heckel, and Vladislav Voroninski. A Provably Convergent Scheme for Compressive Sensing under Random Generative Priors. dec 2018. URL <http://arxiv.org/abs/1812.04176>.
 - [33] Daniel Jiwoong Im, Sungjin Ahn, Roland Memisevic, and Yoshua Bengio. Denoising criterion for variational auto-encoding framework. In *31st AAAI Conference on Artificial Intelligence, AAAI 2017*, pages 2059–2065. AAAI press, nov 2017.
 - [34] Ulugbek S Kamilov, Hassan Mansour, and Brendt Wohlberg. A plug-and-play priors approach for solving nonlinear imaging inverse problems. *IEEE Signal Processing Letters*, 24(12):1872–1876, 2017.
 - [35] Tero Karras, Timo Aila, Samuli Laine, and Jaakko Lehtinen. Progressive Growing of GANs for Improved Quality, Stability, and Variation. *(ICLR) International Conference on Learning Representations*, 10(2):327–331, oct 2017. URL <https://openreview.net/forum?id=Hk99zCeAb>.
 - [36] Diederik P Kingma and Max Welling. Auto-Encoding Variational Bayes. In *(ICLR) International Conference on Learning Representations*, number ML, pages 1–14, dec 2013. ISBN 1312.6114v10. doi: 10.1051/0004-6361/201527329.
 - [37] Alex Krizhevsky, Ilya Sutskever, and Geoffrey E. Hinton. Imagenet classification with deep convolutional neural networks. *(NIPS) Advances in neural information processing systems*, pages 1097–1105, 2012. ISSN 10495258.
 - [38] Fabian Latorre, Armin Eftekhar, Volkan Cevher, Fabian Latorre Gómez, Armin Eftekhar, and Volkan Cevher. Fast and Provable ADMM for Learning with Generative Priors. In H Wallach, H Larochelle, A Beygelzimer, F d\textquotesingle Alché-Buc, E Fox, and R Garnett, editors, *(NeurIPS) Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. URL <https://papers.nips.cc/paper/2019/hash/4559912e7a94a9c32b09d894f2bc3c82-Abstract.html>.
 - [39] Rémi Laumont, Valentin De Bortoli, Andrés Almansa, Julie Delon, Alain Durmus, and Marcelo Pereyra. On Maximum-a-Posteriori estimation with Plug & Play priors and stochastic gradient descent. sep 2021. URL <https://hal.archives-ouvertes.fr/hal-03348735>.
 - [40] Yann Lecun, Leon Bottou, Yoshua Bengio, and P. Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998. ISSN 00189219. doi: 10.1109/5.726791.

- [41] Yang Liu, Saeed Anwar, Zhenyue Qin, Pan Ji, Sabrina Caldwell, and Tom Gedeon. Disentangling Noise from Images: A Flow-Based Image Denoising Neural Network. may 2021. URL <https://arxiv.org/abs/2105.04746v1>.
- [42] Ziwei Liu, Ping Luo, Xiaogang Wang, and Xiaoou Tang. Deep learning face attributes in the wild. In *Proceedings of International Conference on Computer Vision (ICCV)*, December 2015.
- [43] Cécile Louchet and Lionel Moisan. Posterior expectation of the total variation model: Properties and experiments. *SIAM Journal on Imaging Sciences*, 6(4):2640–2684, dec 2013. ISSN 19364954. doi: 10.1137/120902276.
- [44] James Lucas, George Tucker, Roger Grosse, and Mohammad Norouzi. Don’t blame the ELBO! A linear VAE perspective on posterior collapse. In *Advances in Neural Information Processing Systems*, volume 32, nov 2019. URL <https://arxiv.org/abs/1911.02469>.
- [45] Tim Meinhardt, Michael Moller, Caner Hazirbas, and Daniel Cremers. Learning proximal operators: Using denoising networks for regularizing inverse imaging problems. In *(ICCV) International Conference on Computer Vision*, pages 1781–1790, 2017. doi: 10.1109/ICCV.2017.198. URL http://openaccess.thecvf.com/content_iccv_2017/html/Meinhardt_Learning_Proximal_Operators_ICCV_2017_paper.html.
- [46] Sachit Menon, Alexandru Damian, Shijia Hu, Nikhil Ravi, and Cynthia Rudin. PULSE: Self-Supervised Photo Upsampling via Latent Space Exploration of Generative Models. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 2434–2442, 2020. ISSN 10636919. doi: 10.1109/CVPR42600.2020.00251.
- [47] Thomas Oberlin and Mathieu Verm. Regularization via deep generative models: an analysis point of view. jan 2021. URL <http://arxiv.org/abs/2101.08661>.
- [48] George Papamakarios, Eric Nalisnick, Danilo Jimenez Rezende, Shakir Mohamed, and Balaji Lakshminarayanan. Normalizing Flows for Probabilistic Modeling and Inference. 2019.
- [49] Adam Paszke, Sam Gross, Soumith Chintala, Gregory Chanan, Edward Yang, Zachary DeVito, Zeming Lin, Alban Desmaison, Luca Antiga, and Adam Lerer. Automatic differentiation in pytorch. 2017.
- [50] Marcelo Pereyra. Proximal Markov chain Monte Carlo algorithms. *Statistics and Computing*, 26(4):745–760, jul 2016. ISSN 0960-3174. doi: 10.1007/s11222-015-9567-4. URL <http://dx.doi.org/10.1007/s11222-015-9567-4>.
- [51] Jean-Christophe Pesquet, Audrey Repetti, Matthieu Terris, and Yves Wiaux. Learning Maximally Monotone Operators for Image Recovery. 2020. URL <http://arxiv.org/abs/2012.13247>.
- [52] Yunchen Pu, Weiyao Wang, Ricardo Henao, Liquan Chen, Zhe Gan, Chunyuan Li, and Lawrence Carin. Adversarial symmetric variational autoencoder. In *(NIPS) Advances in Neural Information Processing Systems*, pages 4331–4340, 2017.
- [53] Yunchen Pu, Weiyao Wang, Ricardo Henao, Liquan Chen, Zhe Gan, Chunyuan Li, and Lawrence Carin. Adversarial symmetric variational autoencoder. In *(NIPS) Advances in Neural Information Processing Systems*, volume 2017-Decem, pages 4331–4340, 2017.
- [54] Alec Radford, Luke Metz, and Soumith Chintala. Unsupervised representation learning with deep convolutional generative adversarial networks. *arXiv preprint*

