

Turbo Reconstruction of Structured Sparse Signals

Philip Schniter

Dept. ECE, The Ohio State University, Columbus OH 43210, schniter@ece.osu.edu

Abstract—This paper considers the reconstruction of structured-sparse signals from noisy linear observations. In particular, the support of the signal coefficients is parameterized by hidden binary pattern, and a structured probabilistic prior (e.g., Markov random chain/field/tree) is assumed on the pattern. Exact inference is discussed and an approximate inference scheme, based on loopy belief propagation (BP), is proposed. The proposed scheme iterates between exploitation of the observation-structure and exploitation of the pattern-structure, and is closely related to noncoherent turbo equalization, as used in digital communication receivers. An algorithm that exploits the observation structure is then detailed based on approximate message passing ideas. The application of EXIT charts is discussed, and empirical phase transition plots are calculated for Markov-chain structured sparsity.¹

I. INTRODUCTION

The problem of reconstructing sparse signals lies at the heart of many engineering and scientific applications. Here, the main objective is to estimate the sparse signal $\mathbf{x} \in \mathbb{C}^N$ from the noisy linear measurements $\mathbf{y} \in \mathbb{C}^M$,

$$\mathbf{y} = \mathbf{A}\mathbf{x} + \mathbf{w}, \quad (1)$$

where $\mathbf{A} \in \mathbb{C}^{M \times N}$ is a known matrix and $\mathbf{w} \in \mathbb{C}^M$ is additive noise, often modeled as circular white Gaussian, i.e., $\mathbf{w} \sim \mathcal{CN}(\mathbf{0}, \sigma^2 \mathbf{I})$. By “sparse,” we mean that the signal has only a few (say K , where $K \ll N$) non-zero coefficients.

In many cases of interest, the system of equations in (1) is underdetermined, i.e., $M \ll N$, so that, even in the noiseless case, there is no unique inverse. However, when $\mathbf{x}$ is known to be *sparse*, it is possible to accurately reconstruct $\mathbf{x}$ from $\mathbf{y}$ if the columns of $\mathbf{A}$ are sufficiently incoherent. For various sparse reconstruction algorithms, including convex-optimization-based, greedy, and iterative thresholding algorithms, there exist elegant bounds on reconstruction error that hold when $\mathbf{A}$ satisfies a certain *restricted isometry property* (RIP). (See [1] for a recent comprehensive overview.)

In many applications, however, the signal $\mathbf{x}$ has structure beyond simple sparsity. For example, the wavelet transform coefficients of natural scenes are not only approximately sparse, but also exhibit *persistence across scales* [1], which manifests as correlation within the sparsity pattern. Many other forms of structure in the sparsity pattern are also possible, and so we desire a powerful and flexible approach to modeling and exploiting such structure.

One approach to modeling sparsity structure is through the deterministic *union of subspaces* (UoS) approach. (See [1] for a recent comprehensive overview.) There, the K -sparse signal

$\mathbf{x} \in \mathbb{C}^N$ is assumed to live in a UoS $\mathcal{M}_K = \cup_{m=1}^{m_K} \mathcal{X}_m$, where $\mathcal{X}_m$ is one of the $\binom{N}{K}$ canonical subspaces containing signals with support K . Two flavors of UoS that have garnered particular attention are *block-sparsity* and *tree-sparsity* [1]. For both, probabilistic guarantees of a model-based RIP have been derived and, from them, bounds on noisy reconstruction error [1]. The UoS approach has rather strong limitations, however. For example, it is not possible to specify whether it is *more* likely to find a signal component in one of the allowed subspaces $\mathcal{X}_m$ relative to another.

In this paper, we take a probabilistic approach to modeling sparsity structure, allowing the use of, e.g., *Markov chain* (MC), *Markov random field* (MRF), and *Markov tree* (MT) models [2]. Such models have been previously exploited for sparse reconstruction, but only to a limited extent. For example, [3] and [4] proposed Monte-Carlo-based [5] sparse reconstruction algorithms using MRF and MT models, respectively, and [6] and [7] proposed to iterate matching-pursuit with MAP pattern detection based on MRF and MT models, respectively. Monte-Carlo algorithms, while flexible, are typically regarded as computationally too expensive for many problems of interest. Matching-pursuit algorithms are typically much faster, but the schemes in [6], [7] are ad hoc.

We attack the problem of reconstructing structured-sparse signals through the framework of *belief propagation* (BP) [8]. While BP has been successfully used to recover *unstructured* sparse signals (e.g., [9], [10]), we believe that its application to structured sparse signals is novel. As we shall see, the BP framework suggests an *iterative* approach, where sparsity pattern beliefs are exchanged between two blocks, one exploiting observation structure and the other exploiting pattern structure. In this regard, our scheme resembles *turbo equalization* from digital communications [11], where bit beliefs are exchanged between a soft equalizer and a soft decoder. Our two blocks are themselves naturally implemented using BP, and we detail a particularly efficient algorithm based on the *approximate message passing* (AMP) framework recently proposed by Donoho, Maleki, and Montanari [10].

II. STRUCTURED SPARSITY MODEL

Our structured-sparse signal model uses hidden binary indicators $\{s_n\}_{n=1}^N$, where $s_n \in \{0, 1\}$. In particular, $s_n = 1$ indicates that the signal coefficient x_n is active (i.e., $x_n \neq 0$ w.p.1) while $s_n = 0$ indicates that x_n is inactive (i.e., $x_n = 0$ w.p.1). Assuming that the active signal coefficients are independently but non-identically distributed, we can write

$$p(x_n | s_n) = s_n q_n(x_n) + (1 - s_n) \delta(x_n), \quad (2)$$

¹This work was supported in part by the Office of Naval Research under grant N00014-07-1-0209.

