

An Information-Theoretic View of Stochastic Localization

Ahmed El Alaoui and Andrea Montanari

Abstract—Given a probability measure μ over $\mathbb{R}^n$, it is often useful to approximate it by the convex combination of a small number of probability measures, such that each component is close to a product measure. Recently, Ronen Eldan used a stochastic localization argument to prove a general decomposition result of this type. In Eldan’s theorem, the ‘number of components’ is characterized by the entropy of the mixture, and ‘closeness to product’ is characterized by the covariance matrix of each component.

We present an elementary proof of Eldan’s theorem which makes use of an information theory (or estimation theory) interpretation. The proof is analogous to the one of an earlier decomposition result known as the ‘pinning lemma.’

Index Terms—Low complexity decompositions, the pinning lemma, stochastic localization.

I. MOTIVATION AND RESULT

LET μ be an arbitrary Borel probability measure on $\mathbb{R}^n$. A broadly useful approach to understanding μ is to decompose it into simpler components

$$\mu = \int_{\Theta} \mu_{\theta} \rho(d\theta) =: \mathbb{E}_{\theta} \mu_{\theta}. \quad (\text{I.1})$$

Here $(\Theta, \mathcal{F}_{\Theta}, \rho)$ is an abstract probability space, and, for each θ , μ_{θ} is a probability measure on $\mathbb{R}^n$. In this paper, the components μ_{θ} will be simple in the sense of being close to product measures.

Of course, the decomposition (I.1) is always possible: just take $\Theta = \mathbb{R}^n$, $\rho = \mu$, and $\mu_{\theta} = \delta_{\theta}$ (a point mass at θ) for $\theta \in \mathbb{R}^n$. This is a ‘maximum entropy’ decomposition (all the entropy of μ is pushed into ρ), and is not particularly useful. We will be instead interested in decompositions such that ρ has small entropy, and the μ_{θ} ’s are only approximately product measures.

Low-entropy decompositions have found applications in statistical mechanics [COP19], [Aus19], random graph theory [CD16], [Eld18], random constraint satisfaction problems [ACOGM20], high-dimensional statistics [COKPZ18], analysis of Markov chains [EKZ21] to name a few areas. A generic construction consists in letting μ_{θ} be the conditional distribution of $x \sim \mu$ given a small subset of the coordinates of x (see below for further discussion). This leads to the so-called ‘pinning lemma’ which was discovered independently in [Mon08], [RT12]. Recently, Ronen Eldan [Eld20] pointed out that the pinning lemma can be suboptimal and proved a general

decomposition result with better properties. Eldan’s proof follows the general approach of ‘stochastic localization’ which is in turn inspired by ideas in high-dimensional geometry [Eld16].

The purpose of this note is to present a new interpretation of Eldan’s theorem, leading to an elementary proof. The interpretation has an information-theoretic (or estimation-theoretic) nature. We consider a noisy communication channel that takes as input $x \sim \mu$ and outputs θ . The distribution μ_{θ} is then the conditional distribution of the channel input given its output θ . Note that this is the same interpretation as for the proof of the pinning lemma in [Mon08]. The main difference is that while in the pinning lemma the channel $x \rightarrow \theta$ is an erasure channel, here we will use a Gaussian channel.

Turning to our construction, let $Q \in \mathbb{R}^{n \times n}$ be a fixed positive semidefinite matrix, $x \sim \mu$ and define the noisy observation (channel output) y via

$$y = \sqrt{\tau} x + Q^{1/2} z, \quad (\text{I.2})$$

where $z \sim N(0, I_n)$, and τ is uniform in the interval $[1, 2]$. The random variables x , z and τ are independent. For any fixed $\tau = t$, this is a noisy Gaussian channel, with signal-to-noise ratio (SNR) t . Introducing a random τ is convenient for the proof, but the estimates we prove hold also (with possibly worse constants) for all $\tau \in [1, 2]$, except on a set of small measure.

We set $\theta = (y, \tau)$ and $\mu_{\theta}(\cdot) = \mu(\cdot | \theta)$. It is clear that $\mu = \mathbb{E}_{\theta} \mu_{\theta}$ so that the decomposition (I.1) holds.