- arXiv:1511.06434*, 2015.
- [55] Nasim Rahaman, Aristide Baratin, Devansh Arpit, Felix Draxler, Min Lin, Fred Hamprecht, Yoshua Bengio, and Aaron Courville. On the spectral bias of neural networks. In *International Conference on Machine Learning*, pages 5301–5310. PMLR, 2019.
 - [56] Ankit Raj, Yuqi Li, and Yoram Bresler. GAN-Based Projector for Faster Recovery With Convergence Guarantees in Linear Inverse Problems. In *(ICCV) International Conference on Computer Vision*, pages 5601–5610. IEEE, oct 2019. ISBN 978-1-7281-4803-8. doi: 10.1109/ICCV.2019.00570.
 - [57] Edward T Reehorst and Philip Schniter. Regularization by denoising: Clarifications and new interpretations. *IEEE Transactions on Computational Imaging*, 5(1):52–67, 2018. doi: 10.1109/TCI.2018.2880326.
 - [58] Yaniv Romano, Michael Elad, and Peyman Milanfar. The little engine that could: Regularization by denoising (red). *SIAM Journal on Imaging Sciences*, 10(4):1804–1844, 2017.
 - [59] Leonid I. Rudin, Stanley Osher, and Emad Fatemi. Nonlinear total variation based noise removal algorithms. *Physica D: Nonlinear Phenomena*, 60(1-4):259–268, 1992. ISSN 01672789. doi: 10.1016/0167-2789(92)90242-F.
 - [60] Ernest K. Ryu, Jialin Liu, Sicheng Wang, Xiaohan Chen, Zhangyang Wang, and Wotao Yin. Plug-and-play methods provably converge with properly trained denoisers. In *Proceedings of the 36th International Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA*, pages 5546–5557, 2019. URL <http://proceedings.mlr.press/v97/ryu19a.html>.
 - [61] Mehdi S. M. Sajjadi, Olivier Bachem, Mario Lucic, Olivier Bousquet, and Sylvain Gelly. Assessing Generative Models via Precision and Recall. In *(NeurIPS) Neural Information Processing Systems*, may 2018.
 - [62] Eli Schwartz, Raja Giryes, and Alex M Bronstein. Deepisp: Toward learning an end-to-end image processing pipeline. *IEEE Transactions on Image Processing*, 28(2):912–923, 2018.
 - [63] Viraj Shah and Chinmay Hegde. Solving linear inverse problems using gan priors: An algorithm with provable guarantees. In *2018 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 4609–4613. IEEE, 2018.
 - [64] Suhas Sreehari, Singanallur V. Venkatakrishnan, Brendt Wohlberg, Gregory T. Buzzard, Lawrence F. Drummy, Jeffrey P. Simmons, and Charles A. Bouman. Plug-and-Play Priors for Bright Field Electron Tomography and Sparse Interpolation. *IEEE Transactions on Computational Imaging*, 2(4):1–1, 2016. ISSN 2333-9403. doi: 10.1109/TCI.2016.2599778.
 - [65] Afonso M. Teodoro, José M. Bioucas-Dias, and Mário A. T. Figueiredo. Scene-Adapted Plug-and-Play Algorithm with Guaranteed Convergence: Applications to Data Fusion in Imaging, jan 2018.
 - [66] Matthieu Terris, Audrey Repetti, Jean Christophe Pesquet, and Yves Wiaux. Building firmly nonexpansive convolutional neural networks. *ICASSP, IEEE International Conference on Acoustics, Speech and Signal Processing - Proceedings*, 2020-May:8658–8662, 2020. ISSN 15206149. doi: 10.1109/ICASSP40776.2020.9054731.
 - [67] A. N. Tikhonov. On the regularization of ill-posed problems. *Dokl. Akad. Nauk SSSR*,