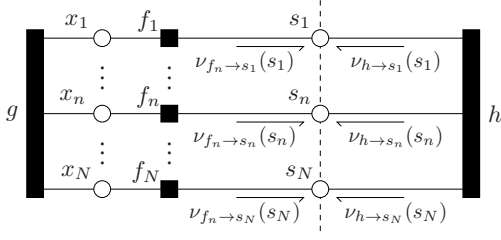


Fig. 1. Factor graph of posterior $p(\mathbf{x}, \mathbf{s} | \mathbf{y} = \mathbf{y}_0)$. The boxes represent constraint nodes, as defined in (4), and the circles represent variable nodes. The dashed line partitions the factor graph into two sub-graphs.

where $q_n(\cdot)$ denotes the pdf of x_n , when active, and $\delta(\cdot)$ denotes the Dirac delta. We refer to $\mathbf{s} = [s_1, s_2, \dots, s_N]^T \in \{0, 1\}^N$ as the *sparsity pattern*, and model structure in $\mathbf{s}$ through an assumed prior pmf $p(\mathbf{s})$. While we place no restrictions on $p(\mathbf{s})$, we note that Markov-chain/field/tree priors often lead to efficient inference algorithms [2].

III. TURBO INFERENCE

Our primary goal is estimating the structured-sparse signal $\mathbf{x}$ given the observations $\mathbf{y} = \mathbf{y}_0$ in model (1). In particular, we are interested in computing minimum mean-squared error (MMSE) estimates of $\{x_n\}$.

A. Exact inference

Our estimation task is facilitated by the following factorization of the posterior pdf shown by the *factor graph* in Fig. 1.

$$\begin{aligned} p(\mathbf{x}, \mathbf{s} | \mathbf{y} = \mathbf{y}_0) &\propto p(\mathbf{y} = \mathbf{y}_0 | \mathbf{x}, \mathbf{s}) p(\mathbf{x}, \mathbf{s}) \\ &= \underbrace{p(\mathbf{s})}_{\triangleq h(\mathbf{s})} \underbrace{p(\mathbf{y} = \mathbf{y}_0 | \mathbf{x})}_{\triangleq g(\mathbf{x})} \prod_{n=1}^N \underbrace{p(x_n | s_n)}_{\triangleq f_n(x_n, s_n)}. \end{aligned} \quad (3)$$

We use $\propto$ to denote equality after scaling to unit area.

The MMSE estimate of x_n is given by the mean of the marginal posterior $p(x_n | \mathbf{y} = \mathbf{y}_0)$, which can be written as

$$\begin{aligned} p(x_n | \mathbf{y} = \mathbf{y}_0) &= \sum_{\mathbf{s} \in \{0, 1\}^N} \int_{\mathbf{x}_{-n}} p(\mathbf{x}, \mathbf{s} | \mathbf{y} = \mathbf{y}_0) \\ &\propto \sum_{s_n=0}^1 f_n(x_n, s_n) p(s_n) \int_{\mathbf{x}_{-n}} g(\mathbf{x}) \prod_{q \neq n} \sum_{s_q=0}^1 f_q(x_q, s_q) \\ &\quad \times \sum_{\mathbf{s}_{-n, q} \in \{0, 1\}^{N-2}} p(\mathbf{s}_{-n} | s_n), \end{aligned} \quad (4)$$

where $\mathbf{z}_{-n}$ denotes vector $\mathbf{z}$ with the n^{th} element omitted, and $\mathbf{z}_{-n, q}$ denotes $\mathbf{z}$ with both the n^{th} and q^{th} elements omitted. Writing $p(\mathbf{s}_{-n} | s_n) = p(\mathbf{s}_{-n, q} | s_q, s_n) p(s_q | s_n)$, the last summation in (6) reduces to $p(s_q | s_n)$, giving

$$p(x_n | \mathbf{y} = \mathbf{y}_0) \propto \nu_{f_n \rightarrow x_n}^{\text{exact}}(x_n) \nu_{g \rightarrow x_n}^{\text{exact}}(x_n) \quad (5)$$

$$\nu_{f_n \rightarrow x_n}^{\text{exact}}(x_n) \triangleq \sum_{s_n=0}^1 f_n(x_n, s_n) p(s_n) \quad (6)$$

$$\nu_{g \rightarrow x_n}^{\text{exact}}(x_n) \triangleq \int_{\mathbf{x}_{-n}} g(\mathbf{x}) \prod_{q \neq n} \sum_{s_q=0}^1 f_q(x_q, s_q) p(s_q | s_n). \quad (7)$$

The notation $\nu_{A \rightarrow B}^{\text{exact}}(\cdot)$ will be explained in the sequel.

B. Approximate inference

Whereas exact posterior calculation via (7)-(9) is computationally prohibitive for typical problem sizes, approximate calculation can be efficiently accomplished using message passing [8], i.e., belief propagation (BP), on the factor graph in Fig. 1. In the sequel, we use $\nu_{A \rightarrow B}(\cdot)$ to denote a message passed from some node A to some adjacent node B in the factor graph. Our messages will either be pdfs on the real line or binary pmfs; which one will be clear from the context. BP proceeds according to the following two rules [8]: i) the message emitted by a variable node along a given edge equals the product of the messages coming into that variable node along the other edges, and ii) the message emitted by a function node along a given edge equals the integral of the product of the constraint function (associated with that node) and all messages coming into that function node along the other edges. (Here, “equals” holds after appropriate scaling.)