Let ν be a background measure on $\mathbb{R}^n$ such that μ is absolutely continuous with respect to ν , i.e. $\mu \ll \nu$. Recall that the relative entropy of μ with respect to ν (Kullback-Leibler divergence) is defined by

$$D(\mu \| \nu) = \mathbb{E}_{\mu} \log \frac{d\mu}{d\nu}(x). \quad (\text{I.3})$$

It is easy to see that $\mu_{\theta} \ll \mu \ll \nu$ a.s., so that $D(\mu_{\theta} \| \nu)$ is well defined.

Theorem 1 ([Eld20]). *We have*

$$\mathbb{E}_{\theta} \text{Cov}(\mu_{\theta}) \preceq Q, \quad (\text{I.4})$$

$$0 \leq \mathbb{E}_{\theta} D(\mu_{\theta} \| \nu) - D(\mu \| \nu) \quad (\text{I.5})$$

$$\leq \frac{1}{2} \log \det(I_n + 2Q^{-1} \text{Cov}(\mu)), \quad (\text{I.6})$$

$$\mathbb{E}_{\theta} \{ \text{Cov}(\mu_{\theta}) Q^{-1} \text{Cov}(\mu_{\theta}) \} \preceq \text{Cov}(\mu). \quad (\text{I.7})$$

Remark I.1. Equations (I.4) and (I.7) bound the covariance of the component measure μ_{θ} . On the other hand, Eq. (I.6) controls the entropy of the variable θ , which is given by the difference between the entropy of the mixture and the

A. El Alaoui is in the Department of Statistics and Data Science, Cornell University, Ithaca, NY 14850, USA. Email: elalaoui@cornell.edu.

A. Montanari is in the Department of Electrical Engineering and the Department of Statistics, Stanford University, CA 94305, USA. Email: montanar@stanford.edu.

Manuscript received September 8, 2021; revised June 1, 2022.

average entropy of the components. More precisely, in information theoretic terms, we control the mutual information $\mathbb{E}_\theta D(\mu_\theta \| \nu) - D(\mu \| \nu) = I(\theta; \mathbf{x})$ (see proof for definitions).

Remark I.2. As mentioned above, the difference between the decomposition presented here and the pinning lemma of [Mon08] is in the choice of the noisy channel. Here we use the Gaussian channel (I.2) (neglecting for a moment the fact that τ is random). In contrast, in [Mon08], the channel is an erasure channel: $y_i = x_i$ with probability ε and $y_i = *$ (an erasure) otherwise, independently across coordinates.

Remark I.3. The expectation with respect to θ in Eqs. (I.4) to (I.7) is an expectation with respect to $(\tau, \mathbf{y})$: $\mathbb{E}_\theta(\cdot) = \mathbb{E}_\tau(\mathbb{E}_{\mathbf{y}|\tau}(\cdot|\tau))$. The expectation with respect to τ can be eliminated using Markov inequality. This implies that there exists a set $T \subseteq [1, 2]$ of Lebesgue measure $|T| \geq 1 - 3M^{-1}$ such that, for any $t \in T$, the inequalities (I.4), (I.6), (I.7) hold up to a factor M for $\tau = t$. For instance, the first inequality is replaced by $\mathbb{E}(\text{Cov}(\mu_\theta)|\tau = t) \preceq M\mathbf{Q}$. Note that the conditional expectation is simply the expectation with respect to $\mathbf{y} = \sqrt{t}\mathbf{x} + \mathbf{Q}^{1/2}\mathbf{z}$.

We finally notice that the above information-theoretic interpretation does not only apply to the final decomposition (I.1), but to the whole measure-valued stochastic process defined in [Eld20]¹. We explain this extension in Section III.

II. PROOF

A. Proof of Eq. (I.4)

We compare the error of the maximum likelihood estimator of $\mathbf{x}$ to that of the Bayes optimal estimator of $\mathbf{x}$ given $\theta = (\mathbf{y}, \tau)$:

$$\hat{\mathbf{x}}_{\text{ML}} = \tau^{-1/2} \mathbf{y} \quad \text{and} \quad \hat{\mathbf{x}}_{\text{Bayes}} = \mathbb{E}[\mathbf{x}|\mathbf{y}, \tau].$$