- 39(1):195–198, 1943.
- [68] Paul Tseng and Dimitri P. Bertsekas. On the convergence of the exponential multiplier method for convex programming. *Mathematical Programming*, 60(1-3):1–19, jun 1993. ISSN 00255610. doi: 10.1007/BF01580598. URL <https://link.springer.com/article/10.1007/BF01580598>.
 - [69] Arash Vahdat and Jan Kautz. Nvae: A deep hierarchical variational autoencoder. *Advances in Neural Information Processing Systems*, 33, 2020.
 - [70] Tuomo Valkonen. A primal-dual hybrid gradient method for nonlinear operators with applications to MRI. *Inverse Problems*, 30(5):1–42, 2014. ISSN 13616420. doi: 10.1088/0266-5611/30/5/055012.
 - [71] Singanallur V. Venkatakrishnan, Charles A. Bouman, and Brendt Wohlberg. Plug-and-Play priors for model based reconstruction. *2013 IEEE Global Conference on Signal and Information Processing, GlobalSIP 2013 - Proceedings*, pages 945–948, 2013. doi: 10.1109/GlobalSIP.2013.6737048.
 - [72] Jay Whang, Qi Lei, and Alexandros G. Dimakis. Compressed Sensing with Invertible Generative Models and Dependent Noise. In *NeurIPS deep-inverse workshop*, 2020.
 - [73] Xiaojian Xu, Yu Sun, Jiaming Liu, Brendt Wohlberg, and Ulugbek S. Kamilov. Provable Convergence of Plug-and-Play Priors with MMSE denoisers. (4):1–10, 2020. URL <http://arxiv.org/abs/2005.07685>.
 - [74] Guoshen Yu, Guillermo Sapiro, and Stéphane Mallat. Solving inverse problems with piecewise linear estimators: From gaussian mixture models to structured sparsity. *IEEE Transactions on Image Processing*, 21(5):2481–2499, 2011.
 - [75] Kai Zhang, Wangmeng Zuo, Yunjin Chen, Deyu Meng, and Lei Zhang. Beyond a gaussian denoiser: Residual learning of deep cnn for image denoising. *IEEE Transactions on Image Processing*, 26(7):3142–3155, 2017.
 - [76] Kai Zhang, Wangmeng Zuo, Shuhang Gu, and Lei Zhang. Learning Deep CNN Denoiser Prior for Image Restoration. In *(CVPR) IEEE Conference on Computer Vision and Pattern Recognition*, pages 2808–2817. IEEE, apr 2017. ISBN 978-1-5386-0457-1. doi: 10.1109/CVPR.2017.300. URL http://openaccess.thecvf.com/content_cvpr_2017/html/Zhang_Learning_Deep_CNN.CVPR.2017_paper.html.
 - [77] Kai Zhang, Wangmeng Zuo, and Lei Zhang. Ffdnet: Toward a fast and flexible solution for cnn-based image denoising. *IEEE Transactions on Image Processing*, 27(9):4608–4622, 2018.
 - [78] Richard Zhang, Phillip Isola, Alexei A Efros, Eli Shechtman, and Oliver Wang. The unreasonable effectiveness of deep features as a perceptual metric. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 586–595, 2018.
 - [79] Zijun Zhang, Ruixiang Zhang, Zongpeng Li, Yoshua Bengio, and Liam Paull. Perceptual Generative Autoencoders. In *(ICLR) International Conference on Learning Representations*, pages 1–7, jun 2019. URL <https://github.com/zj10/PGA>.
 - [80] Daniel Zoran and Yair Weiss. From learning models of natural image patches to whole image restoration. In *2011 International Conference on Computer Vision*, pages 479–486. IEEE, nov 2011. ISBN 978-1-4577-1102-2. doi: 10.1109/ICCV.2011.6126278. URL <http://people.csail.mit.edu/danielzoran/EPLLICCVCameraReady.pdf>.

Appendix A. Properties of J_1 . In this section, we establish that the objective function J_1 fulfills the assumptions required to prove the convergence of Algorithm 3, namely

- $J_1(\cdot, \mathbf{z})$ is convex for any $\mathbf{z}$;
- $J_1(\cdot, \mathbf{z})$ has a unique minimizer for any $\mathbf{z}$;
- J_1 is coercive;
- J_1 is continuously differentiable;

We recall that

$$\begin{aligned} J_1(\mathbf{x}, \mathbf{z}) &= \underbrace{\frac{1}{2\sigma^2} \|\mathbf{A}\mathbf{x} - \mathbf{y}\|^2}_{F(\mathbf{x}, \mathbf{y})} \\ &\quad + \underbrace{\frac{1}{2} \left(Z_\theta(\mathbf{z}) + \|\Sigma_\theta^{-1/2}(\mathbf{z})(\mathbf{x} - \boldsymbol{\mu}_\theta(\mathbf{z}))\|^2 \right)}_{H_\theta(\mathbf{x}, \mathbf{z})} \\ &\quad + \frac{1}{2} \|\mathbf{z}\|^2 \end{aligned}$$

where

$$Z_\theta(\mathbf{z}) = d \log(2\pi) + \log \det \Sigma_\theta(\mathbf{z})$$

Thus, it is the sum of three non-negative terms.

A.1. Convexity and unicity of the minimizer of $J_1(\cdot, \mathbf{z})$. Let $\mathbf{z}$ be fixed. Then there exists a constant $C \in \mathbb{R}$ such that $\forall \mathbf{x}$

$$J_1(\mathbf{x}, \mathbf{z}) = \frac{1}{2\sigma^2} \|\mathbf{A}\mathbf{x} - \mathbf{y}\|^2 + \|\Sigma_\theta^{-1/2}(\mathbf{z})(\mathbf{x} - \boldsymbol{\mu}_\theta(\mathbf{z}))\|^2 + C$$

Being the sum of two quadratic forms, $J_1(\cdot, \mathbf{z})$ is obviously twice differentiable. Its gradient is given by

$$\begin{aligned} \frac{\partial J_1}{\partial \mathbf{x}}(\mathbf{x}, \mathbf{z}) &= \frac{1}{\sigma^2} \mathbf{A}^T (\mathbf{A}\mathbf{x} - \mathbf{y}) \\ &\quad + 2 (\Sigma_\theta^{-1/2}(\mathbf{z}))^T (\Sigma_\theta^{-1/2}(\mathbf{z})(\mathbf{x} - \boldsymbol{\mu}_\theta(\mathbf{z}))) \end{aligned}$$

and its Hessian is

$$\text{Hess}_{\mathbf{x}} J_1(\mathbf{x}, \mathbf{z}) = \frac{1}{\sigma^2} \mathbf{A}^T \mathbf{A} + 2 (\Sigma_\theta^{-1/2}(\mathbf{z}))^T \Sigma_\theta^{-1/2}(\mathbf{z})$$

Since $\Sigma_\theta(\mathbf{z}) = \gamma^2 I$ the Hessian is positive definite (without the need to assume that $\mathbf{A}$ is full rank), and we have that

Lemma A.1. $J_1(\cdot, \mathbf{z})$ is strictly convex for any $\mathbf{z}$.

An immediate consequence is the unicity of the minimizer of the partial function $J_1(\cdot, \mathbf{z})$.

A.2. Coercivity of J_1 .

Lemma A.2. J_1 is coercive.