Using the framework of BP, the functions $\nu_{f_n \rightarrow x_n}^{\text{exact}}(\cdot)$ and $\nu_{g \rightarrow x_n}^{\text{exact}}(\cdot)$ from Section III-A can be approximated (up to a scaling factor) by steady-state versions of the messages

$$\nu_{f_n \rightarrow x_n}^{(t)}(x_n) \propto \sum_{s_n=0}^1 f_n(x_n, s_n) \nu_{s_n \rightarrow f_n}^{(t)}(s_n) \quad (8)$$

$$\nu_{g \rightarrow x_n}^{(t)}(x_n) \propto \int_{\mathbf{x}_{-n}} g(\mathbf{x}) \prod_{q \neq n} \sum_{s_q=0}^1 f_q(x_q, s_q) \nu_{s_q \rightarrow f_q}^{(t)}(s_q), \quad (9)$$

$$\underbrace{\nu_{f_q \rightarrow x_q}^{(t)}(x_n)}_{= \nu_{x_q \rightarrow g}^{(t)}(x_n)}$$

which depend on the other messages

$$\nu_{s_n \rightarrow f_n}^{(t)}(s_n) = \nu_{h \rightarrow s_n}^{(t)}(s_n) \quad (10)$$

$$\propto \sum_{\mathbf{s}_{-n} \in \{0, 1\}^{N-1}} h(\mathbf{s}) \prod_{q \neq n} \underbrace{\nu_{s_q \rightarrow h}^{(t-1)}(s_q)}_{= \nu_{f_q \rightarrow s_q}^{(t-1)}(s_q)} \quad (11)$$

$$\nu_{f_n \rightarrow s_n}^{(t)}(s_n) \propto \int_{x_n} f_n(x_n, s_n) \underbrace{\nu_{x_n \rightarrow f_n}^{(t)}(x_n)}_{= \nu_{g \rightarrow f_n}^{(t)}(x_n)}. \quad (12)$$

We use the superscript-(t) to denote turbo iteration. These messages can then be combined for marginal inference:

$$\hat{p}^{(t)}(x_n | \mathbf{y} = \mathbf{y}_0) \propto \nu_{f_n \rightarrow x_n}^{(t)}(x_n) \nu_{g \rightarrow x_n}^{(t)}(x_n) \quad (13)$$

$$\hat{p}^{(t)}(s_n | \mathbf{y} = \mathbf{y}_0) \propto \nu_{f_n \rightarrow s_n}^{(t)}(s_n) \nu_{h \rightarrow s_n}^{(t)}(s_n), \quad (14)$$

where $\hat{p}^{(t)}$ denotes the iteration- t approximation to the pdf.

Because the factor graph in Fig. 1 contains many loops, exact inference is known to be NP-hard [12], and thus BP can only be claimed as an approximation. However, loopy BP has demonstrated very accurate results in similar settings (e.g., in LDPC and turbo decoding [13] and in inference on MRFs [14]). Furthermore, in the large-system limit (i.e., $M, N \rightarrow \infty$ with M/N fixed), BP has recently been shown to enjoy various optimality properties [10], [15].

With practical implementation in mind, we now partition our factor graph into the two sub-graphs separated by the dashed line in Fig. 1. The messages $\{\nu_{f_n \rightarrow s_n}^{(t)}(\cdot)\}_{n=1}^N$ form the outputs of the left sub-graph and the inputs to the right one, while the messages $\{\nu_{h \rightarrow s_n}^{(t)}(\cdot)\}_{n=1}^N$ form the outputs of the right sub-graph and the inputs to the left one. From this, we can interpret the BP scheme as iterating between two blocks, one which performs inference on the left sub-graph (which models structure in the observation) and the other which performs inference on the right sub-graph (which models structure in the sparsity pattern), with message-passing between blocks.

C. Relation to noncoherent turbo equalization

The iterative approach described in Section III-B mimics the *turbo equalization* [11] procedure that has become popular for the reception of digital communication signals, and in particular, *noncoherent turbo equalization* (e.g., [16]). There, the goal is to infer a sequence of information bits $\mathbf{b}$ that are transformed into coded bits $\mathbf{s}$, linearly modulated, and transmitted over a channel, yielding the noisy observations $\mathbf{y} = \mathbf{H}\mathbf{s} + \mathbf{w}$. In the noncoherent case, $\mathbf{H}$ depends on an *unknown* channel fading realization. Because optimal joint channel-estimation/equalization/decoding is computationally overwhelming, the bit inference task is often split into two sub-tasks, *soft noncoherent equalization* and *soft decoding*, which are then iterated. The equalizer calculates soft estimates of the coded bits $\mathbf{s}$ from $\mathbf{y}$ using knowledge of the channel structure and prior beliefs on $\mathbf{s}$ (as pilots or from the decoder). It then passes these $\mathbf{s}$ -estimates to the decoder, which treats them as priors when exploiting its knowledge of the code structure to refine the $\mathbf{s}$ -estimates. The decoder's estimates are then passed back to the equalizer, and so on, until they agree. To avoid self-reinforcement, the two blocks pass *extrinsic* information (i.e., a bit estimate may not employ the most recently assumed prior for that bit). When used in conjunction with powerful codes, turbo equalization facilitates practical communication at near-Shannon capacity [11].

To see the similarities between noncoherent turbo equalization and our BP-based approach to structured sparse signal recovery, we write the sparse coefficients as $x_n = \theta_n s_n$, where θ_n has prior pdf $q_n(\cdot)$, so that $\mathbf{x} = \mathcal{D}(\boldsymbol{\theta})\mathbf{s}$, where $\mathcal{D}(\boldsymbol{\theta})$ is the diagonal matrix constructed from $\boldsymbol{\theta} = [\theta_1, \theta_2, \dots, \theta_N]^T$. In this case, the sparse observations (1) can be written as $\mathbf{y} = \mathbf{H}\mathbf{s} + \mathbf{w}$ with $\mathbf{H} = \mathbf{A}\mathcal{D}(\boldsymbol{\theta})$. Thus, the structured sparsity pattern $\mathbf{s}$ is analogous to “coded bits,” the active-coefficient vector $\boldsymbol{\theta}$ is analogous to the “unknown fading realization,” and $\mathbf{A}$ contributes known structure to the unknown “channel matrix” $\mathbf{H}$. Inference on the left sub-graph of Fig. 1 is analogous to “soft noncoherent equalization,” while that on the right sub-graph corresponds to “soft decoding.”