Let $\mathbf{R} \in \mathbb{R}^{n \times n}$ be a fixed PSD matrix, and define $\|\mathbf{v}\|_{\mathbf{R}}^2 = \langle \mathbf{v}, \mathbf{R}\mathbf{v} \rangle$. Since $\hat{\mathbf{x}}_{\text{Bayes}}$ minimizes the mean squared error $\mathbb{E}\{\|\mathbf{x} - \hat{\mathbf{x}}(\theta)\|_{\mathbf{R}}^2\}$ among all measurable estimators $\hat{\mathbf{x}}$, we have

$$\mathbb{E}\text{Tr}[\mathbf{R}(\mathbf{x} - \hat{\mathbf{x}}_{\text{Bayes}})(\mathbf{x} - \hat{\mathbf{x}}_{\text{Bayes}})^\top] \leq \mathbb{E}\text{Tr}[\mathbf{R}(\mathbf{x} - \hat{\mathbf{x}}_{\text{ML}})(\mathbf{x} - \hat{\mathbf{x}}_{\text{ML}})^\top]. \quad (\text{II.1})$$

The left-hand side is equal to $\mathbb{E}\text{Tr}(\mathbf{R}\text{Cov}(\mu_\theta))$, and the right-hand side is equal to

$$\mathbb{E}[\tau^{-1} \langle \mathbf{z}\mathbf{Q}^{1/2}, \mathbf{R}\mathbf{Q}^{1/2}\mathbf{z} \rangle] = \mathbb{E}[\tau^{-1}] \cdot \text{Tr}(\mathbf{R}\mathbf{Q}). \quad (\text{II.2})$$

Since the inequality holds for all $\mathbf{R} \succeq 0$, we conclude that

$$\mathbb{E}\text{Cov}(\mu_\theta) \preceq \mathbb{E}[\tau^{-1}] \cdot \mathbf{Q} \preceq \mathbf{Q}. \quad (\text{II.3})$$

B. Proof of (I.5) and (I.6)

This claim follows immediately using some basic inequalities from information theory [CT91], [DCT91].

¹After the completion and submission of this work, we learned that this extension has been independently discovered in a concurrent work of Klartag and Putterman [KP21].

The mutual information of two random variables X, Y is $I(X; Y) := D(\mu_{X,Y} \| \mu_X \times \mu_Y)$, where we denote by $\mu_{X,Y}, \dots$ the joint law of $X, Y, \dots$. If $\mu_X, \mu_{X|Y}(\cdot|y) \ll \nu_X$, we have

$$I(X; Y) := \mathbb{E}_y D(\mu_{X|Y}(\cdot|y) \| \nu_X) - D(\mu_X \| \nu_X). \quad (\text{II.4})$$

We consider $X = \mathbf{x}$, $Y = \mathbf{y}$, under their joint distribution given $\tau = t$. In other words $\mathbf{x} \sim \mu$ and $\mathbf{y}$ is given by Eq. (I.2) with $\tau = t$. We have $\mu_\theta = \mu_{X|Y}(\cdot|\mathbf{y})$. Using non-negativity of the mutual information, we have

$$0 \leq I(\mathbf{x}; \mathbf{y}) = \mathbb{E}_y D(\mu_\theta \| \nu) - D(\mu \| \nu). \quad (\text{II.5})$$

which yields the inequality Eq. (I.5). As for Eq. (I.6), inverting the role of X, Y in the definition of mutual information, and letting $\bar{\nu}$ be, for instance, the standard Gaussian measure we get

$$I(\mathbf{x}; \mathbf{y}) = \mathbb{E}_x D(\mu_{\mathbf{y}|\mathbf{x}}(\cdot|\mathbf{x}) \| \bar{\nu}) - D(\mu_{\mathbf{y}} \| \bar{\nu}). \quad (\text{II.6})$$

Recall the definition of differential entropy of a random variable X with density f with respect to the Lebesgue measure: $h(X) := -\int f(x) \log f(x) dx$, and $h(X|Y) = h(X, Y) - h(Y)$. We then have from the last display

$$I(\mathbf{x}; \mathbf{y}) = h(\mathbf{y}) - h(\mathbf{y}|\mathbf{x}). \quad (\text{II.7})$$

The differential entropy of an n -dimensional Gaussian vector $\mathbf{g}$ with covariance Σ is $h(\mathbf{g}) = (n/2) \log(2\pi e) + (1/2) \text{Tr} \log \Sigma$. Therefore,

$$h(\mathbf{y}|\mathbf{x}) = \frac{n}{2} \log(2\pi e) + \frac{1}{2} \text{Tr} \log \mathbf{Q}.$$