Proof. First, let us note that J_1 is the sum of three non-negative terms. If it was not coercive, then we could find a sequence $(\mathbf{x}_k, \mathbf{z}_k) \rightarrow \infty$ such that $J_1(\mathbf{x}_k, \mathbf{z}_k)$ is bounded. As a consequence all three terms are bounded. In particular the last term $\|\mathbf{z}_k\|$ is bounded, which means that $\mathbf{x}_k \rightarrow \infty$. From Property 1, $\{\boldsymbol{\mu}_\theta(\mathbf{z}_k)\}$ and $\{\boldsymbol{\Sigma}_\theta(\mathbf{z}_k)\}$ are bounded for bounded $\{\mathbf{z}_k\}$. Now, from the definition of the second term of J_1 , we get that, $\{\boldsymbol{\mu}_\theta(\mathbf{z}_k)\}$ and $\{\boldsymbol{\Sigma}_\theta(\mathbf{z}_k)\}$ being bounded and $\mathbf{x}_k$ going to ∞ yield that $H_\theta(\mathbf{x}_k, \mathbf{z}_k)$ goes to infinity, while being bounded. This leads to a contradiction and thus proves that J_1 is coercive.

A.3. Regularity of J_1 . In the sequel we adopt the common assumption that all neural networks used in this work are composed of a finite number d of layers, each layer being composed of: (a) a linear operator (e.g. convolutional or fully connected layer), followed by (b) a non-linear L -Lipschitz component-wise activation function with $0 < L < \infty$.

Therefore we have the following property:

Property 1. For any neural network f_θ with parameters θ having the structure described above:

There exists a constant C_θ such that $\forall \mathbf{u}$,

$$\|f_\theta(\mathbf{u})\|_2 \leq C_\theta \|\mathbf{u}\|_2.$$

Concerning activation functions we use two kinds:

- continuously differentiable activations like ELU, or
- continuous but non-differentiable activations like ReLU

Hence, by composition, we have that

Lemma A.3. For continuously differentiable activation functions, J_1 is continuously differentiable.

Appendix B. MAP-x and MAP-z for deterministic generative models. Assume that the stochastic γ -generative model is

$$p_{X_\gamma|Z_\gamma}(\mathbf{x}|\mathbf{z}) = \mathcal{N}(\mathbf{G}(\mathbf{z}), \gamma^2 \mathbf{I})$$

meaning that when $\gamma \rightarrow 0$

$$p_{X|Z}(\mathbf{x}|\mathbf{z}) = \delta(\mathbf{x} - \mathbf{G}(\mathbf{z}))$$

We now analyze the MAP- $\mathbf{z}$ and MAP- $\mathbf{x}$ estimators for the limit case when $\gamma = 0$. This is what we call a deterministic generative model, and it includes GANs for instance.

B.1. MAP-z. By definition the MAP- $\mathbf{z}$ estimator is obtained by maximising the posterior with respect to $\mathbf{z}$:

$$\begin{aligned} \hat{\mathbf{z}}_{\text{MAP-z}} &= \arg \max_{\mathbf{z}} \{p_{Z|Y}(\mathbf{z}|\mathbf{y})\} \\ &= \arg \max_{\mathbf{z}} \{p_{Y|Z}(\mathbf{y}|\mathbf{z}) p_Z(\mathbf{z})\}. \end{aligned} \tag{B.1}$$

In the last line we used Bayes rule to rewrite this posterior in more simple terms. However, this expression still involves the unknown conditional $p_{Y|Z}(\mathbf{y} | \mathbf{z})$.

Let us express this maximization in terms of $p_{Y|X}(\mathbf{y} | \mathbf{x})$.

To do so we recall the relation between the conditionals and the joint:

$$(B.2) \quad p_{Y|Z}(\mathbf{y} | \mathbf{z}) p_Z(\mathbf{z}) = p_{Y,Z}(\mathbf{y}, \mathbf{z}) = p_{Z|Y}(\mathbf{z} | \mathbf{y}) p_Y(\mathbf{y})$$

We can also compute the joint distribution $p_{Y,Z}(\mathbf{y}, \mathbf{z})$ by marginalization on a third random variable X :

$$(B.3) \quad \begin{aligned} p_{Y,Z}(\mathbf{y}, \mathbf{z}) &= \int p_{X,Y,Z}(\mathbf{x}, \mathbf{y}, \mathbf{z}) d\mathbf{x} \\ &= \int p_{Y|X,Z}(\mathbf{y} | \mathbf{x}, \mathbf{z}) p_{X|Z}(\mathbf{x} | \mathbf{z}) p_Z(\mathbf{z}) d\mathbf{x} \\ &= \int p_{Y|X}(\mathbf{y} | \mathbf{x}) \delta(\mathbf{x} - \mathbf{G}(\mathbf{z})) p_Z(\mathbf{z}) d\mathbf{x} \\ &= p_{Y|X}(\mathbf{y} | \mathbf{G}(\mathbf{z})) p_Z(\mathbf{z}) \end{aligned}$$

The third line follows from our graphical model $Z \rightarrow X \rightarrow Y$ which implies that once we know $X = \mathbf{x}$, then Z provides no additional information, therefore

$$p_{Y|X,Z}(\mathbf{y} | \mathbf{x}, \mathbf{z}) = p_{Y|X}(\mathbf{y} | \mathbf{x}).$$

The last line follows simply from the integration on $\mathbf{x}$ of a delta function.

From equations (B.2) and (B.3) we can derive an expression of $p_{Z|Y}(\mathbf{z} | \mathbf{y})$ in terms of $p_{Y|X}(\cdot | \cdot)$ and the generator $\mathbf{G}$ namely:

$$p_{Z|Y}(\mathbf{z} | \mathbf{y}) = \frac{1}{p_Y(\mathbf{y})} p_{Y|X}(\mathbf{y} | \mathbf{G}(\mathbf{z})) p_Z(\mathbf{z})$$

This proves the main result of this section:

Proposition B.1 (MAP- $\mathbf{z}$ estimator for deterministic generative models). *Assume we have*

- *a deterministic generative model where $X = \mathbf{G}(Z)$ and*
- *an inverse problem characterised by the log conditional distribution $\log p_{Y|X}(\mathbf{y} | \mathbf{x}) = -F(\mathbf{x}, \mathbf{y})$.*