Taking cues from turbo equalization, we will henceforth refer to inference on the left sub-graph of Fig. 1 as “*sparsity pattern equalization*” (SPE) and inference on the right sub-graph as “*sparsity pattern decoding*” (SPD). We now formally decouple these subtasks and represent each of them using a separate factor graph, as in Fig. 2. For this, we define two

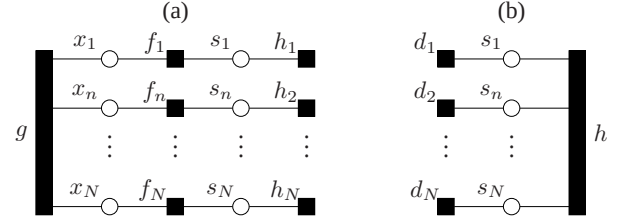


Fig. 2. Decoupling of partitioned factor graph from Fig. 1 into (a) sparsity pattern equalization and (b) sparsity pattern decoding.

additional t^{th} turbo-iteration constraint functions,

$$h_n^{(t)}(s_n) \triangleq \nu_{h \rightarrow s_n}^{(t)}(s_n) \quad (17)$$

$$d_n^{(t)}(s_n) \triangleq \nu_{f_n \rightarrow s_n}^{(t-1)}(s_n), \quad (18)$$

recalling (13)-(14), initialized using $\nu_{f_n \rightarrow s_n}^{(0)}(s_n) = 0.5$.

For t^{th} -iteration SPE, (16) implies an output message $\nu_{f_n \rightarrow s_n}^{(t)}(s_n) \propto \hat{p}^{(t)}(s_n | \mathbf{y} = \mathbf{y}_0) / h_n^{(t)}(s_n)$, which equals the (approximate) likelihood of s_n , since $h_n^{(t)}(\cdot)$ acts as SPE's prior on s_n . As such, the SPE output messages are extrinsic. Similar reasoning applies to SPD outputs as well.

IV. SPARSITY PATTERN EQUALIZATION

As previously discussed, SPE accepts independent but non-identical probabilities on the indicators $\{s_n\}$ and returns likelihoods on the same. One can imagine a number of ways to accomplish this task, e.g., Monte Carlo methods [5], “soft” Bayesian matching pursuit [17], expectation maximization, or BP. Below we outline a BP-based technique that follows the “approximate message passing” (AMP) framework recently proposed by Donoho, Maleki, and Montanari [10], [18]. Since we focus on a single turbo iteration t , we suppress the superscript- (t) notation on messages in this section.

For BP-based SPE, we expand the g node in Fig. 2(a), yielding the loopy factor graph in Fig. 3, with constraints

$$g_m(\mathbf{x}) \triangleq \mathcal{CN}(\mathbf{y}_m; \mathbf{a}_m^H \mathbf{x}, \sigma^2), \quad (19)$$

where $\mathbf{a}_m^H$ denotes the m^{th} row of $\mathbf{A}$. Noting that SPE will require several iterations of message passing between nodes $\{g_m\}$ and $\{x_n\}$, we will henceforth use $\nu_{x_n \rightarrow g_m}^i$ and $\nu_{g_m \rightarrow x_n}^i$ to denote the SPE-iteration- i messages. In addition, we will assume Gaussian active-coefficients, i.e.,

$$q_n(x_n) = \mathcal{CN}(x_n; 0, \sigma_n^2), \quad (20)$$

though other distributions could be handled using similar techniques. Since (20) allows non-uniform coefficient variance, we can assume w.l.o.g that all columns of $\mathbf{A}$ have unit ℓ_2 -norm (after absorbing any variations into $\{\sigma_n^2\}$). In the sequel, we use λ_n to abbreviate $h_n(1)$, the prior probability of $s_n = 1$ assumed by SPE. Thus, the coefficient prior implicit to our incarnation of AMP is Bernoulli-Gaussian, with the form

$$\nu_{f_n \rightarrow x_n}(x_n) = \lambda_n \mathcal{CN}(x_n; 0, \sigma_n^2) + (1 - \lambda_n) \delta(x_n). \quad (21)$$

In the sequel, we will make use of the abbreviations

$$\alpha_n(c) \triangleq \frac{\sigma_n^2}{c + \sigma_n^2}, \quad \beta_n(c) \triangleq \frac{1 - \lambda_n}{\lambda_n} \frac{c + \sigma_n^2}{c}, \quad \zeta_n(c) \triangleq \frac{\sigma_n^2}{c(c + \sigma_n^2)}.$$

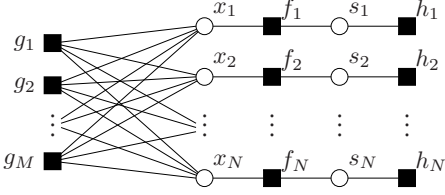


Fig. 3. Factor graph for BP-based implementation of SPE.

A. BP approximation via the large-system limit

Due to the mixture prior (21), exact calculation of $\nu_{g_m \rightarrow x_n}^i(x_n)$ would involve the integration of 2^{N-1} terms, which is clearly impractical. However, in the *large system limit* (i.e., $M, N \rightarrow \infty$ with M/N fixed), the central limit theorem motivates the treatment of $\nu_{g_m \rightarrow x_n}^i(x_n)$ as Gaussian [18]. In this case, it is sufficient to parameterize the inputs to g_m via

$$\mu_{nm}^i \triangleq \int x_n \nu_{x_n \rightarrow g_m}^i(x_n), \quad v_{nm}^i \triangleq \int (x_n - \mu_{nm}^i)^2 \nu_{x_n \rightarrow g_m}^i(x_n), \quad (22)$$

which yields outputs from g_m that take the form

$$\nu_{g_m \rightarrow x_n}^i(x_n) \propto \mathcal{CN}(A_{mn}x_n; z_{mn}^i, c_{mn}^i) \quad (23)$$

$$z_{mn}^i \triangleq y_m - \sum_{q \neq n} A_{mq} \mu_{qm}^i \quad (24)$$

$$c_{mn}^i \triangleq \sigma^2 + \sum_{q \neq n} |A_{mq}|^2 v_{qm}^i \quad (25)$$

as can be shown using the fact that

$$\prod_q \mathcal{CN}(x; \mu_q, v_q) \propto \mathcal{CN}\left(x; \frac{\sum_q \mu_q / v_q}{\sum_q v_q^{-1}}, \frac{1}{\sum_q v_q^{-1}}\right). \quad (26)$$