Moreover, the Gaussian distribution maximizes the differential entropy among all those distributions with the same covariance, hence

$$h(\mathbf{y}) \leq \frac{n}{2} \log(2\pi e) + \frac{1}{2} \text{Tr} \log \text{Cov}(\mu_{\mathbf{y}}). \quad (\text{II.8})$$

Since $\text{Cov}(\mu_{\mathbf{y}}) = t\text{Cov}(\mu) + \mathbf{Q}$, we obtain

$$I(\mathbf{x}; \mathbf{y}) \leq \frac{1}{2} \text{Tr} \log \text{Cov}(\mu_{\mathbf{y}}) - \frac{1}{2} \text{Tr} \log \mathbf{Q} \quad (\text{II.9})$$

$$= \frac{1}{2} \log \det (\mathbf{I}_n + t\mathbf{Q}^{-1} \text{Cov}(\mu)). \quad (\text{II.10})$$

The claim then follows since $t \leq 2$.

C. Proof of (I.7)

Fixing $\tau = t$ a non-random value, we write $\mathbf{y}_t$ for the corresponding output of channel (I.2), namely $\mathbf{y}_t = \sqrt{t}\mathbf{x} + \mathbf{Q}^{1/2}\mathbf{z}$. We denote by $\mu_{\mathbf{y}_t}(\text{d}\mathbf{x})$ for the conditional distribution of $\mathbf{x}$ given $\mathbf{y}_t$. In order to emphasize its dependence on t , we will write $\mu_{\mathbf{x}_0, \mathbf{z}, t} := \mu_{\mathbf{y}_t = \sqrt{t}\mathbf{x}_0 + \mathbf{Q}^{1/2}\mathbf{z}}$. Simplifying Bayes formula, we get

$$\mu_{\mathbf{x}_0, \mathbf{z}, t}(\text{d}\mathbf{x}) = \exp \left\{ -\frac{t}{2} \|\mathbf{x} - \mathbf{x}_0\|_{\mathbf{Q}^{-1}}^2 + \sqrt{t} \langle \mathbf{z}, \mathbf{Q}^{-1/2} \mathbf{x} \rangle \right\} \frac{\mu(\text{d}\mathbf{x})}{Z_{\mathbf{x}_0, \mathbf{z}, t}}. \quad (\text{II.11})$$

Here we use the notation $\|\mathbf{v}\|_{\mathbf{A}}^2 := \langle \mathbf{v}, \mathbf{A}\mathbf{v} \rangle$.

Throughout this proof, given a measure $\nu(d\mathbf{x})$ and function $\psi_1(\mathbf{x}), \psi_2(\mathbf{x})$, we use the shorthands

$$\nu(\psi(\mathbf{x})) := \int \psi(\mathbf{x})\nu(d\mathbf{x}) \quad \text{and,}$$

$$\nu(\psi_1(\mathbf{x}); \psi_2(\mathbf{x})) := \nu(\psi_1(\mathbf{x}) \cdot \psi_2(\mathbf{x})) - \nu(\psi_1(\mathbf{x}))\nu(\psi_2(\mathbf{x})).$$

We then have

$$\begin{aligned} \frac{d}{dt}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}) &= -\frac{1}{2}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}; \|\mathbf{x} - \mathbf{x}_0\|_{\mathbf{Q}^{-1}}^2) \\ &\quad + \frac{1}{2\sqrt{t}}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}; \langle \mathbf{z}, \mathbf{Q}^{-1/2}\mathbf{x} \rangle). \end{aligned} \quad (\text{II.12})$$

(Note that $\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x})$ is the mean of the vector $\mathbf{x}$.)