Then the MAP- $\mathbf{z}$ estimator is computed as $\hat{\mathbf{x}}_{\text{MAP-}\mathbf{z}} = \mathbf{G}(\hat{\mathbf{z}}_{\text{MAP-}\mathbf{z}})$ where

$$(B.4) \quad \begin{aligned} \hat{\mathbf{z}}_{\text{MAP-}\mathbf{z}} &= \arg \max_{\mathbf{z}} \{p_{Y|X}(\mathbf{y} | \mathbf{G}(\mathbf{z})) p_Z(\mathbf{z})\} \\ &= \arg \min_{\mathbf{z}} \{F(\mathbf{G}(\mathbf{z}), \mathbf{y}) - \log p_Z(\mathbf{z})\}. \end{aligned}$$

B.2. MAP- $\mathbf{x}$. The MAP- $\mathbf{x}$ estimator is obtained by maximizing the posterior with respect to $\mathbf{x}$. The generative model induces a prior on X via the push-forward measure $p_X = \mathbf{G}_\# p_Z$, which following [48, section 5] can be developed as

$$p_X(\mathbf{x}) = \frac{p_Z(\mathbf{G}^{-1}(\mathbf{x}))}{\sqrt{\det S(\mathbf{G}^{-1}(\mathbf{x}))}} \delta_{\mathcal{M}}(\mathbf{x})$$

where $S = \left(\frac{\partial \mathbf{G}}{\partial \mathbf{z}}\right)^T \left(\frac{\partial \mathbf{G}}{\partial \mathbf{z}}\right)$ is the squared Jacobian and the manifold $\mathcal{M} = \{\mathbf{x} : \exists \mathbf{z}, \mathbf{x} = \mathbf{G}(\mathbf{z})\}$ represents the image of the generator $\mathbf{G}$.

With such a prior p_X , the $\mathbf{x}$ -optimization (1.2) required to obtain $\hat{\mathbf{x}}_{\text{MAP}}$ becomes intractable (in general), for various reasons:

- the computation of $\det S$,
- the inversion of $\mathbf{G}$, and
- the hard constraint $\mathbf{x} \in \mathcal{M}$.

These operations are all memory and/or computationally intensive, except when they are partially addressed by the use of a normalizing flow like in [31, 72].

Appendix C. Joint MAP- $\mathbf{x}$ - $\mathbf{z}$, Continuation Scheme and convergence to MAP- $\mathbf{z}$.

The functional $J_{1,\beta}$ introduced in Equation 1.6 can be seen from two different perspectives.

From a machine learning perspective it corresponds to the joint log-posterior J_1 in the case where $\Sigma_\theta(\mathbf{z}) = \frac{1}{\beta} I$ and $\mu_\theta = \mathbf{G}$, namely:

$$\begin{aligned} J_{1,\beta}(\mathbf{x}, \mathbf{z}) &= \underbrace{\frac{1}{2\sigma^2} \|\mathbf{A}\mathbf{x} - \mathbf{y}\|^2}_{F(\mathbf{x}, \mathbf{y})} \\ &+ \underbrace{\frac{\beta}{2} \|\mathbf{x} - \mu_\theta(\mathbf{z})\|^2}_{H_\theta(\mathbf{x}, \mathbf{z}) = \varphi_\beta(\mathbf{x}, \mathbf{z})} \\ &+ \frac{1}{2} \|\mathbf{z}\|^2 + C_\beta \end{aligned}$$

From an optimization standpoint it can be considered as an inexact penalisation procedure: We want to solve the constrained problem

$$\min_{(\mathbf{x}, \mathbf{z}) \in \mathcal{C}} \underbrace{F(\mathbf{x}, \mathbf{y}) + \frac{1}{2} \|\mathbf{z}\|^2}_{=J_{1,0}(\mathbf{x}, \mathbf{z})}$$

with $\mathcal{C} = \{(\mathbf{x}, \mathbf{z}) \mid \mathbf{x} = \mu_\theta(\mathbf{z})\}$ whose solution provides the MAP- $\mathbf{z}$ estimator

$$(C.1) \quad (\mathbf{x}^*, \mathbf{z}^*) \in \arg \min_{(\mathbf{x}, \mathbf{z}) \in \mathcal{C}} J_{1,0}(\mathbf{x}, \mathbf{z}).$$

To do so, we introduced the family of unconstrained problems

$$\min_{\mathbf{x}, \mathbf{z}} J_{1,\beta}(\mathbf{x}, \mathbf{z})$$

and their corresponding minimizers

$$(\hat{\mathbf{x}}_\beta, \hat{\mathbf{z}}_\beta) \in \arg \min_{\mathbf{x}, \mathbf{z}} J_{1,\beta}(\mathbf{x}, \mathbf{z})$$

which for $\beta = \frac{1}{\gamma^2}$ provide the MAP- $\mathbf{x}$ - $\mathbf{z}$ estimator.

We can show that the MAP- $\mathbf{x}$ - $\mathbf{z}$ estimator converges to the MAP- $\mathbf{z}$ estimator when $\beta \rightarrow \infty$ (or equivalently $\gamma \rightarrow 0$).

Proposition C.1. *The unconstrained functional tends to the constrained functional plus the constraint:*

$$(C.2) \quad J_{1,\beta}(\mathbf{x}, \mathbf{z}) \xrightarrow{\beta \rightarrow \infty} J_{1,\infty}(\mathbf{x}, \mathbf{z}) = F(\mathbf{x}, \mathbf{z}) + \iota_{\mathbf{x}=\mathbf{G}(\mathbf{z})}(\mathbf{x}, \mathbf{z}) + \frac{1}{2} \|\mathbf{z}\|^2$$

and the unconstrained minimizers tend to the constrained minimizer as $\beta \rightarrow \infty$:

$$(C.3) \quad \lim_{\beta \rightarrow \infty} (\hat{\mathbf{x}}_\beta, \hat{\mathbf{z}}_\beta) \in \arg \min_{(\mathbf{x}, \mathbf{z}) \in \mathcal{C}} J_{1,0}(\mathbf{x}, \mathbf{z}) = \arg \min_{\mathbf{x}, \mathbf{z}} J_{1,\infty}(\mathbf{x}, \mathbf{z}).$$

Proof. The pointwise convergence of φ_β to $\iota_{\mathbf{x}=\mathbf{G}(\mathbf{z})}$ as β goes to ∞ is straightforward.