From (22), we see that μ_{nm}^{i+1} and v_{nm}^{i+1} are then determined by the mean and variance, respectively, of the pdf

$$\nu_{x_n \rightarrow g_m}^{i+1}(x_n) \propto \nu_{f_n \rightarrow x_n}(x_n) \prod_{l \neq m} \nu_{g_l \rightarrow x_n}^i(x_n). \quad (27)$$

Using (26), the product term in (27) reduces to

$$\mathcal{CN}\left(x_n; \frac{\sum_{l \neq m} A_{ln}^* z_{ln}^i / c_{ln}^i}{\sum_{l \neq m} |A_{ln}|^2 / c_{ln}^i}, \frac{1}{\sum_{l \neq m} |A_{ln}|^2 / c_{ln}^i}\right), \quad (28)$$

and so, under the large-system-limit approximations

$$c_{ln}^i \approx c_n^i \triangleq \frac{1}{M} \sum_{m=1}^M c_{mn}^i \quad (29)$$

and $\sum_{l \neq m} |A_{ln}|^2 \approx \sum_{l=1}^M |A_{ln}|^2 = 1$, (27) simplifies to

$$\nu_{x_n \rightarrow g_m}^{i+1}(x_n) \propto (\lambda_n \mathcal{CN}(x_n; 0, \sigma_n^2) + (1 - \lambda_n) \delta(x_n)) \times \mathcal{CN}(x_n; \sum_{l \neq m} A_{ln}^* z_{ln}^i, c_n^i). \quad (30)$$

Applying (26) to (30), we find, after some algebra, that

$$\mu_{nm}^{i+1} = \alpha_n(c_n^i) \theta_{nm}^i / (1 + \gamma_{nm}^i) \quad (31)$$

$$v_{nm}^{i+1} = \gamma_{nm}^i |\mu_{nm}^{i+1}|^2 + \mu_{nm}^{i+1} c_n^i / \theta_{nm}^i \quad (32)$$

$$\theta_{nm}^i \triangleq \sum_{l \neq m} A_{ln}^* z_{ln}^i \quad (33)$$

$$\gamma_{nm}^i \triangleq \beta_n(c_n^i) \exp(-\zeta_n(c_n^i) |\theta_{nm}^i|^2). \quad (34)$$

For the first turbo iteration, we desire that $\nu_{g_m \rightarrow x_n}^0(x_n) \propto 1$, and so we set $c_n^0 \gg \sigma_n^2$ and $z_{mn}^0 = y_m$ for all n . For second

and later turbo iterations, we set $\{c_n^0\}$ and $\{z_{mn}^0\}$ equal to their final values from the previous turbo iteration.

The i^{th} SPE iteration yields the x_n -posterior approximation

$$\hat{p}^i(x_n | \mathbf{y} = \mathbf{y}_0) \propto \nu_{f_n \rightarrow x_n}(x_n) \prod_{l=1}^M \nu_{g_l \rightarrow x_n}^{i-1}(x_n). \quad (35)$$

The mean and variance of (35) constitute the MMSE estimate of x_n and its MSE. Noting that (35) differs from (27) only in the inclusion of the m^{th} product term, analysis similar to (28)-(34) yields the following iteration- $(i+1)$ MMSE quantities:

$$\mu_n^{i+1} \triangleq \hat{\mathbb{E}}^i\{x_n | \mathbf{y} = \mathbf{y}_0\} = \alpha_n(c_n^i) \theta_n^i / (1 + \gamma_n^i) \quad (36)$$

$$v_n^{i+1} \triangleq \widehat{\text{var}}^i\{x_n | \mathbf{y} = \mathbf{y}_0\} = \gamma_n^i |\mu_n^{i+1}|^2 + \mu_n^{i+1} c_n^i / \theta_n^i \quad (37)$$

$$\theta_n^i \triangleq \sum_{l=1}^M A_{ln}^* z_{ln}^i \quad (38)$$

$$\gamma_n^i \triangleq \beta_n(c_n^i) \exp(-\zeta_n(c_n^i) |\theta_n^i|^2). \quad (39)$$

The i^{th} SPE iteration approximation to the s_n -posterior is

$$\hat{p}^i(s_n | \mathbf{y} = \mathbf{y}_0) \propto \nu_{f_n \rightarrow s_n}^{i-1}(s_n) \nu_{h_n \rightarrow s_n}(s_n), \quad (40)$$

$$\text{where } \nu_{f_n \rightarrow s_n}^i(s_n) \propto \int_{x_n} f_n(x_n, s_n) \prod_{l=1}^M \nu_{g_l \rightarrow x_n}^i(x_n). \quad (41)$$

From $f_n(x_n, s_n) = s_n \mathcal{CN}(x_n; 0, \sigma_n^2) + (1 - s_n) \delta(x_n)$, it is easy to show that the output log-likelihood ratio (LLR) is

$$L_n^i \triangleq \ln \frac{\nu_{f_n \rightarrow s_n}^i(1)}{\nu_{f_n \rightarrow s_n}^i(0)} = \ln \frac{c_n^i}{c_n^i + \sigma_n^2} + \zeta_n(c_n^i) |\theta_n^i|^2. \quad (42)$$

B. Approximate message passing

The approximate BP algorithm outlined in Section IV-A updates $\mathcal{O}(NM)$ variables per iteration: $\{z_{mn}^i\}$, $\{\mu_{nm}^i\}$, $\{v_{nm}^i\}$, and $\{c_n^i\}$ for $m = 1 \dots M$ and $n = 1 \dots N$. When N and M are large, the resulting complexity may be undesirably high, motivating us to find a simpler scheme.