Given a matrix $\mathbf{R} \succeq 0$, we define the minimum mean square error

$$\begin{aligned} \text{mmse}(t) &:= \mathbb{E}\left\{\|\mathbf{x} - \mathbb{E}(\mathbf{x}|\mathbf{y}_t)\|_{\mathbf{R}}^2\right\} \\ &= \text{Tr}(\mathbf{R} \text{Cov}(\mu)) - \mathbb{E}_{\mathbf{x}_0, \mathbf{z}}\langle \mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}), \mathbf{R}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}) \rangle. \end{aligned} \quad (\text{II.13})$$

Differentiating, and using the formula above to differentiate $\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x})$, we get

$$\begin{aligned} \frac{d}{dt}\text{mmse}(t) &= -2 \cdot \mathbb{E}_{\mathbf{x}_0, \mathbf{z}}\langle \mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}), \mathbf{R} \frac{d}{dt}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}) \rangle \\ &=: A - B_1 + B_2, \end{aligned} \quad (\text{II.14})$$

where

$$\begin{aligned} A &:= \mathbb{E}_{\mathbf{x}_0, \mathbf{z}}\left\{\langle \mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}), \mathbf{R}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}; \|\mathbf{x} - \mathbf{x}_0\|_{\mathbf{Q}^{-1}}^2) \rangle\right\}, \\ B_1 &:= \frac{1}{\sqrt{t}}\mathbb{E}_{\mathbf{x}_0, \mathbf{z}}\left\{\langle \mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}), \mathbf{R}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}\langle \mathbf{z}, \mathbf{Q}^{-1/2}\mathbf{x} \rangle) \rangle\right\}, \\ B_2 &:= \frac{1}{\sqrt{t}}\mathbb{E}_{\mathbf{x}_0, \mathbf{z}}\left\{\langle \mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}), \mathbf{R}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\mathbf{x}) \rangle \right. \\ &\quad \left. \cdot \mu_{\mathbf{x}_0, \mathbf{z}, t}(\langle \mathbf{z}, \mathbf{Q}^{-1/2}\mathbf{x} \rangle) \right\}. \end{aligned}$$

We next use Gaussian integration by parts (Stein's Lemma) to simplify the terms B_1, B_2 . It is useful to note that

$$\frac{1}{\sqrt{t}}\nabla_{\mathbf{z}}\mu_{\mathbf{x}_0, \mathbf{z}, t}(\psi(\mathbf{x})) = \mu_{\mathbf{x}_0, \mathbf{z}, t}(\psi(\mathbf{x}); \mathbf{Q}^{-1/2}\mathbf{x}). \quad (\text{II.15})$$

Writing for simplicity $\mu_t := \mu_{\mathbf{x}_0, \mathbf{z}, t}$ and $\mathbb{E} = \mathbb{E}_{\mathbf{x}_0, \mathbf{z}}$, we get

$$\begin{aligned} B_1 &= \mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x}|\|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2)\} \\ &\quad - \mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x}\mathbf{x}^\top)\mathbf{Q}^{-1}\mu_t(\mathbf{x})\} \\ &\quad + \mathbb{E}\text{Tr}\{\mathbf{Q}^{-1}[\mu_t(\mathbf{x}\mathbf{x}^\top) - \mu_t(\mathbf{x})\mu_t(\mathbf{x}^\top)]\mathbf{R}\mu_t(\mathbf{x}\mathbf{x}^\top)\}, \end{aligned} \quad (\text{II.16})$$

$$\begin{aligned} B_2 &= \mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x})\mu_t(\|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2)\} \\ &\quad - \mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x})\mu_t(\mathbf{x})^\top \mathbf{Q}^{-1}\mu_t(\mathbf{x})\} \\ &\quad + 2\mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{Q}^{-1}[\mu_t(\mathbf{x}\mathbf{x}^\top) - \mu_t(\mathbf{x})\mu_t(\mathbf{x}^\top)]\mathbf{R}\mu_t(\mathbf{x})\}. \end{aligned} \quad (\text{II.17})$$

Finally, by the tower property of conditional expectation

$$\begin{aligned} A &= \mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x}; \|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2)\} \\ &\quad - 2\mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x}; \mathbf{x}^\top)\mathbf{Q}^{-1}\mathbf{x}_0\} \\ &= \mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}\mu_t(\mathbf{x}; \|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2)\} \\ &\quad - 2\mathbb{E}\{\mu_t(\mathbf{x})^\top \mathbf{R}[\mu_t(\mathbf{x}\mathbf{x}^\top) - \mu_t(\mathbf{x})\mu_t(\mathbf{x}^\top)]\mathbf{Q}^{-1}\mu_t(\mathbf{x})\}. \end{aligned} \quad (\text{II.18})$$