Let us first prove that for any sequence $(\beta_n)_n$ that goes to ∞ , the quantity $\|\hat{\mathbf{x}}_{\beta_n} - \mathbf{G}(\hat{\mathbf{z}}_{\beta_n})\|$ goes to zero. Otherwise, for any $\varepsilon > 0$, there exists a subsequence $(\beta_{n_j})_j$ such that $\|\hat{\mathbf{x}}_{\beta_{n_j}} - \mathbf{G}(\hat{\mathbf{z}}_{\beta_{n_j}})\| > \varepsilon$. In this case, for any $\mathbf{z}$, one has by optimality

$$J_{1,0}(\mathbf{G}(\mathbf{z}), \mathbf{z}) = J_{1,\beta_{n_j}}(\mathbf{G}(\mathbf{z}), \mathbf{z}) \geq J_{1,\beta_{n_j}}(\hat{\mathbf{x}}_{\beta_{n_j}}, \hat{\mathbf{z}}_{\beta_{n_j}}) > J_{1,0}(\hat{\mathbf{x}}_{\beta_{n_j}}, \hat{\mathbf{z}}_{\beta_{n_j}}) + \frac{\beta_{n_j}}{2} \varepsilon^2$$

As a result, the nonnegative quantity $J_{1,0}(\hat{\mathbf{x}}_{\beta_{n_j}}, \hat{\mathbf{z}}_{\beta_{n_j}})$ goes to $-\infty$, which leads to a contradiction. Thus, one has $\hat{\mathbf{x}}_\infty = \mathbf{G}(\hat{\mathbf{z}}_\infty)$ for any limit point $(\hat{\mathbf{x}}_\infty, \hat{\mathbf{z}}_\infty)$ of $(\hat{\mathbf{x}}_{\beta_n}, \hat{\mathbf{z}}_{\beta_n})$. Assume that $J_{1,0}(\hat{\mathbf{x}}_\infty, \hat{\mathbf{z}}_\infty) > J_{1,0}(\mathbf{x}^*, \mathbf{z}^*)$. Since

$$J_{1,\beta_n}(\hat{\mathbf{x}}_{\beta_n}, \hat{\mathbf{z}}_{\beta_n}) \leq J_{1,\beta_n}(\mathbf{x}^*, \mathbf{z}^*) = J_{1,0}(\mathbf{x}^*, \mathbf{z}^*) < J_{1,0}(\hat{\mathbf{x}}_\infty, \hat{\mathbf{z}}_\infty) \quad \blacksquare$$

this leads to another contradiction.

The previous result motivates Algorithm 1.1.

Consider Algorithm 1.1 in the ideal case ($\text{maxiter}=\infty$) where the internal loop converges.

Proposition C.2 (Convergence of Algorithm 1.1). *Let $(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k)_k$ be a sequence generated by Algorithm 1.1 when $\text{maxiter}=\infty$. If $(\mathbf{z}_\infty^k)_k$ is bounded, then any limit point of $(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k)_k$ is in $\mathcal{C}$. Moreover, any limit point of $(\mathbf{z}_\infty^k)_k$ is a stationary point of*

$$(C.4) \quad f(\mathbf{z}) = F(\mathbf{G}(\mathbf{z}), \mathbf{y}) + \frac{1}{2} \|\mathbf{z}\|^2$$

Proof. Note that, for any k , $(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k)$ is a limit point of the sequence generated by the k -th subloop in Algorithm 1.1 if it does not converge. Let $(\beta_k)_k$ a sequence that converges to ∞ . Let $k \in \mathbb{N}$. We consider the sequence $(\mathbf{x}_n^k, \mathbf{z}_n^k)_n$ generated by

$$\forall n \in \mathbb{N}, \quad \mathbf{z}_{n+1}^k \in \arg \min_{\mathbf{z}} J_{1,\beta_k}(\mathbf{x}_n^k, \mathbf{z}) \text{ and } \mathbf{x}_{n+1}^k = \arg \min_{\mathbf{x}} J_{1,\beta_k}(\mathbf{x}, \mathbf{z}_{n+1}^k)$$

with $\mathbf{x}_0^k = \mathbf{x}_{\infty}^{k-1}$. Since J_{1,β_k} corresponds to a particular instance of J_1 , and since Algorithm 1.1 can be seen asymptotically as a particular instance of Algorithm 2.3, one can use all the results established in Proposition 2.1. In particular, the sequence $(\mathbf{x}_n^k, \mathbf{z}_n^k)_n$ admits a limit point $(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k)$ and we have

$$\frac{\partial J_{1,\beta_k}}{\partial \mathbf{x}}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) = \frac{\partial J_{1,0}}{\partial \mathbf{x}}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) + \beta_k(\mathbf{x}_{\infty}^k - \mathbf{G}(\mathbf{z}_{\infty}^k)) = 0$$

and

$$\frac{\partial J_{1,\beta_k}}{\partial \mathbf{z}}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) = \frac{\partial J_{1,0}}{\partial \mathbf{z}}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) + \beta_k(D\mathbf{G}(\mathbf{z}_{\infty}^k))^*(\mathbf{G}(\mathbf{z}_{\infty}^k) - \mathbf{x}_{\infty}^k) = 0$$

By convexity, $\mathbf{x}_{\infty}^k$ is the (unique) minimizer of $J_{1,\beta_k}(\cdot, \mathbf{z}_{\infty}^k)$.