Recently, Donoho, Maleki and Montanari proposed a family of so-called *first-order approximate message passing* (AMP) algorithms [10], [18] that greatly simplify BP algorithms of the form outlined in Section IV-A by tracking only $\mathcal{O}(N)$ variables. In particular, [18] describes how the AMP approach can be applied to generic independent (but not necessarily identical) real-coefficient priors. Paraphrasing [18, Sec. V], AMP initializes $\mu_n^0 = 0$, $z_m^0 = y_m$, and $c^0 \gg \sigma_n^2$ for all n and m , and then iterates (43)-(47) for $i = 0, 1, 2, \dots$

$$\theta_n^i = \sum_{m=1}^M A_{mn}^* z_m^i + \mu_n^i \quad (43)$$

$$\mu_n^{i+1} = F_n(\theta_n^i; c^i) \quad (44)$$

$$v_n^{i+1} = G_n(\theta_n^i; c^i) \quad (45)$$

$$c^{i+1} = \sigma^2 + \frac{1}{M} \sum_{n=1}^N v_n^{i+1} \quad (46)$$

$$z_m^{i+1} = y_m - \sum_{n=1}^N A_{mn} \mu_n^{i+1} + \frac{z_m^i}{M} \sum_{n=1}^N F_n'(\theta_n^i; c^i). \quad (47)$$

Above, $F_n(\cdot; \cdot)$, $G_n(\cdot; \cdot)$, and $F_n'(\cdot; \cdot)$ are nonlinear functions that depend on the coefficient prior. (See [18] for definitions.)

Using analyses similar to those in Section IV-A, one can show that, for the Bernoulli-Gaussian prior (21), these nonlin-

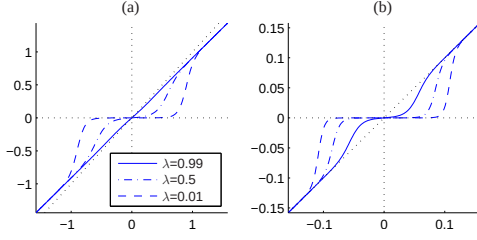


Fig. 4. The Bernoulli-Gaussian soft thresholding function $F_n(\theta; c)$ versus θ for several values of λ_n and for (a) $c = 0.1$ and (b) $c = 0.001$.

ear functions take the following form:

$$F_n(\theta; c) = \frac{\alpha_n(c)}{1 + \beta_n(c)e^{-\zeta_n(c)|\theta|^2}} \theta \quad (48)$$

$$G_n(\theta; c) = \beta_n(c)e^{-\zeta_n(c)|\theta|^2} |F_n(\theta; c)|^2 + \frac{c}{\theta} F_n(\theta; c) \quad (49)$$

$$F'_n(\theta; c) = \frac{\alpha_n(c)}{1 + \beta_n(c)e^{-\zeta_n(c)|\theta|^2}} \times \left[1 + \frac{\zeta_n(c)|\theta|^2}{1 + (\beta_n(c)e^{-\zeta_n(c)|\theta|^2} - 1)} \right]. \quad (50)$$

The function $F_n(\cdot; \cdot)$ plays the role of a soft threshold in (44), suppressing small values in the projected residual θ_n^i , and passing large ones. (See Fig. 4.) Here “small” and “large” are determined by the values of λ_n and c^i , where the latter has been observed to decrease with SPE iteration i .

AMP (43)-(50) follows from the approximate BP algorithm of Section IV-A via the intuition that, in the large system limit, both the n dependence of z_{mn}^i in (24) and the m dependence of θ_{nm}^i in (33) are weak. However, it is not sufficient to simply ignore these dependencies. Comparing (43) to (33), we see that the extra term $A_{mn}^* z_m^i$ is offset by μ_n^i , and comparing (47) to (24), we see that the extra term $A_{mn} \mu_n^{i+1}$ is offset by the last term in (47), which is analogous to the Onsager reaction term from statistical physics [10]. As is evident from above, the complexity of AMP is dominated by the two matrix multiplies (43) and (47) which, in some applications, can be efficiently implemented using a fast algorithm like the FFT.

V. SPARSITY PATTERN DECODING

The implementation of SPD depends strongly on the structure of $p(s)$. When $p(s)$ has a Markov structure, the inference task is a well-studied one that can be efficiently implemented using message passing algorithms [8], [13], [14], in some cases (e.g., with Markov chains and trees) optimally [2].

VI. EXIT CHARTS

As discussed earlier, BP-based implementations of SPE and of SPD are expected to work very well on their own. The question remains: How well do SPE and SPD work as a pair? In the turbo literature, *extrinsic information transfer (EXIT) charts* [19] are used to predict the extent to which two soft-input/output blocks can help each other in decoding. The same ideas can be used to predict the extent to which SPE and SPD can help each other learn the sparsity pattern.

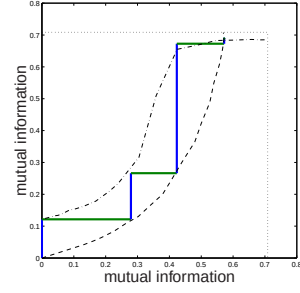


Fig. 5. EXIT chart, for $N = 512$, $M = 128$, $\lambda = 0.2$, $\gamma = 0.2$, and $\text{SNR} = 30$ dB, showing SPE EXIT curve (top), SPD EXIT curve (bottom), turbo trajectory (middle), and i.i.d entropy rate (dotted).