We next substitute Eqs. (II.16), (II.17), and (II.18) in Eq. (II.14) to get

$$\frac{d}{dt}\text{mmse}(t) = -\mathbb{E}\text{Tr}\{\text{Cov}(\mu_t)\mathbf{Q}^{-1}\text{Cov}(\mu_t)\mathbf{R}\}. \quad (\text{II.19})$$

We next integrate this identity for $t \in [1, 2]$, to get

$$\begin{aligned} \int_1^2 \mathbb{E}\text{Tr}\{\text{Cov}(\mu_t)\mathbf{Q}^{-1}\text{Cov}(\mu_t)\mathbf{R}\}dt &= \text{mmse}(1) - \text{mmse}(2) \\ &\leq \text{mmse}(0) = \text{Tr}(\text{Cov}(\mu)\mathbf{R}). \end{aligned}$$

Finally, the integral over t can be interpreted as an expectation over $\tau \sim \text{Unif}([1, 2])$, whence

$$\mathbb{E}\text{Tr}\{\text{Cov}(\mu_\theta)\mathbf{Q}^{-1}\text{Cov}(\mu_\theta)\mathbf{R}\} \leq \text{Tr}(\text{Cov}(\mu)\mathbf{R}).$$

Since this holds for any $\mathbf{R} \succeq 0$, we proved Eq. (I.7).

III. EXTENSION TO THE STOCHASTIC LOCALIZATION PROCESS

Another way of writing the output of the Gaussian channel is as follows: Let $\mathbf{x} \sim \mu$ and $(\mathbf{B}_t)_{t \geq 0}$ be a standard Brownian motion in $\mathbb{R}^n$. Further, let

$$\bar{\mathbf{y}}_t = t\mathbf{x} + \mathbf{Q}^{1/2}\mathbf{B}_t. \quad (\text{III.1})$$

Then it is clear that $\bar{\mathbf{y}}_t$ has the same law as the vector $\sqrt{t}\mathbf{y}$, so the conditional distribution of $\mathbf{x}$ given $\mathbf{y}$ has the same law as the probability measure

$$\mu_t(d\mathbf{x}) = \frac{1}{Z_t} \exp\left\{\langle \bar{\mathbf{y}}_t, \mathbf{x} \rangle_{\mathbf{Q}^{-1}} - \frac{t}{2}\|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2\right\}\mu(d\mathbf{x}). \quad (\text{III.2})$$

We will show that the measure-valued process $(\mu_t)_{t \geq 0}$ satisfies the stochastic localization SDE of Eldan [Eld20], as stated below.

Theorem 2. Write L_t for the likelihood ratio process of μ_t with respect to μ :

$$L_t(\mathbf{x}) := \frac{d\mu_t}{d\mu}(\mathbf{x}). \quad (\text{III.3})$$

Then there exist a Brownian motion $(\mathbf{W}_t)_{t \geq 0}$ adapted to the filtration generated by $(\bar{\mathbf{y}}_t)_{t \geq 0}$, such that for all $\mathbf{x} \in \mathbb{R}^n$ and $t \geq 0$, we have

$$dL_t(\mathbf{x}) = L_t(\mathbf{x})\langle \mathbf{x} - \mathbf{a}_t, \mathbf{Q}^{-1/2}d\mathbf{W}_t \rangle, \quad \text{and} \quad L_0(\mathbf{x}) = 1, \quad (\text{III.4})$$

where

$$\mathbf{a}_t = \mathbb{E}\{\mathbf{x}|\bar{\mathbf{y}}_t\} = \int \mathbf{x} \mu_t(d\mathbf{x}). \quad (\text{III.5})$$

Proof. We use the representation (III.2) to write

$$d \log L_t(\mathbf{x}) = \langle d\bar{\mathbf{y}}_t, \mathbf{x} \rangle_{\mathbf{Q}^{-1}} - \frac{1}{2}\|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2 dt - d \log Z_t. \quad (\text{III.6})$$