Assume that the sequence $(\mathbf{z}_{\infty}^k)_k$ is bounded. By optimality, one has

$$\min J_{1,0} \leq J_{1,0}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) \leq J_{1,\beta_k}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) \leq J_{1,\beta_k}(\mathbf{G}(\mathbf{z}_{\infty}^k), \mathbf{z}_{\infty}^k) = J_{1,0}(\mathbf{G}(\mathbf{z}_{\infty}^k), \mathbf{z}_{\infty}^k)$$

Since $(J_{1,0}(\mathbf{G}(\mathbf{z}_{\infty}^k), \mathbf{z}_{\infty}^k))_k$ is bounded, so is $(J_{1,0}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k))_k$. By coercivity, the sequence $(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k)_k$ is also bounded. Then it admits a limit point denoted $(\hat{\mathbf{x}}_{\infty}, \hat{\mathbf{z}}_{\infty})$. Let $(\mathbf{x}_{\infty}^{k_j}, \mathbf{z}_{\infty}^{k_j})_j$ be a convergent subsequence of limit $(\hat{\mathbf{x}}_{\infty}, \hat{\mathbf{z}}_{\infty})$. Let us assume that $\hat{\mathbf{x}}_{\infty} \neq \mathbf{G}(\hat{\mathbf{z}}_{\infty})$. Then, there exists $a > 0$ and $j_0 \in \mathbb{N}$ such that

$$\forall j \geq j_0, \quad \|\mathbf{x}_{\infty}^{k_j} - \mathbf{G}(\mathbf{z}_{\infty}^{k_j})\|^2 > a$$

Hence, one has

$$J_{1,\beta_{k_j}}(\mathbf{x}_{\infty}^{k_j}, \mathbf{z}_{\infty}^{k_j}) \geq J_{1,0}(\mathbf{x}_{\infty}^{k_j}, \mathbf{z}_{\infty}^{k_j}) + \beta_{k_j} a \geq \min J_{1,0} + \beta_{k_j} a \xrightarrow{j \rightarrow +\infty} \infty$$

which leads to a contradiction. This proves that $\hat{\mathbf{x}}_{\infty} = \mathbf{G}(\hat{\mathbf{z}}_{\infty})$. Otherwise said, $(\mathbf{x}_{\infty}^k - \mathbf{G}(\mathbf{z}_{\infty}^k))_k$ goes to zero.

Since we have for any k

$$\frac{\partial J_{1,0}}{\partial \mathbf{x}}(\mathbf{x}_{\infty}^k, \mathbf{z}_{\infty}^k) + \beta_k(\mathbf{x}_{\infty}^k - \mathbf{G}(\mathbf{z}_{\infty}^k)) = 0$$

the continuity of $\frac{\partial J_{1,0}}{\partial \mathbf{x}}$ ensures that $\left(\frac{\partial J_{1,0}}{\partial \mathbf{x}}(\mathbf{x}_{\infty}^{k_j}, \mathbf{z}_{\infty}^{k_j})\right)_j$ converges; thus, so is $(\beta_{k_j}(\mathbf{x}_{\infty}^{k_j} - \mathbf{G}(\mathbf{z}_{\infty}^{k_j})))_j$. ■

Then there exists $\lambda^* \in \mathbb{R}^d$ such that

$$\frac{\partial J_{1,0}}{\partial \mathbf{x}}(\mathbf{x}_{\infty}^{k_j}, \mathbf{z}_{\infty}^{k_j}) = -\beta_{k_j}(\mathbf{x}_{\infty}^{k_j} - \mathbf{G}(\mathbf{z}_{\infty}^{k_j})) \xrightarrow{j \rightarrow +\infty} \lambda^* = \frac{\partial J_{1,0}}{\partial \mathbf{x}}(\hat{\mathbf{x}}_{\infty}, \hat{\mathbf{z}}_{\infty})$$

and

$$\frac{\partial J_{1,0}}{\partial \mathbf{z}}(\mathbf{x}_{\infty}^{k_j}, \mathbf{z}_{\infty}^{k_j}) = -\beta_{k_j}(D\mathbf{G}(\mathbf{z}_{\infty}^{k_j}))^*(\mathbf{G}(\mathbf{z}_{\infty}^{k_j}) - \mathbf{x}_{\infty}^{k_j}) \xrightarrow{j \rightarrow +\infty} -(D\mathbf{G}(\hat{\mathbf{z}}_{\infty}))^*(\lambda^*)$$

Note that $f(\mathbf{z}) = J_{1,0}(\mathbf{G}(\mathbf{z}), \mathbf{z})$. One can check that f is differentiable and that

$$\nabla f(\mathbf{z}) = (D\mathbf{G}(\mathbf{z}))^* \left(\frac{\partial J_{1,0}}{\partial \mathbf{x}}(\mathbf{G}(\mathbf{z}), \mathbf{z}) \right) + \frac{\partial J_{1,0}}{\partial \mathbf{z}}(\mathbf{G}(\mathbf{z}), \mathbf{z})$$

Hence, we have proved that

$$\nabla f(\hat{\mathbf{z}}_{\infty}) = 0$$

Conclusion: If $(\mathbf{z}_{\infty}^k)_k$ is bounded, any limit point of $(\mathbf{z}_{\infty}^k)_k$ is a stationary point of (C.4). ■

In general, we can only prove that the limit points of the sequences generated by Algorithm 1.1 are stationary points of C.4. However, if the growth of β is sufficiently slow, then we obtain the optimality of the limit points. Indeed, given that, in Algorithm 1.1, each subloop is an exact BCD scheme, one has for any n and any j

$$\forall \mathbf{z}, \quad J_{1,\beta_{k_j}}(\mathbf{x}_{n-1}, \mathbf{z}_n) \leq J_{1,\beta_{k_j}}(\mathbf{x}_{n-1}, \mathbf{z})$$