The *SPE EXIT curve* plots the mutual information between output LLRs $\{L_n^i\}$ and true pattern $\{s_n\}$, versus the mutual information between soft inputs $\{\ln \frac{\lambda_n}{1-\lambda_n}\}$ and true pattern $\{s_n\}$. Similarly, the *SPD EXIT curve* plots SPD’s output mutual information versus its input mutual information. From these two input-output maps, one can predict whether the information about $\{s_n\}$ can be increased through another turbo iteration, or whether learning has saturated. Sparsity-pattern EXIT charts differ from conventional EXIT charts in one small detail: the mutual information is plotted on the interval $[0, I_{\max}]$, where $I_{\max}$ is generally less than 1. This is because the entropy rate of i.i.d $\{s_n\}$ is < 1 when $\{s_n\}$ is non-equiprobable, as occurs in the sparse problem setting.

Figure 5 shows an example EXIT chart. (The simulation setup will be described in Section VII.) The turbo trajectory, starting in the lower left corner (with zero information), moves up with every application of SPE, and to the right with every application of SPD, progressing towards the upper right corner (representing perfect pattern recovery). The trajectory is (approximately) bounded above by the SPE EXIT curve, and (approximately) bounded to the right by the SPD EXIT curve, and thus cannot progress beyond the point at which those curves intersect. In this example, the SPE/SPD together do not allow perfect pattern recovery, though they come close.

VII. NUMERICAL EXPERIMENTS

Numerical experiments were conducted for the observation model (1), where the elements of $\mathbf{A}$ were independently drawn from a $\mathcal{CN}(0, M^{-1})$ distribution and where the signal coefficients were generated via $p(x_n | s_n) = s_n \mathcal{CN}(x_n; 0, 1) + (1 - s_n) \delta(x_n)$ using a Markov chain (MC)-generated binary sparsity pattern $\{s_n\}$. Such a MC is fully described by the transition probabilities $p_{01} \triangleq \Pr\{s_n = 0 | s_{n-1} = 1\}$ and $p_{10} \triangleq \Pr\{s_n = 1 | s_{n-1} = 0\}$, yielding a stationary distribution with *activity rate* $\lambda \triangleq \Pr\{s_n = 1\} = (1 + p_{01}/p_{10})^{-1}$. To generate patterns with a desired $\lambda \in (0, 1]$, we set $p_{01} = p_{10}(\lambda^{-1} - 1)$ with $p_{10} = \gamma\lambda$, for $\gamma \in (0, 1]$ called the *Markov independence parameter*. Note that, as γ increases, the pattern becomes less correlated, with $\gamma = 1$ corresponding to an i.i.d pattern. In fact, the pattern’s coherence time T is inversely proportional to γ , i.e., $T = (2\gamma\lambda(1 - \lambda))^{-1}$. To ensure that we used only “typical” pattern realizations, we

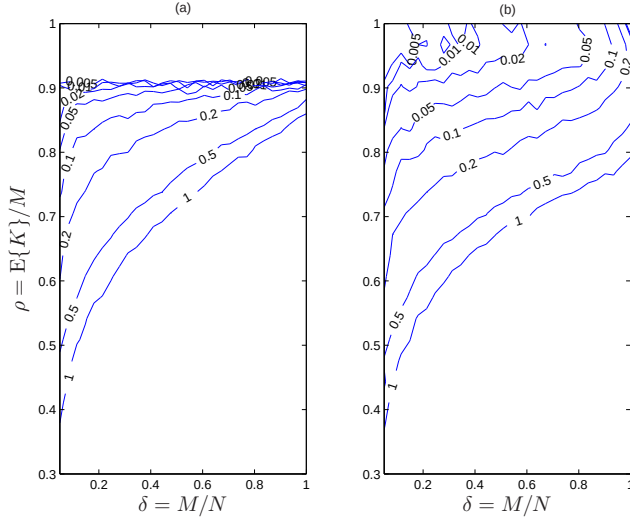


Fig. 6. Empirical 50% phase transition curves for the events (a) $\text{NMSE} \leq -20\text{dB}$ and (b) $\text{NMSE} \leq \text{NMSE}_{\text{genie}} + 1\text{dB}$, for various values of the Markov independence parameter γ (shown by the labels on each curve).

discarded any $\{s_n\}$ for which $\hat{\lambda} \triangleq \frac{1}{N} \sum_{n=1}^N s_n$ differed from λ by more than 5%. All experiments used $N = 512$ and σ^2 such that $\text{SNR} = 30\text{dB}$, where $\text{SNR} \triangleq \mathbb{E}\{\|\mathbf{A}\mathbf{x}\|_2^2\} / \mathbb{E}\{\|\mathbf{w}\|_2^2\}$.

SPD was implemented using the standard forward-backward (or sum-product) algorithm [8], [13], which performs exact inference for our MC pattern, and SPE was implemented using the AMP technique outlined in Section IV-B, allowing 10 SPE iterations. The proposed turbo scheme was allowed 5 iterations. We measured performance using $\text{NMSE} \triangleq \mathbb{E}\{\|\hat{\mathbf{x}} - \mathbf{x}_0\|_2^2\} / \mathbb{E}\{\|\mathbf{x}_0\|_2^2\}$, where $\mathbf{x}_0$ denotes the true parameter vector and $\hat{\mathbf{x}}$ denotes the final turbo estimate. Average performance was tested under different problem settings, parameterized by γ , the undersampling ratio $\delta \triangleq M/N$, and the normalized sparsity $\rho \triangleq \mathbb{E}\{K\}/M = \lambda N/M$, where K denotes the number of non-zero coefficients in $\{x_k\}$ (and $\{s_n\}$).