Let us write $h_t(\mathbf{x})$ for the expression appearing in the exponent in Eq. (III.2). By Itô's formula we have

$$\begin{aligned} dZ_t &= \int_{\mathbb{R}^n} \left(\langle \mathbf{x}, \mathbf{Q}^{-1}d\bar{\mathbf{y}}_t \rangle - \frac{1}{2}\|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2 dt \right) e^{h_t(\mathbf{x})} \mu(d\mathbf{x}) \\ &\quad + \frac{1}{2} \left(\int_{\mathbb{R}^n} \|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2 e^{h_t(\mathbf{x})} \mu(d\mathbf{x}) \right) dt \\ &= \left\langle \int_{\mathbb{R}^n} \mathbf{x} e^{h_t(\mathbf{x})} \mu(d\mathbf{x}), \mathbf{Q}^{-1}d\bar{\mathbf{y}}_t \right\rangle. \end{aligned} \quad (\text{III.7})$$

With the notation $\mathbf{a}_t = \mathbb{E}\{\mathbf{x}|\bar{\mathbf{y}}_t\} = \int \mathbf{x} \mu_t(d\mathbf{x})$ (and writing $[Z]_t$ for the quadratic variation process), we obtain

$$d \log Z_t = \frac{dZ_t}{Z_t} - \frac{1}{2} \frac{d[Z]_t}{Z_t^2} \quad (\text{III.8})$$

$$= \langle \mathbf{a}_t, d\bar{\mathbf{y}}_t \rangle_{\mathbf{Q}^{-1}} - \frac{1}{2} \|\mathbf{a}_t\|_{\mathbf{Q}^{-1}}^2 dt. \quad (\text{III.9})$$

Substituting this in Eq. (III.6) we have

$$\begin{aligned} d \log L_t(\mathbf{x}) &= \langle \mathbf{x} - \mathbf{a}_t, d\bar{\mathbf{y}}_t \rangle_{\mathbf{Q}^{-1}} - \frac{1}{2} \left(\|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2 - \|\mathbf{a}_t\|_{\mathbf{Q}^{-1}}^2 \right) dt \\ &= \langle \mathbf{x} - \mathbf{a}_t, d\bar{\mathbf{y}}_t - \mathbf{a}_t dt \rangle_{\mathbf{Q}^{-1}} - \frac{1}{2} \|\mathbf{x} - \mathbf{a}_t\|_{\mathbf{Q}^{-1}}^2 dt. \end{aligned} \quad (\text{III.10})$$

It remains to understand the law of the process $\bar{\mathbf{y}}_t - \int_0^t \mathbf{a}_s ds$. We let $\mathcal{F}_t = \sigma(\{\bar{\mathbf{y}}_s : 0 \leq s \leq t\})$ and $\mathbb{F} = (\mathcal{F}_t)_{t \geq 0}$. It is known that the process $(\mathbf{W}_t)_{t \geq 0}$ defined by

$$\mathbf{W}_t := \mathbf{Q}^{-1/2} \left(\bar{\mathbf{y}}_t - \int_0^t \mathbb{E}\{\mathbf{x}|\mathcal{F}_s\} ds \right), \quad (\text{III.11})$$

is an $\mathbb{F}$ -adapted Brownian motion. See for instance Theorem 7.12 in [LS77], or Theorem 5.13 in [LG16]. We conclude by arguing that $\mathbb{E}\{\mathbf{x}|\mathcal{F}_t\} = \mathbf{a}_t$.

Indeed, for a fixed $t \geq 0$, we let

$$p_{t,\mathbf{x}} := \mathcal{L}((\bar{\mathbf{y}}_s)_{s \in [0,t]} | \mathbf{x})$$

denote the law of the path $(\bar{\mathbf{y}}_s)_{s \in [0,T]}$ conditional on $\mathbf{x}$ as per Eq. (III.1), and let $q_t = p_{t,\mathbf{x}=0}$. We obtain by Girsanov's theorem—see also Theorem 7.1 and its Corollary in [LS77]—that

$$\frac{dp_{t,\mathbf{x}}}{dq_t}(\mathbf{y}) = \exp \left\{ \langle \mathbf{y}_t, \mathbf{x} \rangle_{\mathbf{Q}^{-1}} - \frac{t}{2} \|\mathbf{x}\|_{\mathbf{Q}^{-1}}^2 \right\}, \quad (\text{III.12})$$

for all $\mathbf{y} \in C([0,t], \mathbb{R}^n)$ (the space of continuous functions on $[0,t]$ endowed with the topology of uniform convergence). We observe that the above only depends on the path $\mathbf{y}$ through the endpoint $\mathbf{y}_t$. Therefore by the Bayes rule, the conditional distribution of $\mathbf{x}$ given $\mathcal{F}_t$ is equal to μ_t (III.2), and we have $\mathbb{E}\{\mathbf{x}|\mathcal{F}_t\} = \mathbb{E}\{\mathbf{x}|\mathbf{y}_t\} = \mathbf{a}_t$ as a consequence. Continuing from Eq. (III.10), we have