By considering the subsequence $(\mathbf{x}_{n_\ell}, \mathbf{z}_{n_\ell})$, which converges to $\mathbf{x}_\infty^{k_j}, \mathbf{z}_\infty^{k_j}$ (we recall that $\mathbf{x}_{n_\ell}$ and $\mathbf{x}_{n_\ell-1}$ have same limit), we can prove that

$$\forall \mathbf{z}, \quad J_{1,\beta_{k_j}}(\mathbf{x}_\infty^{k_j}, \mathbf{z}_\infty^{k_j}) \leq J_{1,\beta_{k_j}}(\mathbf{x}_\infty^{k_j}, \mathbf{z})$$

that is, $\mathbf{z}_\infty^{k_j}$ is a minimizer of $J_{1,\beta_{k_j}}(\mathbf{x}_\infty^{k_j}, \cdot)$. Hence, we have

$$\forall \mathbf{z}, \quad J_{1,0}(\mathbf{x}_\infty^{k_j}, \mathbf{z}_\infty^{k_j}) + \frac{\beta_{k_j}}{2} \|\mathbf{x}_\infty^{k_j} - \mathbf{G}(\mathbf{z}_\infty^{k_j})\|^2 \leq J_{1,0}(\mathbf{x}_\infty^{k_j}, \mathbf{z}) + \frac{\beta_{k_j}}{2} \|\mathbf{x}_\infty^{k_j} - \mathbf{G}(\mathbf{z})\|^2$$

Assume that $\beta_{k_j} \|\mathbf{x}_\infty^{k_j} - \hat{\mathbf{x}}_\infty\|^2 \xrightarrow{j \rightarrow +\infty} 0$. By letting j to ∞ , we get that, for any $\mathbf{z}$ such that $\mathbf{G}(\mathbf{z}) = \mathbf{x}_\infty^{k_j}$,

$$J_{1,0}(\hat{\mathbf{x}}_\infty, \hat{\mathbf{z}}_\infty) \leq J_{1,0}(\hat{\mathbf{x}}_\infty, \mathbf{z}) + \lim_{j \rightarrow +\infty} \frac{\beta_{k_j}}{2} \|\mathbf{x}_\infty^{k_j} - \mathbf{x}_\infty^{k_j}\|^2 = J_{1,0}(\hat{\mathbf{x}}_\infty, \mathbf{z})$$

that is, $\hat{\mathbf{z}}_\infty$ is a minimizer of $J_{1,0}(\hat{\mathbf{x}}_\infty, \cdot) + \chi_C(\hat{\mathbf{x}}_\infty, \cdot)$. By definition of f , this also means that $\hat{\mathbf{z}}_\infty$ is a minimizer of f . However, one has to note that the growth control for β depends on the convergence speed of $\mathbf{x}_\infty^k$, which cannot be estimated.

Algorithm 2.4 is a particular (truncated) case of Algorithm 1.1 with an adaptive choice of β that does not need to go to ∞ .

Proposition C.3 (Convergence of Algorithm 2.4).

Proof. Let us write the Lagrangian of the problem solved in Algorithm 2.4:

$$\forall \lambda \geq 0, \quad \mathcal{L}(\mathbf{x}, \mathbf{z}; \lambda) = J_{1,0}(\mathbf{x}, \mathbf{z}) + \lambda (\|\mathbf{x} - \mathbf{G}(\mathbf{z})\|^2 - \varepsilon)$$

KKT conditions ensure that any solution $(\mathbf{x}^*, \mathbf{z}^*)$ of the constrained problem is associated to at least one Lagrange multiplier $\lambda^* \geq 0$ such that

$$\frac{\partial \mathcal{L}}{\partial (\mathbf{x}, \mathbf{z})}(\mathbf{x}^*, \mathbf{z}^*; \lambda^*) = 0 = \begin{pmatrix} \frac{\partial J_{1,0}}{\partial \mathbf{x}}(\mathbf{x}^*, \mathbf{z}^*) + 2\lambda^* (\mathbf{x}^* - \mathbf{G}(\mathbf{z}^*)) \\ \frac{\partial J_{1,0}}{\partial \mathbf{z}}(\mathbf{x}^*, \mathbf{z}^*) + 2\lambda^* (D\mathbf{G}(\mathbf{z}^*))^*(\mathbf{x}^* - \mathbf{G}(\mathbf{z}^*)) \end{pmatrix}$$

According to the calculus above, this proves that $(\mathbf{x}^*, \mathbf{z}^*)$ is a stationary point of $J_{1,2\lambda^*}$. Note that, if $\lambda^* = 0$, then $(\mathbf{x}^*, \mathbf{z}^*)$ is a minimizer of $J_{1,0}$. Otherwise, one has $\|\mathbf{x}^* - \mathbf{G}(\mathbf{z}^*)\|^2 = \varepsilon$.

Hence, if we consider Algorithm 1.1 with the update rule for β_k as in Algorithm 2.4 and a stopping rule saying that the iterations stop as soon as, for any given k ,

$$\|\mathbf{x}_\infty^k - \mathbf{G}(\mathbf{z}_\infty^k)\|^2 \leq \varepsilon$$

there are two possible cases:

1. **case** $\lambda^* = 0$: then $(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k)$ is a solution of the constraint problem iff $\nabla J_{1,0}(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k) = 0$ (that is, $\mathbf{x}_\infty^k = \mathbf{G}(\mathbf{z}_\infty^k)$);
2. **case** $\lambda^* > 0$: unless $\|\mathbf{x}_\infty^k - \mathbf{G}(\mathbf{z}_\infty^k)\|^2$ exactly equals ε , $(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k)$ is **not** a solution of the constraint problem

However, in general, $(\mathbf{x}_\infty^k, \mathbf{z}_\infty^k)$ is a solution of the following constraint problem

$$\min_{\|\mathbf{x}_\infty^k - \mathbf{G}(\mathbf{z}_\infty^k)\|^2 \leq \tilde{\varepsilon}} J_{1,0}(\mathbf{x}, \mathbf{z})$$

with $\tilde{\varepsilon} = \|\mathbf{x}_\infty^k - \mathbf{G}(\mathbf{z}_\infty^k)\|^2 \leq \varepsilon$. Hence, if we stop the iterations when $\|\mathbf{x}_\infty^k - \mathbf{G}(\mathbf{z}_\infty^k)\|^2 \leq \varepsilon$, we will get a solution of

$$\min_{\|\mathbf{x}_\infty^k - \mathbf{G}(\mathbf{z}_\infty^k)\|^2 \leq \tilde{\varepsilon}} J_{1,0}(\mathbf{x}, \mathbf{z}), \quad \tilde{\varepsilon} \leq \varepsilon$$

which provides an error control as well. ■