In Fig. 6(a), for various choices of Markov independence parameter γ , we show empirically estimated phase transition curves [10] for the event $\text{NMSE} \leq -20\text{dB}$, labeling each curve with the corresponding value of γ . The curves were constructed as follows. Defining $\text{NMSE} \leq -20\text{dB}$ as “success,” we empirically estimated the probability of success at each triplet (γ, δ, ρ) by comparing the NMSE for 200 independent realizations of $(\mathbf{A}, \mathbf{x}, \mathbf{s}, \mathbf{w})$. Matlab’s `contour` command was then used to draw the 50% success-probability contour, which always manifested as a curve bisecting the (δ, ρ) space: points northwest of the curve had $< 50\%$ success, and points southeast of the curve had $> 50\%$ success. Figure 6(a) shows that, as γ decreases (i.e., the sparsity becomes more structured), the successful- (δ, ρ) domain expands. This domain is bounded above by $\rho \approx 0.91$, however, since above that line the NMSE calculated under genie-aided perfect sparsity-pattern knowledge itself exceeded -20dB . Thus, in Fig. 6(b), we show 50% phase transition curves for the event $\text{NMSE} \leq \text{NMSE}_{\text{genie}} + 1\text{dB}$. There, the successful- (δ, ρ) domain appears to expand with every decrease in γ , i.e., increase in structure.

VIII. CONCLUSION

In this paper, we considered the reconstruction of structured-sparse signals from noisy linear observations, using probabilistic priors to model sparsity structure. Exact inference was examined and found to be computationally impractical, and so an approximate inference scheme, based on belief propagation, was proposed that iterates between exploitation of observation structure (via “SPE”) and sparsity pattern structure (via “SPD”), with connections to noncoherent turbo equalization. Using the AMP framework recently proposed by Donoho, Maleki, and Montanari, an efficient implementation of SPE was outlined, and EXIT charts were proposed as a means of predicting the interaction between SPE and SPD. Finally, performance was quantified using NMSE-based empirical phase transition curves, demonstrating that sparsity pattern structure can be successfully exploited by our turbo scheme.

REFERENCES

- [1] R. G. Baraniuk, V. Cevher, M. F. Duarte, and C. Hegde, “Model-based compressive sensing,” *IEEE Trans. Inform. Theory*, vol. 56, pp. 1982–2001, Apr. 2010.
- [2] C. A. Bouman, “Markov random fields and stochastic image models,” in *IEEE Int. Conf. Image Processing Tutorial*, Oct. 1995.
- [3] P. J. Wolfe, S. J. Godsill, and W.-J. Ng, “Bayesian variable selection and regularization for time-frequency surface estimation,” *J. R. Statist. Soc. B*, vol. 66, no. 3, pp. 575–589, 2004.
- [4] L. He and L. Carin, “Exploiting structure in wavelet-based Bayesian compressive sensing,” *IEEE Trans. Signal Process.*, vol. 57, pp. 3488–3497, Sept. 2009.
- [5] C. P. Robert and G. Casella, *Monte Carlo Statistical Methods*. New York: Springer, 2nd ed., 2004.
- [6] V. Cevher, M. F. Duarte, C. Hedge, and R. G. Baraniuk, “Sparse signal recovery using Markov random fields,” in *Proc. Neural Inform. Process. Syst. Conf.*, (Vancouver, B.C.), Dec. 2008.
- [7] M. F. Duarte, M. B. Wakin, and R. G. Baraniuk, “Wavelet-domain compressive signal reconstruction using a hidden Markov tree model,” in *Proc. IEEE Int. Conf. Acoust. Speech & Signal Process.*, (Las Vegas, NV), pp. 5137–5140, Apr. 2008.
- [8] J. Pearl, *Probabilistic Reasoning in Intelligent Systems*. San Mateo, CA: Morgan Kaufman, 1988.
- [9] D. Baron, S. Sarvotham, and R. G. Baraniuk, “Bayesian compressive sensing via belief propagation,” *IEEE Trans. Signal Process.*, vol. 58, pp. 2269–280, Jan. 2010.
- [10] D. L. Donoho, A. Maleki, and A. Montanari, “Message passing algorithms for compressed sensing,” *Proc. National Academy of Sciences*, vol. 106, pp. 18914–18919, Nov. 2009.
- [11] R. Koetter, A. C. Singer, and M. Tüchler, “Turbo equalization,” *IEEE Signal Process. Mag.*, vol. 21, pp. 67–80, Jan. 2004.
- [12] G. F. Cooper, “The computational complexity of probabilistic inference using Bayesian belief networks,” *Artificial Intelligence*, vol. 42, 1990.
- [13] D. J. C. MacKay, *Information Theory, Inference, and Learning Algorithms*. New York: Cambridge University Press, 2003.
- [14] W. T. Freeman, E. C. Pasztor, and O. T. Carmichael, “Learning low-level vision,” *Intl. J. Computer Vision*, vol. 40, pp. 25–47, Oct. 2000.
- [15] D. Guo, D. Baron, and S. Shamai, “A single-letter characterization of optimal noisy compressed sensing,” in *Proc. Allerton Conf. Commun. Control Comput.*, (Monticello, IL), Oct. 2009.
- [16] S.-J. Hwang and P. Schniter, “Efficient multicarrier communication for highly spread underwater acoustic channels,” *IEEE J. Sel. Areas Commun.*, vol. 26, pp. 1674–1683, Dec. 2008.
- [17] P. Schniter, L. C. Potter, and J. Ziniel, “Fast Bayesian matching pursuit,” in *Proc. Inform. Theory & Appl. Workshop*, Feb. 2008.
- [18] D. L. Donoho, A. Maleki, and A. Montanari, “Message passing algorithms for compressed sensing: I. Motivation and construction,” in *Proc. Inform. Theory Workshop*, (Cairo, Egypt), Jan. 2010.
- [19] S. ten Brink, “Convergence behavior of iteratively decoded parallel concatenated codes,” *IEEE Trans. Commun.*, vol. 49, pp. 1727–1737, Oct. 2001.