$$d \log L_t(\mathbf{x}) = \langle \mathbf{x} - \mathbf{a}_t, \mathbf{Q}^{-1/2} d\mathbf{W}_t \rangle - \frac{1}{2} \|\mathbf{x} - \mathbf{a}_t\|_{\mathbf{Q}^{-1}}^2 dt, \quad (\text{III.13})$$

and we obtain (III.4) by applying Itô's formula once more. $\square$

ACKNOWLEDGMENT

We thank Ronen Eldan for a stimulating discussion. This work was carried out while the authors were (virtually) participating in the Fall 2020 program on Probability, Geometry and Computation in High Dimensions at the Simons Institute for the Theory of Computing. A.E.A. was supported by the Simons Berkeley Richard M. Karp Research Fellowship. A.M. was partially supported by the NSF grant CCF-2006489 and the ONR grant N00014-18-1-2729.

REFERENCES

- [ACOGM20] Peter Ayre, Amin Coja-Oghlan, Pu Gao, and Noëla Müller, *The satisfiability threshold for random linear equations*, *Combinatorica* **40** (2020), no. 2, 179–235.
- [Aus19] Tim Austin, *The structure of low-complexity gibbs measures on product spaces*, *The Annals of Probability* **47** (2019), no. 6, 4002–4023.
- [CD16] Sourav Chatterjee and Amir Dembo, *Nonlinear large deviations*, *Advances in Mathematics* **299** (2016), 396–450.
- [COKPZ18] Amin Coja-Oghlan, Florent Krzakala, Will Perkins, and Lenka Zdeborová, *Information-theoretic thresholds from the cavity method*, *Advances in Mathematics* **333** (2018), 694–795.
- [COP19] Amin Coja-Oghlan and Will Perkins, *Bethe states of random factor graphs*, *Communications in Mathematical Physics* **366** (2019), no. 1, 173–201.
- [CT91] Thomas M. Cover and Joy A. Thomas, *Elements of Information Theory*, Wiley, 1991.
- [DCT91] Amir Dembo, Thomas M Cover, and Joy A Thomas, *Information theoretic inequalities*, *IEEE Transactions on Information theory* **37** (1991), no. 6, 1501–1518.
- [EKZ21] Ronen Eldan, Frederic Koehler, and Ofer Zeitouni, *A spectral condition for spectral gap: fast mixing in high-temperature ising models*, *Probability Theory and Related Fields* (2021), 1–17.
- [Eld16] Ronen Eldan, *Skorokhod embeddings via stochastic flows on the space of gaussian measures*, *Annales de l'Institut Henri Poincaré, Probabilités et Statistiques*, vol. 52, Institut Henri Poincaré, 2016, pp. 1259–1280.
- [Eld18] ———, *Gaussian-width gradient complexity, reverse log-sobolev inequalities and nonlinear large deviations*, *Geometric and Functional Analysis* **28** (2018), no. 6, 1548–1596.
- [Eld20] ———, *Taming correlations through entropy-efficient measure decompositions with applications to mean-field approximation*, *Probability Theory and Related Fields* **176** (2020), no. 3, 737–755.
- [KP21] Bo'az Klartag and Eli Putterman, *Spectral monotonicity under Gaussian convolution*, arXiv preprint arXiv:2107.09496, 2021.
- [LG16] Jean-François Le Gall, *Brownian motion, martingales, and stochastic calculus*, Springer, 2016.
- [LS77] Robert Shevilevich Liptser and Al'bert Nikolaevich Shiriaev, *Statistics of random processes: General theory*, vol. 394, Springer, 1977.
- [Mon08] Andrea Montanari, *Estimating random variables from random sparse observations*, *European Transactions on Telecommunications* **19** (2008), 385–403.
- [RT12] Prasad Raghavendra and Ning Tan, *Approximating csps with global cardinality constraints using sdh hierarchies*, *Proceedings of the twenty-third annual ACM-SIAM symposium on Discrete Algorithms*, SIAM, 2012, pp. 373–387.